

CS 784: Computational Linguistics

Lecture 16: Grounded Semantics

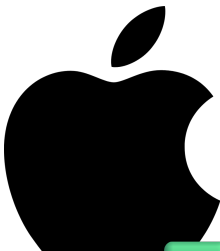
Freda Shi

School of Computer Science, University of Waterloo
fhs@uwaterloo.ca

March 13, 2025

Meanings in the Real World

My favorite fruit is **apple**.



Experience Grounds Language (Bisk et al., 2020)

*We posit that the present success of representation learning approaches trained on large, text-only corpora requires the parallel tradition of research on the broader **physical and social context** of language to address the deeper questions of communication.*

[Bisk, Y. et al., 2020. *Experience Grounds Language*. In *EMNLP*.]

What is Grounding?

Given the primary data source $\mathcal{X}$ and the ground $\mathcal{Y}$, **grounding** is the process of establishing a meaningful relationship between them.

It is implied that the mutual information $I(\mathcal{X}; \mathcal{Y}) > 0$.

Usually, we also have $H(\mathcal{Y} \mid \mathcal{X}) > 0$ —the ground is not specifically determined by the data source.

Examples

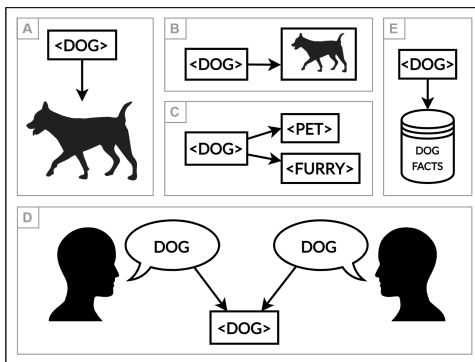
- $\mathcal{X}$: text, $\mathcal{Y}$: image (represent meaning of text with image)
- $\mathcal{X}$: utterance from person A , $\mathcal{Y}$: mental state of person B (understanding the communication intention)
- $\mathcal{X}$: text, $\mathcal{Y}$: audio (connecting text with corresponding audio)
- $\mathcal{X}$: image, $\mathcal{Y}$: text (image understanding with textual supervision)

[Chai, J. Y. et al. 2018. *Language to action: Towards interactive task learning with physical agents*. In *IJCAI*.]

[Shi, H. F. 2024. *Learning language structures through grounding*. Ph.D. Thesis. Toyota Technological Institute at Chicago.]

Grounding: A Comprehensive Taxonomy

Grounding can be categorized into A. referential grounding, B. sensorimotor grounding, C. relational grounding, D. communicative grounding, and E. epistemic grounding.



[Figure credit: Mollo and Millière. The Vector Grounding Problem.

<https://arxiv.org/pdf/2304.01481>]

Grounding: This Lecture

Grounding is a broad topic that goes beyond semantics—communicative grounding is the key of pragmatics.

However, in this lecture, we focus on **semantic grounding** (or more specifically, sensorimotor grounding): representing meanings of text with data from other modalities (e.g., images).

Recap: (Ungrounded) Pure-Text Language Models

Two popular types of (ungrounded) pure-text language models:

- Autoregressive models (e.g., GPT):

$$P_{\Theta}(w_i \mid w_1, \dots, w_{i-1})$$

- Masked language models (e.g., BERT):

$$P_{\Theta}(w_i \mid w_1, \dots, w_{i-1}, w_{i+1}, \dots, w_n)$$

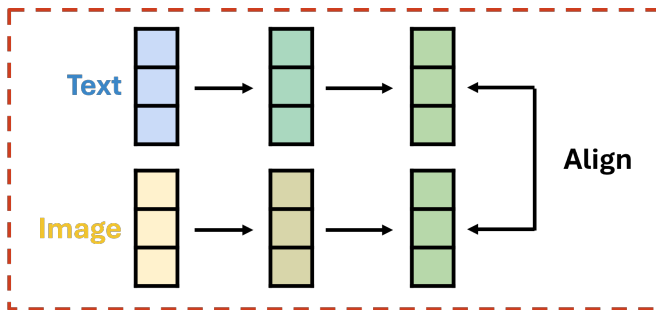
Whether these pure-text language models encode meaning, and to what extent, is still under debate.

[Radford, A. et al. 2018. *Improving language understanding by generative pretraining.*]

[Devlin, J. et al. 2019. *BERT: Pre-training of deep bidirectional transformers for language understanding.* In *NAACL*.]

Joint Visual-Semantic Embedding Space

Encode visual and textual information into a **shared space**.



Learning Joint Visual Semantic Space

Training data: pairs of images and text descriptions.

Core idea: encode images and text into a **joint embedding space** by minimizing the hinge-based triplet loss.

$$\Theta^* = \arg \min_{\Theta} \sum_{(I^+, T^+, T^-)} \max(0, \alpha - \text{sim}(I_{\Theta}^+, T_{\Theta}^+) + \text{sim}(I_{\Theta}^+, T_{\Theta}^-)) \\ + \sum_{(T^+, I^+, I^-)} \max(0, \alpha - \text{sim}(T_{\Theta}^+, I_{\Theta}^+) + \text{sim}(T_{\Theta}^+, I_{\Theta}^-))$$

I^+ :



T^+ : *There is a cat standing on the lawn.*

I^- :



T^- : *There is an apple on the table.*

[Kiros, R. et al. 2014. *Unifying visual-semantic embeddings with multimodal neural language models.*]

Properties of the Joint Space

Images and text descriptions are close in the joint space if they are semantically related.

Example Applications:

- Bidirectional image-caption retrieval: encode the query (image or text), and the “database” into the joint space and retrieve the closest neighbors.
- Image captioning: encode the image into the joint space, and train a decoder to generate text conditioned on the image encoding.

Text in the training corpus can be at any level of granularity (e.g., word, phrase, sentence, paragraph).

Variations of Training Objective: Hard Negative Mining

Original:

$$\Theta^* = \arg \min_{\Theta} \sum_{(I^+, T^+, T^-)} [\alpha - \text{sim}(I_{\Theta}^+, T_{\Theta}^+ + \text{sim}(I_{\Theta}^+, T_{\Theta}^-))]_{+} \\ + \sum_{(T^+, I^+, I^-)} [\alpha - \text{sim}(T_{\Theta}^+, I_{\Theta}^+) + \text{sim}(T_{\Theta}^+, I_{\Theta}^-)]_{+}$$

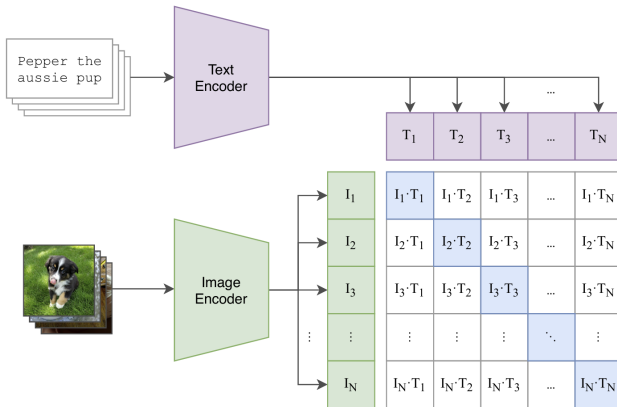
$$[\cdot]_{+} = \max(0, \cdot)$$

Modified:

$$\Theta^* = \arg \min_{\Theta} \sum_{(I^+, T^+)} \max_{T^-} [\alpha - \text{sim}(I_{\Theta}^+, T_{\Theta}^+) + \text{sim}(I_{\Theta}^+, T_{\Theta}^-)]_{+} \\ + \sum_{(T^+, I^+)} \max_{I^-} [\alpha - \text{sim}(T_{\Theta}^+, I_{\Theta}^+) + \text{sim}(T_{\Theta}^+, I_{\Theta}^-)]_{+}$$

[Faghri, F. et al. 2017. *VSE++: Improving visual-semantic embeddings with hard negatives*. In *BMVC*.]

Variations of Training Objective: Contrastive Learning



[Radford, A. et al. 2021. *Learning transferable visual models from natural language supervision.*]

Variations of Training Objective: Contrastive Learning

$$\Theta^* = \arg \min_{\Theta} \mathbb{E}_{[(I_1, T_1), \dots, (I_n, T_n)]} \left[\sum_i -\log P_{\Theta}(T_i \mid I_i [T_{1\dots n}]) - \log P_{\Theta}(I_i \mid T_i [I_{1\dots n}]) \right]$$

$[(I_1, T_1), \dots, (I_n, T_n)]$: a batch of image-text pairs.

There exists a probabilistic interpretation of the training loss.

- **Query**: image I_i .
- **Database**: text descriptions $T_1, \dots, T_n$.
- **Ground truth**: T_i .

$P_{\Theta}(T_i \mid I_i [T_{1\dots n}])$: the probability of T_i being the correct retrieval result in the above settings.

Variations of Training Objective: Contrastive Learning

The softmax function converts a list of real values (e.g., $\mathbf{x} \in \mathbb{R}^n$) to a probability distribution.

$$\text{softmax}(\mathbf{x})_i = \frac{\exp(x_i)}{\sum_{j=1}^n \exp(x_j)}$$

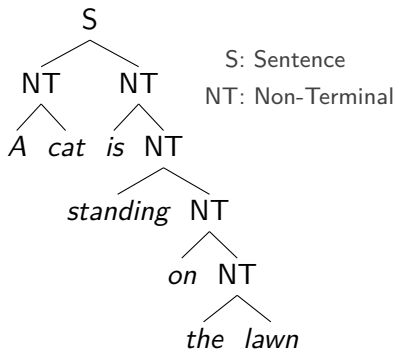
$$P_{\Theta}(T_i \mid I_i, [T_{1\dots n}]) = \frac{\exp(\langle I_i, T_i \rangle)}{\sum_{j=1}^n \exp(\langle I_i, T_j \rangle)} = [\text{softmax}(I_i^T [T_{1\dots n}])]_i$$

Visually Grounded Grammar Induction

Input: Captioned images.

Output: Linguistically plausible structure for captions.

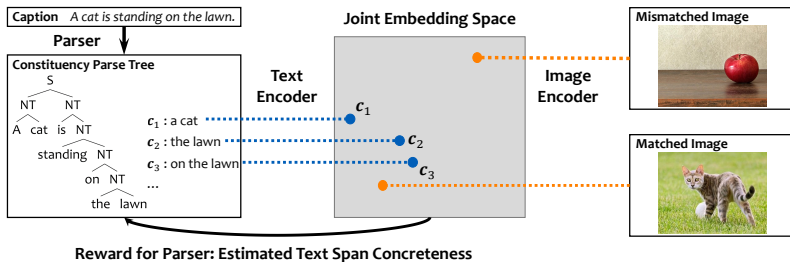
A cat is standing on the lawn.



[Source: Shi et al. 2019. Visually Grounded Neural Syntax Acquisition. In ACL.]

Grounded Signals for Syntax Acquisition

Hypothesis: more visually concrete word spans are more likely to be constituents.



Concreteness Estimation with the Joint Embedding Space

Image i



Candidate
Constituent c

a cat
on the

Another Image i'



$$\ell(c; i, i') = \text{sim}(i', c) - \text{sim}(i, c)$$

Value of ℓ

$$\text{sim}\left(\text{img of apple}, a \text{ cat}\right) = 0.2 \quad \text{sim}\left(\text{img of cat}, a \text{ cat}\right) = 0.9 \quad \ell = -0.7$$

$$\text{sim}\left(\text{img of apple}, on \text{ the}\right) = 0.4 \quad \text{sim}\left(\text{img of cat}, on \text{ the}\right) = 0.4 \quad \ell = 0$$

Key Idea: Smaller $\ell(c)$ $\iff$ c is more visually concrete.

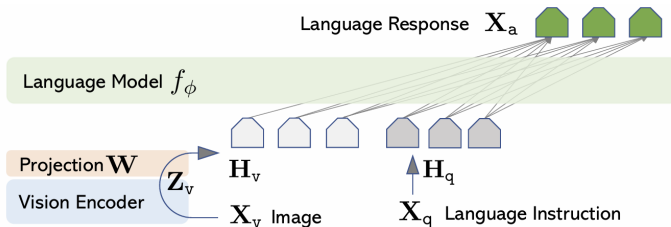
Quantify *visual concreteness* of word spans using loss values.

LLaVA: Visual Instruction Tuning

Use GPT-style language modeling objective.

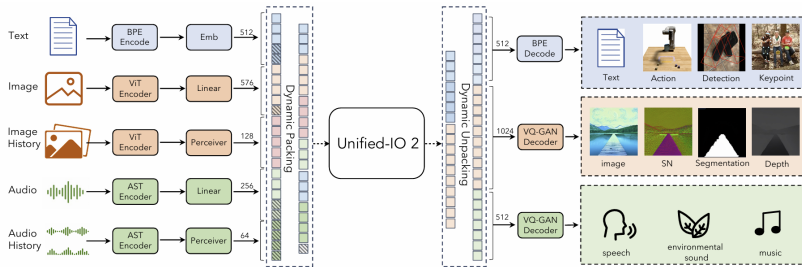
Encode images with different resolutions into “visual tokens.”

Project the visual tokens into the textual (joint) space.



[Liu, H. et al. 2023. *Visual instruction tuning*. In *NeurIPS*.]

Towards Encoding Everything in the World



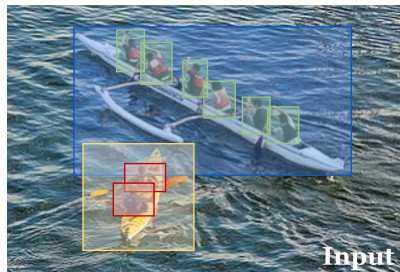
[Lu, J. et al. 2024. Unified-IO 2: Scaling autoregressive multimodal models with vision language audio and action. In *CVPR*]

Finer-Grained Vision-Language Tasks

Object retrieval

Note: the object bounding boxes are given in both training and testing.
This is a reasonable assumption, as cognitive scientists have shown that 5-month infants recognize objects well.

a smaller yellow boat



a smaller yellow boat



[Baillargeon, R. et al. 1985. *Object permanence in five-month-old infants*. In *Cognition*.]

Finer-Grained Vision-Language Tasks

Multimodal coreference resolution

a smaller yellow boat



a smaller yellow boat



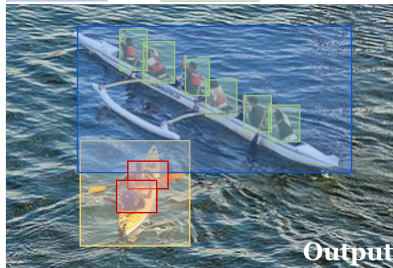
Finer-Grained Vision-Language Tasks

Phrase grounding

Two boats of people kayaking, a smaller yellow boat with two people and a larger white boat with six people.

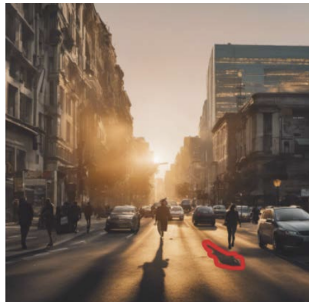


Two boats of people kayaking, a smaller yellow boat with two people and a larger white boat with six people.



Limitation of Current Vision-Language Models

- Lack of full understanding of the physical world.



[Sarkar, A. et al. 2024. *Shadows don't lie and lines can't bend! Generative models don't know projective geometry...for now.* In *CVPR*.]

Limitation of Current Vision-Language Models

- Poor in recognizing spatial relations, especially poor adapting different spatial frames of reference.

Is the basketball to the right of the car?

- Yes, from the camera's viewpoint
- Yes, from the woman's viewpoint
- Yes, from the car's viewpoint



[Zhang Z. et al. 2024. Do vision-language models represent space and how?
Evaluating spatial frame of reference under ambiguities.]

Limitation of Current Vision-Language Models

- Highly biased towards cultures with more presence in the training data.

wedding						
CLIP	100.00					
	31.09					
CoCA	82.77					
	59.04					
BLIP-2	82.77					
	31.09					

[Bhatia, M. et al. 2024. *From local concepts to universals: Evaluating the multicultural understanding of vision-language models.*]

Next

Pragmatics