



Language
Technologies
Institute

Carnegie
Mellon
University

Multimodal Machine Learning

Lecture 4.2: Multimodal Alignment

Louis-Philippe Morency

** Original course co-developed with Tadas Baltrusaitis.
Major revision co-developed with Paul Liang. Co-lecturers
included Yonatan Bisk, Daniel Fried and Carlos Busso.*

Administrative Stuff

Upcoming Course Assignments

Proposal, aka First Assignment (due Sunday 9/20 at 8pm ET)

- Literature review + initial research ideas

Week 5 Reading Assignment (due Friday 9/11 at 8pm ET)

- Two papers (with a focus on multimodal alignment)

Proposal presentations (next week)

- Be sure to signup for a 15-mins timeslot!
- Designed to give feedback about your research ideas

Second project assignment (due Sunday 10/4 at 8pm ET)

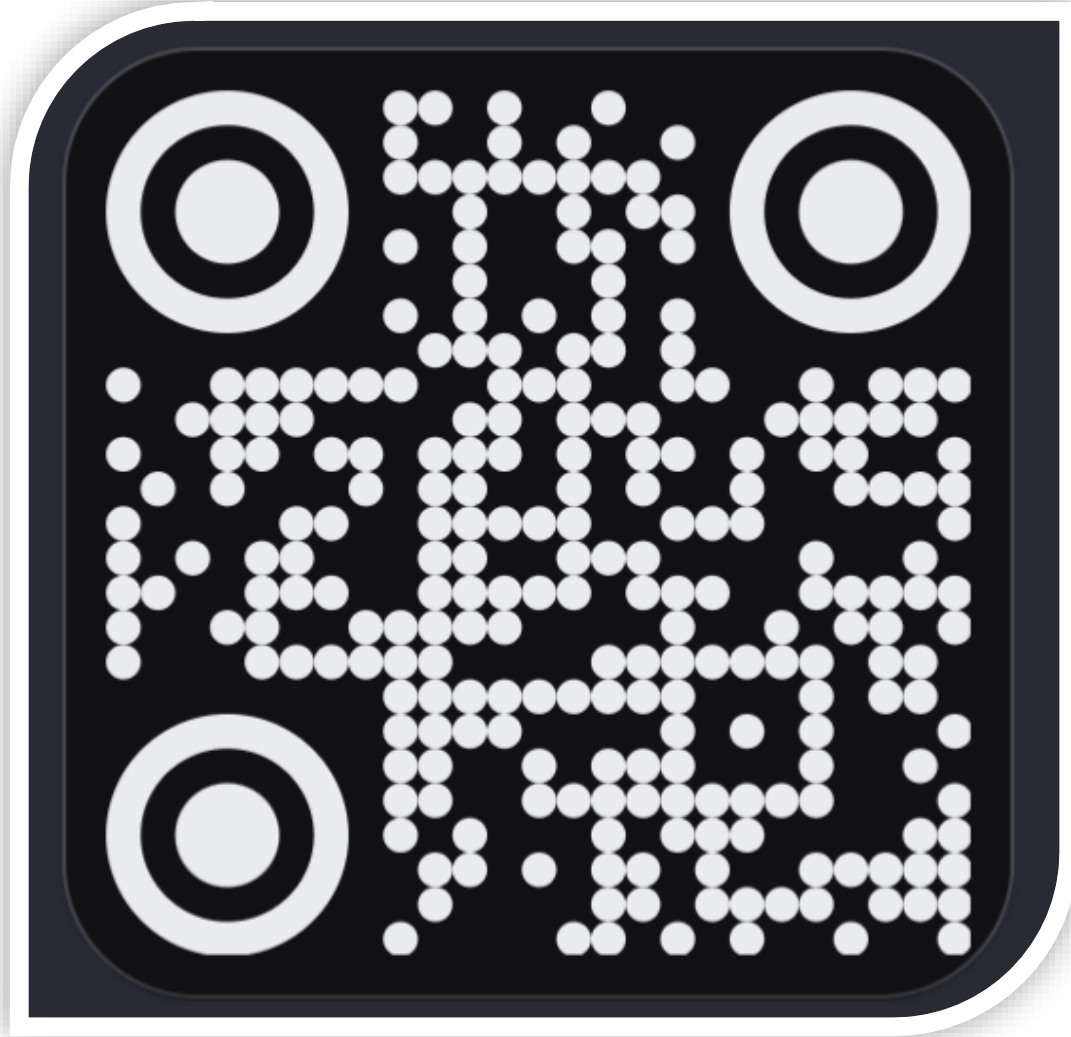
- Look at your data, before modeling it! 😊

No lectures next week!

No lectures on Tuesday 9/22 and Thursday 9/24



Attendance using Poll Everywhere



- 1 Scan QR Code
or
<https://pe.app/mmm1>
- 2 Enter your AndrewID
(as your “name tag”)
-  *Do not try to sign-in*
- 3 Click “Share location”



Language
Technologies
Institute

Carnegie
Mellon
University

Multimodal Machine Learning

Lecture 4.2: Multimodal Alignment

Louis-Philippe Morency

** Original course co-developed with Tadas Baltrusaitis.
Major revision co-developed with Paul Liang. Co-lecturers
included Yonatan Bisk, Daniel Fried and Carlos Busso.*

Lecture objectives

- Discrete alignment
 - Global alignment
 - Assignment problem and optimal transport
- Continuous alignment
 - Continuous warping
 - Dynamic time warping
 - Discretization and segmentation
- Multimodal Alignment + Representation
 - Cross-modal and concatenated architectures

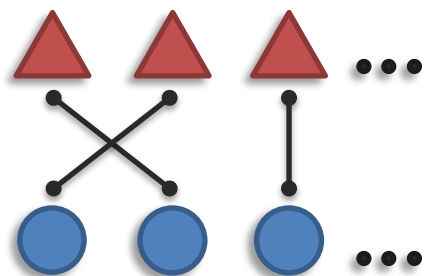
Discrete Alignment

Challenge 2: Alignment

Definition: Identifying and modeling cross-modal connections between all elements of multiple modalities, building from the data structure

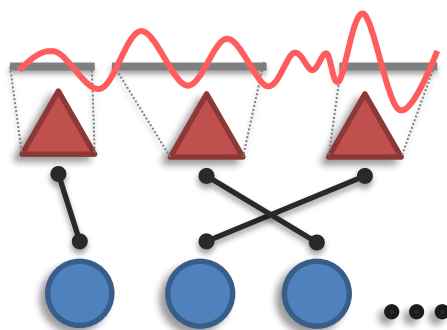
Sub-challenges:

Discrete Alignment



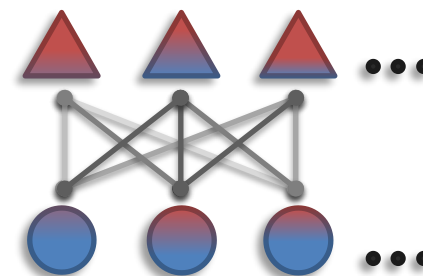
Discrete elements
and connections

Continuous Alignment



Segmentation and
continuous warping

Contextualized Representation



Alignment + representation

Global Alignment

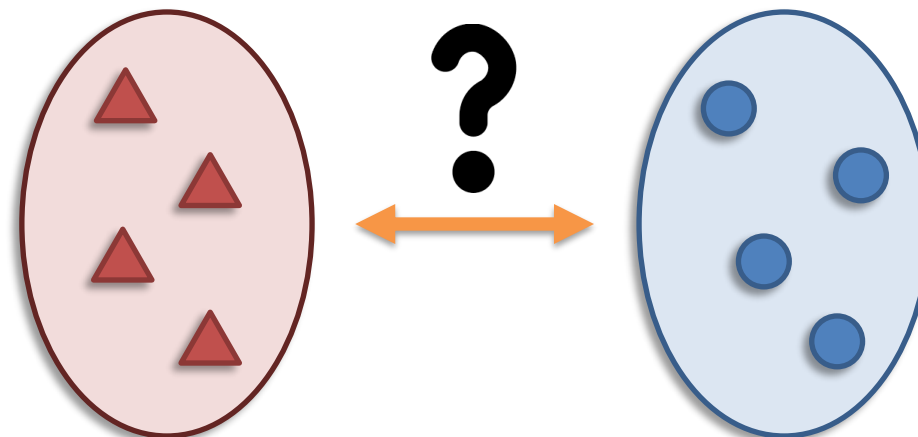
Visual



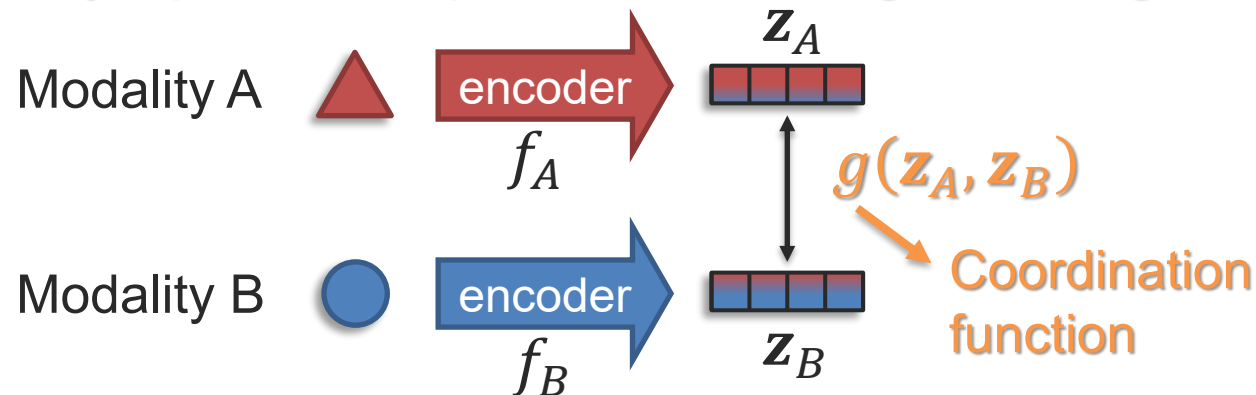
Language



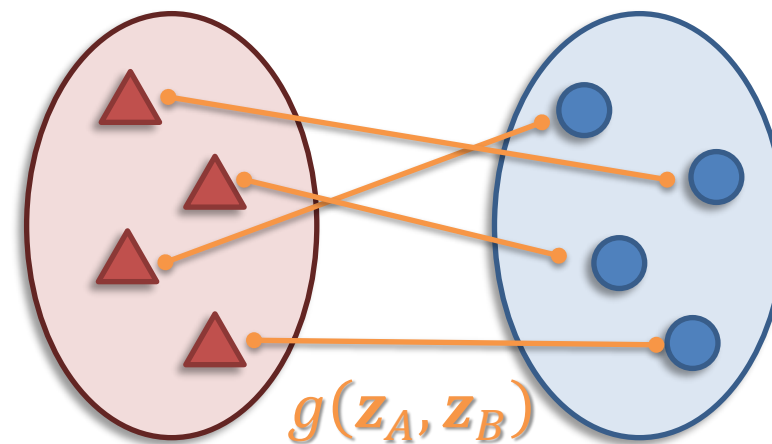
Latent pairing information



Jointly optimize representation + global alignment:

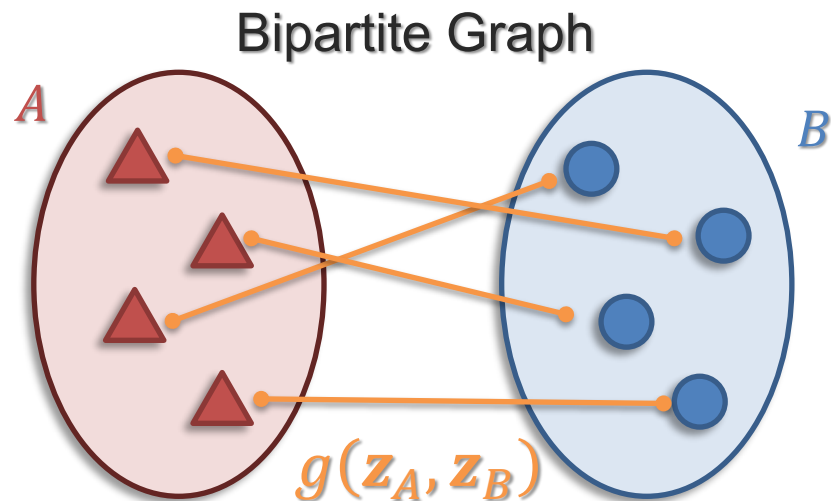


(representation)



(global alignment)

Assignment Problem



Initial assumptions:

- Same number of elements in A and B modalities
- 1-to-1 “hard” alignment between elements
- All elements assigned (aka “perfect matching”)

➡ How to solve?

Naive solution: check all assignments

Better solution: Linear Programming

Assignment: ~~$f: A \rightarrow B$~~
(vector of indices)

$x_{ij} = 1$ when matching connection, otherwise 0

Similarity weights: ~~$w_{(i,f(i))} = g(\mathbf{z}_A^i, \mathbf{z}_B^{f(i)})$~~

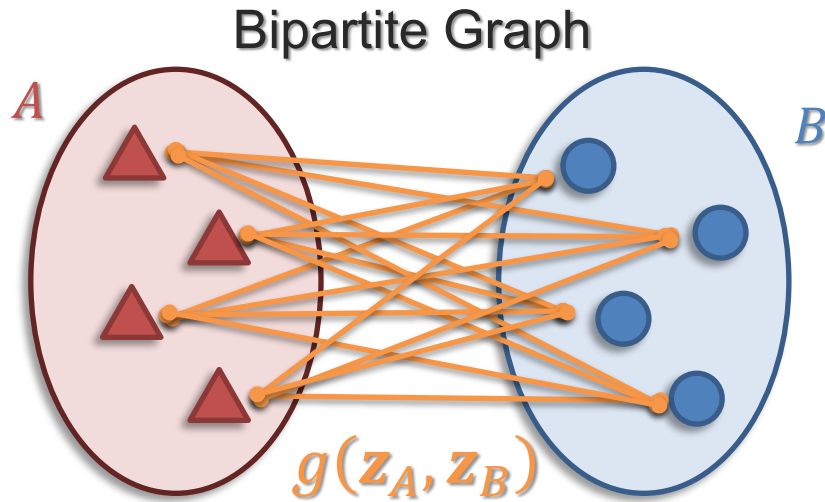
$w_{(i,j)} = g(\mathbf{z}_A^i, \mathbf{z}_B^j)$

Maximize: ~~$\max_{f \in \text{Perm}(N)} \sum_{i=1}^N w_{i,f(i)}$~~

$\max_{\{x_{ij}\}} \sum_{(i,j) \in A \times B} w_{i,j} x_{ij}$

➡ Can be solved with
simplex algorithm

Optimal transport



New assumptions:

- Different number of elements in A and B modalities
- Many-to-many “soft” alignment between elements

➡ It can be seen as “transporting” elements from modality A to modality B (and vice-versa)

Assignments: $x_{(i,j)}$: soft alignment between $\mathbf{z}_A^i$ and $\mathbf{z}_B^j$

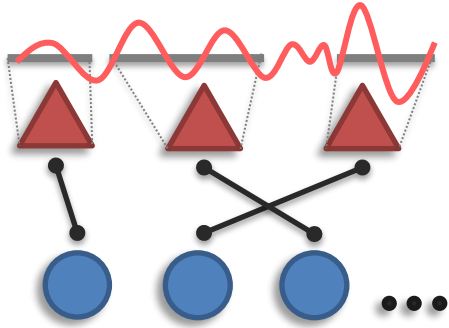
Similarity weights: $w_{(i,j)} = g(\mathbf{z}_A^i, \mathbf{z}_B^j)$

Maximize:
$$\max_{\{x_{ij}\}} \sum_{(i,j) \in A \times B} w_{i,j} x_{ij}$$

➡ Wassertein distance gives optimal transport

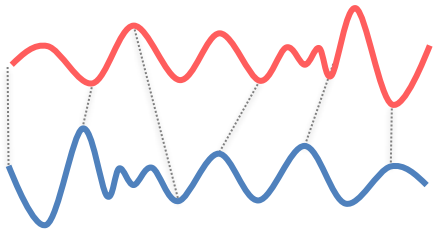
Continuous Alignment

Challenge 2b: Continuous Alignment

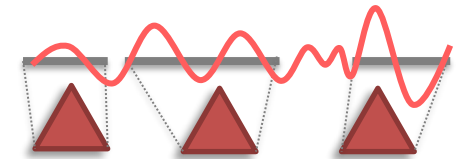


Definition: Model alignment between modalities with continuous signals and no explicit elements

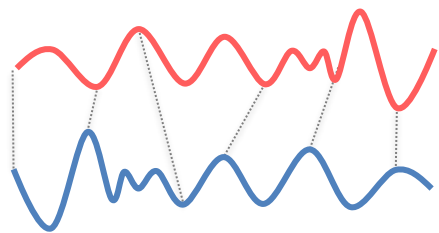
Continuous
warping



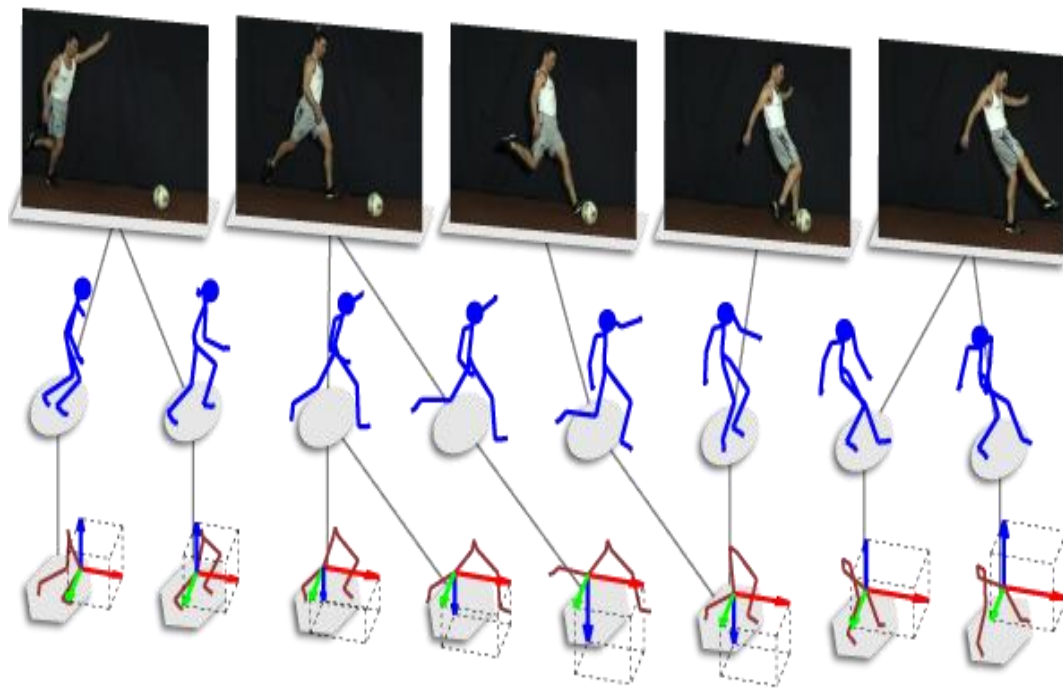
Discretization
(segmentation)



Continuous Warping – Example



➡ Aligning video sequences



Dynamic Time Warping (DTW)

We have two unaligned temporal unimodal signals

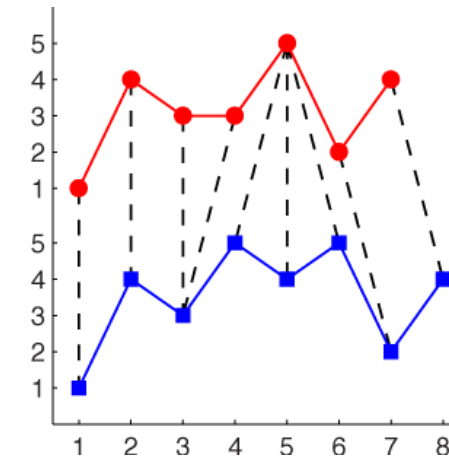
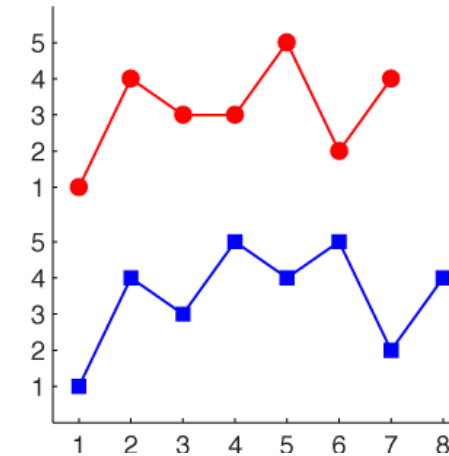
- $\mathbf{X} = [x_1, x_2, \dots, x_{n_x}] \in \mathbb{R}^{d \times n_x}$
- $\mathbf{Y} = [y_1, y_2, \dots, y_{n_y}] \in \mathbb{R}^{d \times n_y}$

Find set of indices to minimize the alignment difference:

$$L(\mathbf{p}^x, \mathbf{p}^y) = \sum_{t=1}^l \left\| x_{p_t^x} - y_{p_t^y} \right\|_2^2$$

where $\mathbf{p}^x$ and $\mathbf{p}^y$ are index vectors of same length

Dynamic Time Warping is designed to find these index vectors!



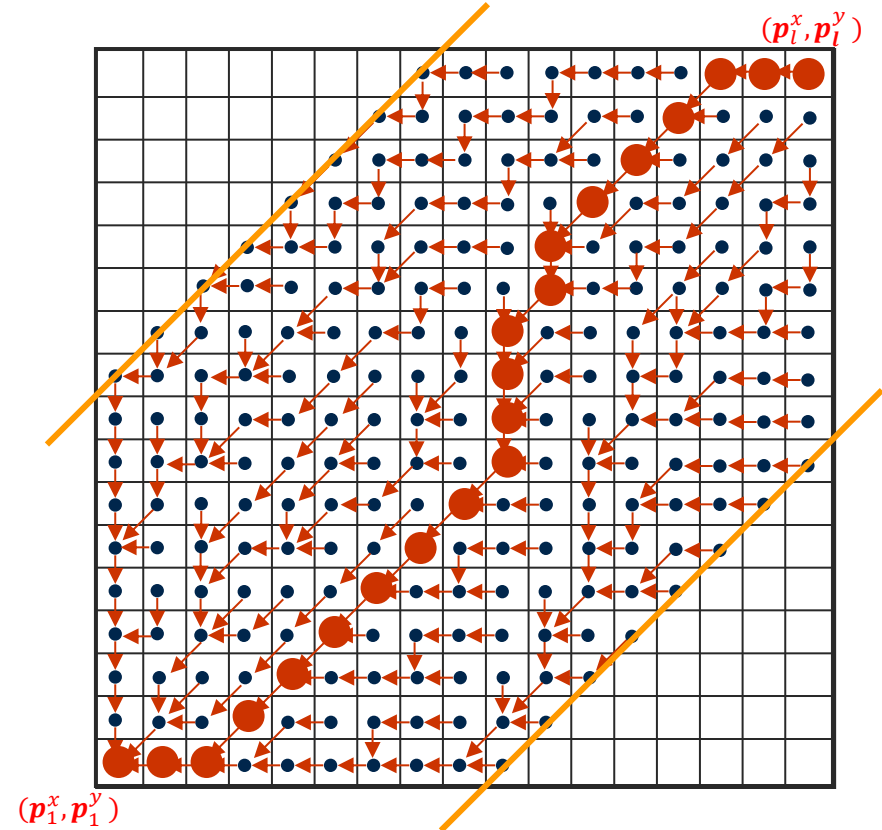
Dynamic Time Warping (DTW)

Lowest cost path in a cost matrix

- Restrictions?

- Monotonicity – no going back in time
- Continuity - no gaps
- Boundary conditions - start and end at the same points
- Warping window - don't get too far from diagonal
- Slope constraint – do not insert or skip too much

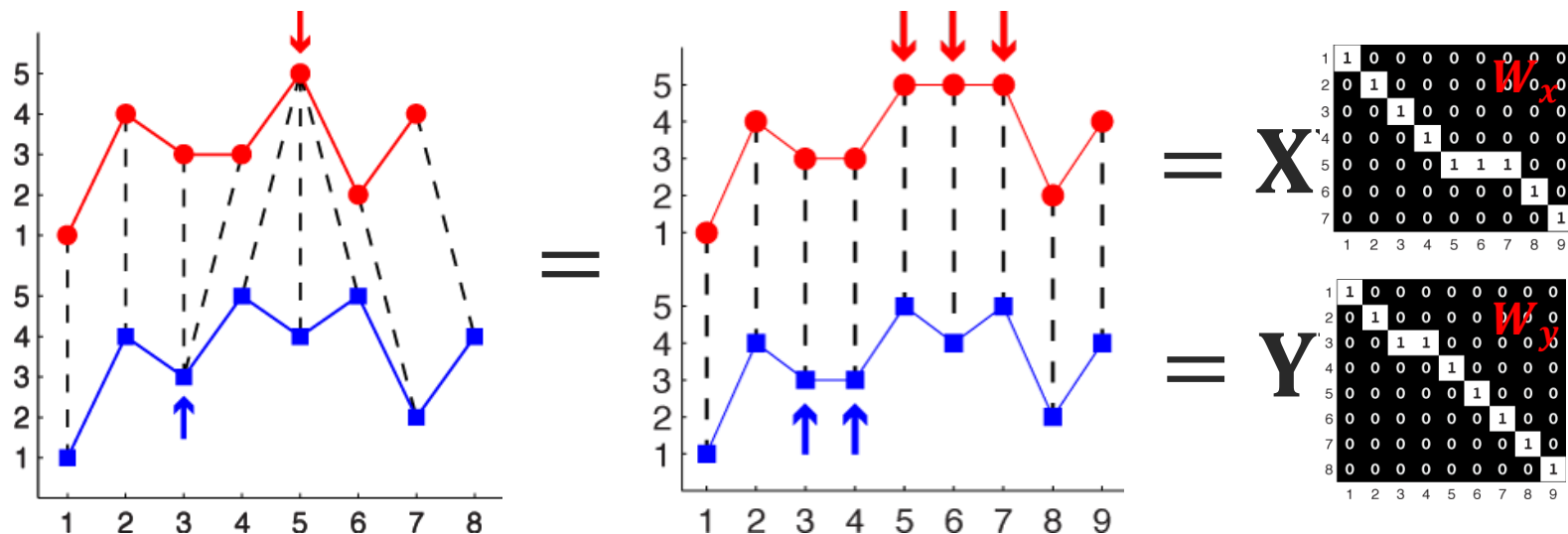
Solved using dynamic programming while respecting the restrictions



DTW alternative formulation

$$L(\mathbf{p}^x, \mathbf{p}^y) = \sum_{t=1}^l \left\| \mathbf{x}_{\mathbf{p}_t^x} - \mathbf{y}_{\mathbf{p}_t^y} \right\|_2^2$$

Replication doesn't change the objective!



Alternative objective:

$$L(\mathbf{W}_x, \mathbf{W}_y) = \left\| \mathbf{X}\mathbf{W}_x - \mathbf{Y}\mathbf{W}_y \right\|_F^2$$

Frobenius norm $\|\mathbf{A}\|_F^2 = \sum_i \sum_j |a_{i,j}|^2$

$\mathbf{X}, \mathbf{Y}$ – original signals (same #rows, possibly different #columns)

$\mathbf{W}_x, \mathbf{W}_y$ - alignment matrices

A differentiable version of DTW also exists...

<https://arxiv.org/pdf/1703.01541.pdf>

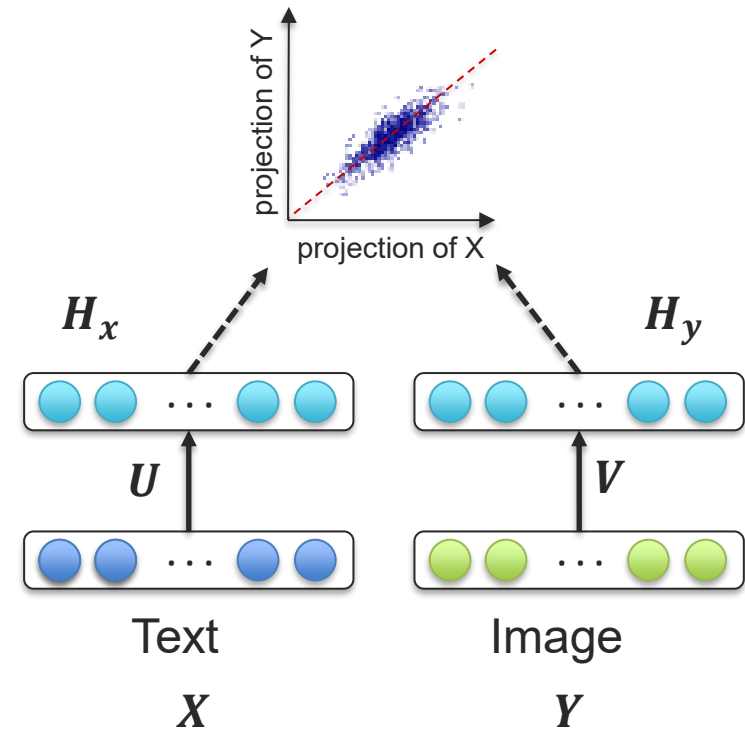
Canonical Correlation Analysis – Reminder

CCA loss can also be re-written as:

$$L(U, V) = \|\mathbf{U}^T \mathbf{X} - \mathbf{V}^T \mathbf{Y}\|_F^2$$

subject to:

$$\mathbf{U}^T \Sigma_{YY} \mathbf{U} = \mathbf{V}^T \Sigma_{YY} \mathbf{V} = \mathbf{I}, \mathbf{u}_{(j)}^T \Sigma_{XY} \mathbf{v}_{(i)} = 0$$



Canonical Time Warping

Dynamic Time Warping + Canonical Correlation Analysis = Canonical Time Warping

$$L(\mathbf{U}, \mathbf{V}, \mathbf{W}_x, \mathbf{W}_y) = \|\mathbf{U}^T \mathbf{X} \mathbf{W}_x - \mathbf{V}^T \mathbf{Y} \mathbf{W}_y\|_F^2$$

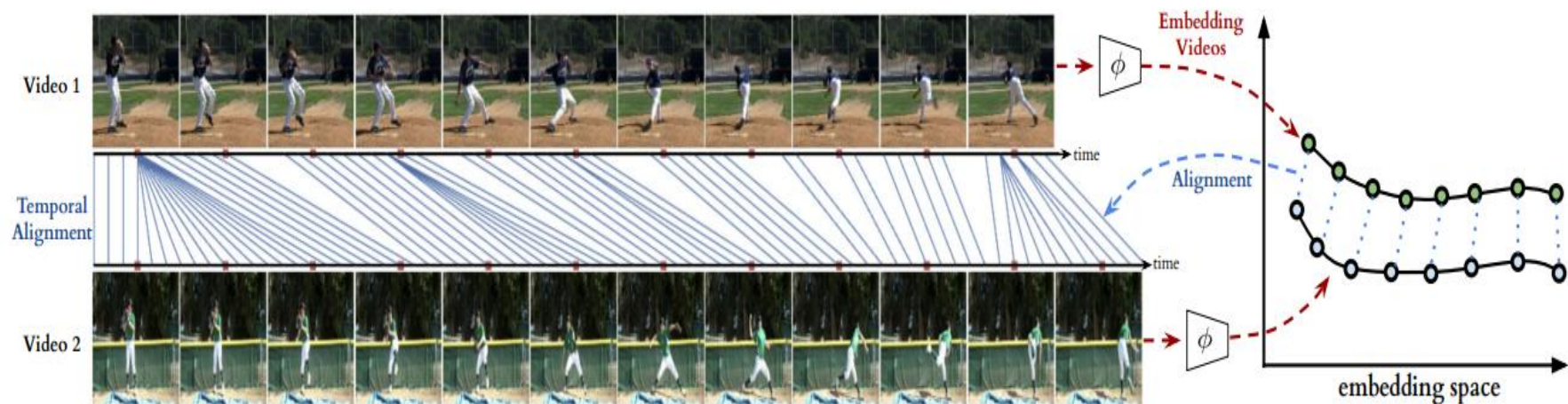
Allows to align multi-modal or multi-view (same modality but from a different point of view)

- $\mathbf{W}_x, \mathbf{W}_y$ – temporal alignment
- $\mathbf{U}, \mathbf{V}$ – cross-modal (spatial) alignment

[Canonical Time Warping for Alignment of Human Behavior, Zhou and De la Tore, 2009]

Temporal Alignment and Neural Representation Learning

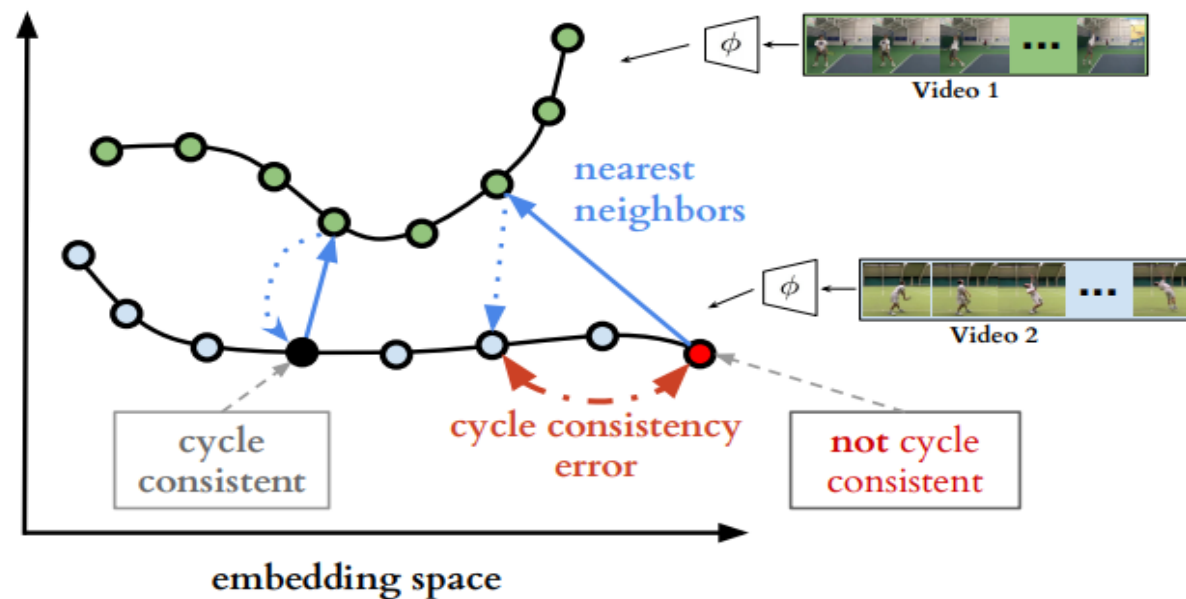
Premise: we have paired video sequences that can be temporally aligned



How can we define a loss function to enforce the alignment between sequences while at the same time learning good representations?

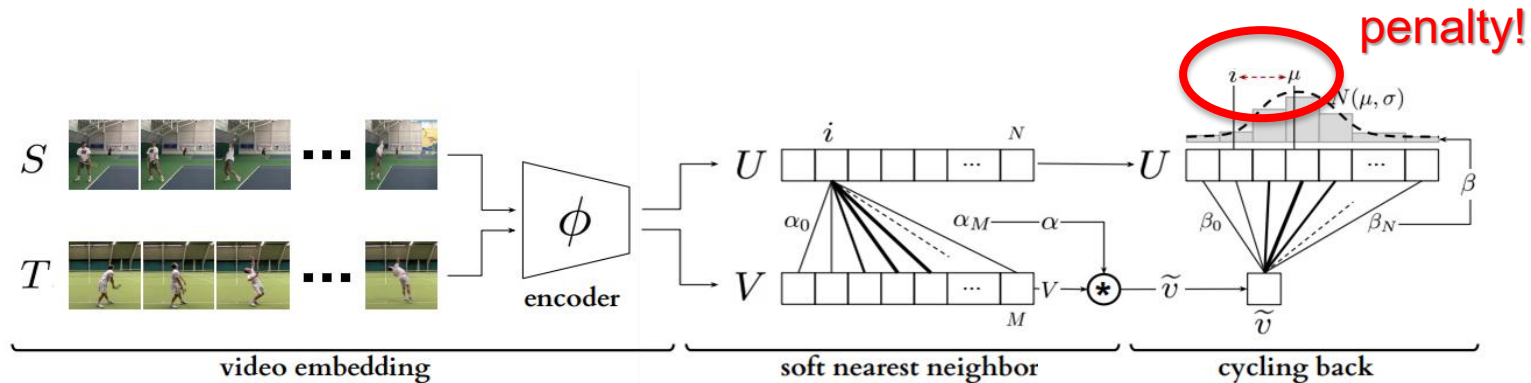
Temporal Cycle-Consistency Learning

Solution: Representation learning by enforcing **Cycle consistency**



Main idea: My closest neighbor also views me as their closest neighbor

Temporal Cycle-Consistency Learning



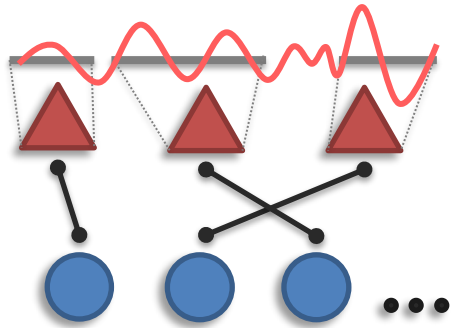
Compute “soft” / “weighted” nearest neighbour:

distances: $\alpha_j = \frac{e^{-||u_i - v_j||^2}}{\sum_k^M e^{-||u_i - v_k||^2}}$ Soft nearest neighbor: $\tilde{v} = \sum_j^M \alpha_j v_j,$

Find the nearest neighbor the other way and then penalize the distance:

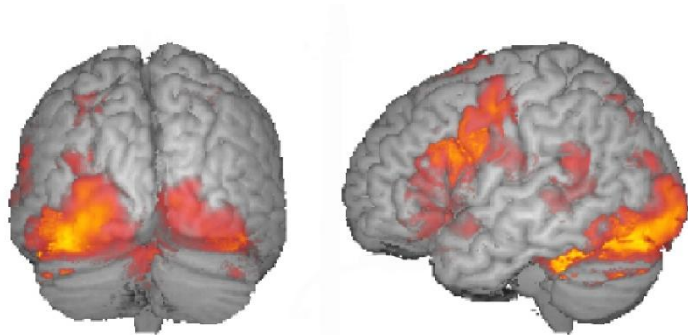
$$\beta_k = \frac{e^{-||\tilde{v} - u_k||^2}}{\sum_j^N e^{-||\tilde{v} - u_j||^2}} \quad L_{cbr} = \frac{|i - \mu|^2}{\sigma^2} + \lambda \log(\sigma)$$

Discretization (aka Segmentation)

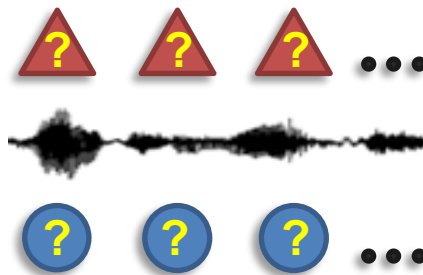


Common assumptions: ① Segmented elements

Examples:



Medical imaging



Signals

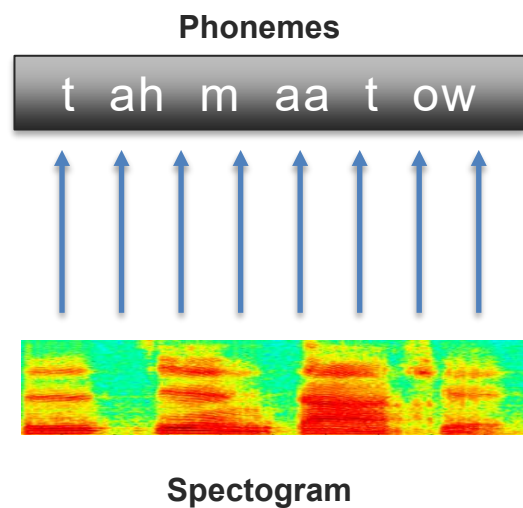


Images

objects

Discretization – Example

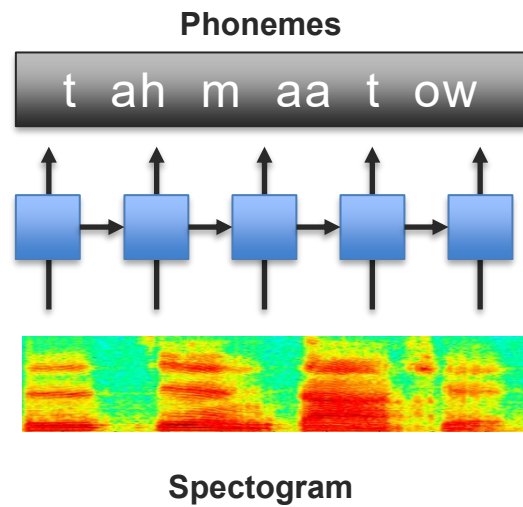
Sequence Labeling and Alignment



How can we predict the sequence
of phoneme labels?

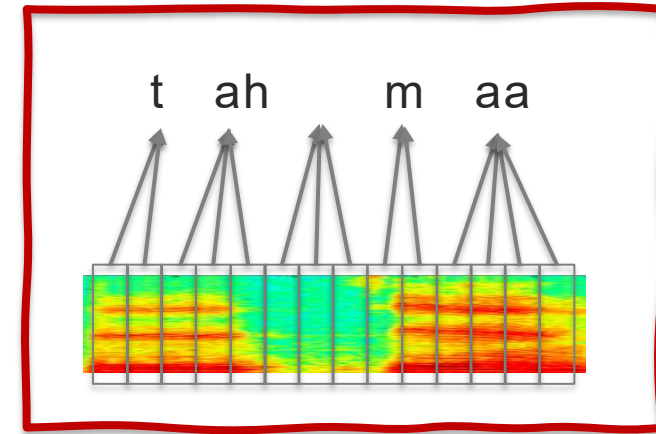
Discretization – Example

Sequence Labeling and Alignment



How can we predict the sequence of phoneme labels?

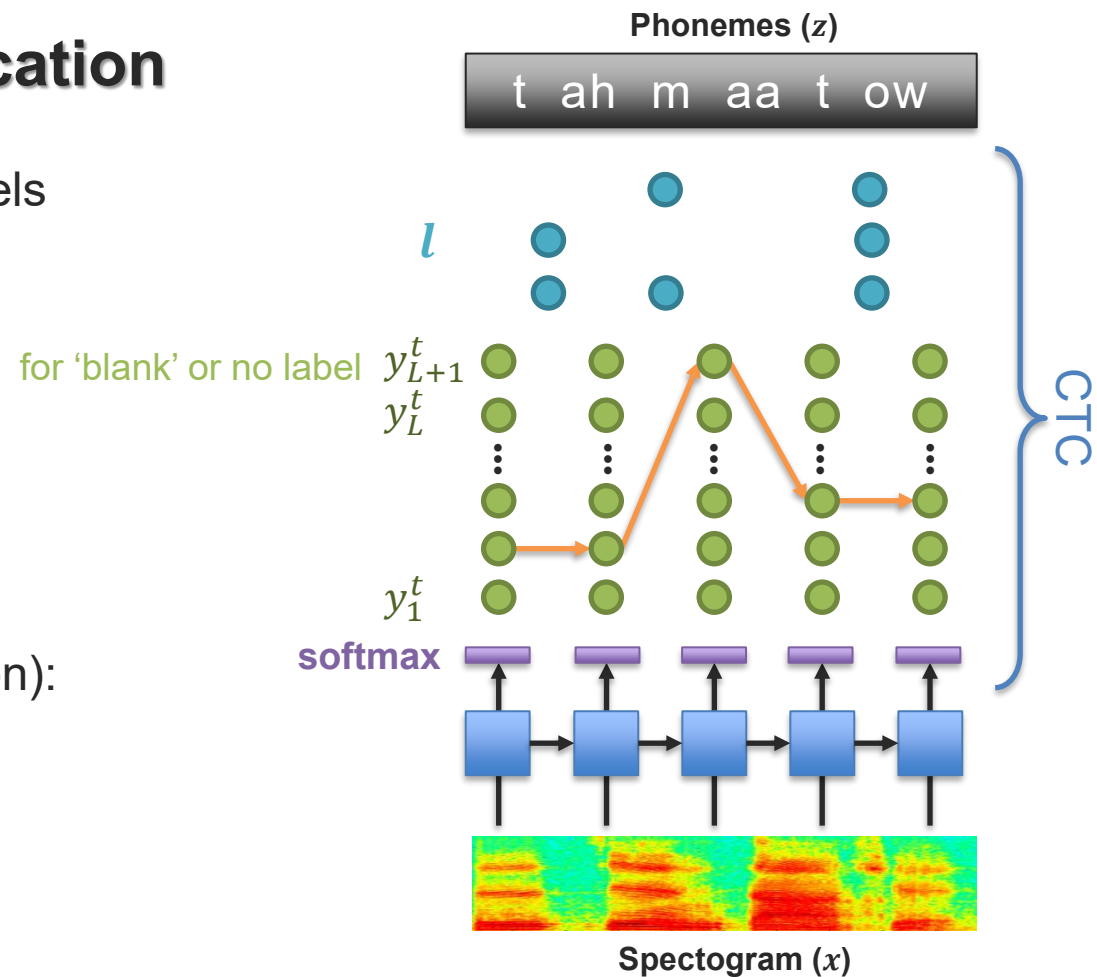
Challenge: many-to-1 alignment



Discretization – A Classification Approach

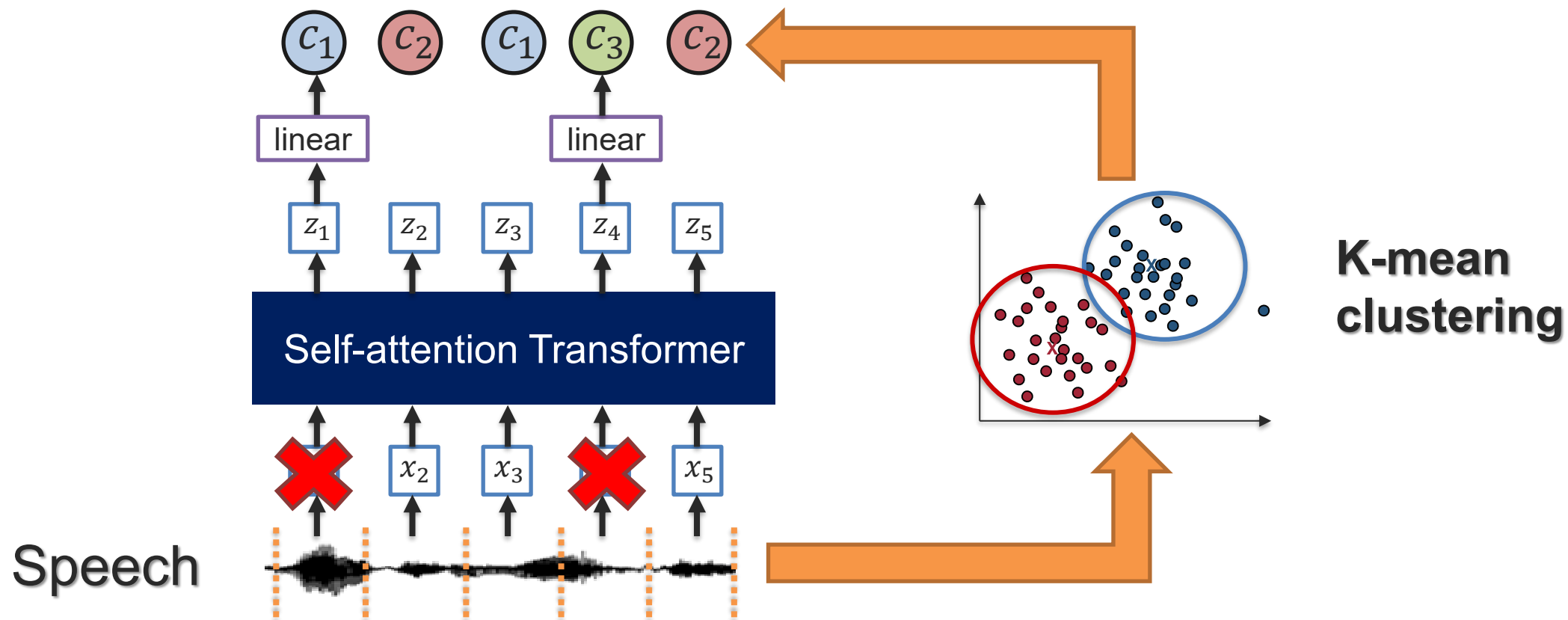
Connectionist Temporal Classification

- ④ Most probable sequence labels
- ③ Predicted labels l
- ② Path π over the activations:
- ① Output activations (distribution):



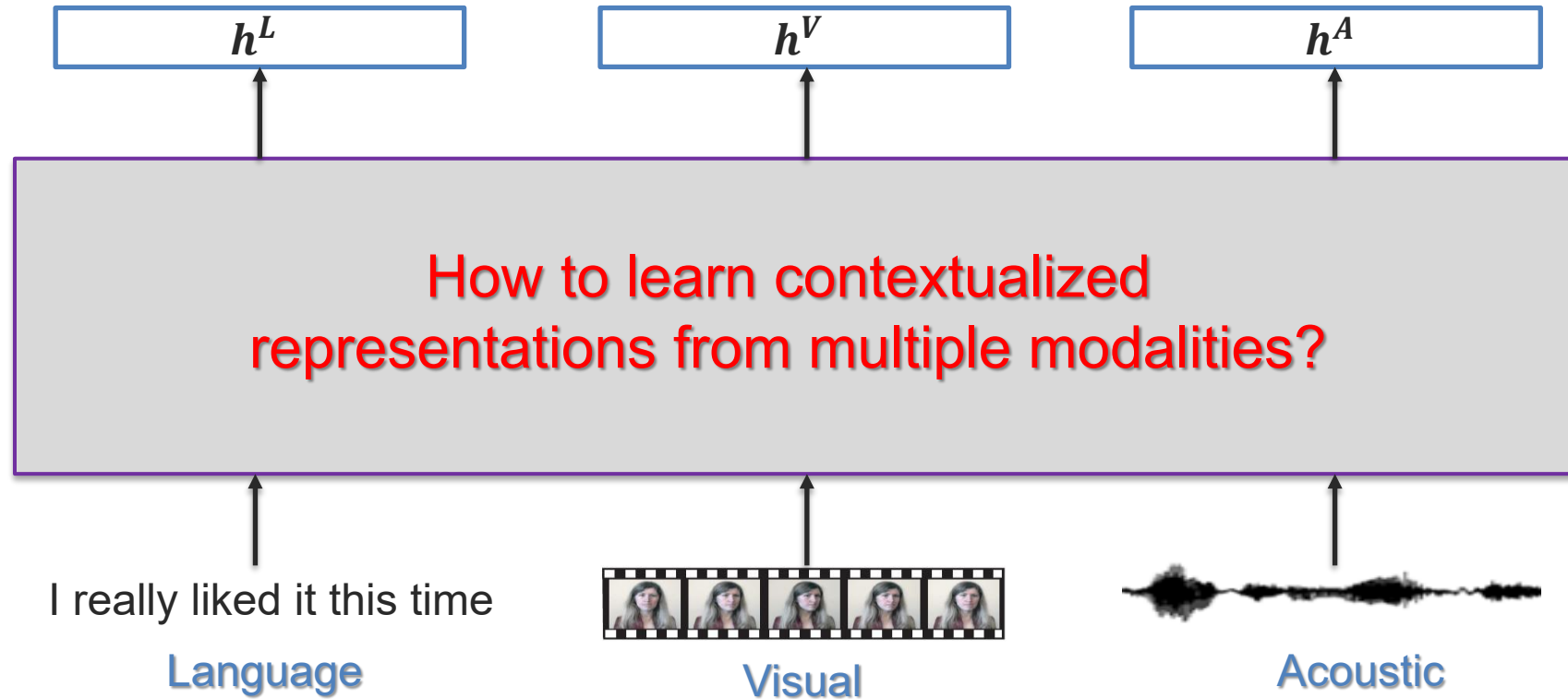
Discretization and Representation – Cluster-based Approaches

HUBERT: Hidden-Unit BERT



Alignment + Representation

Multimodal Alignment and Representation

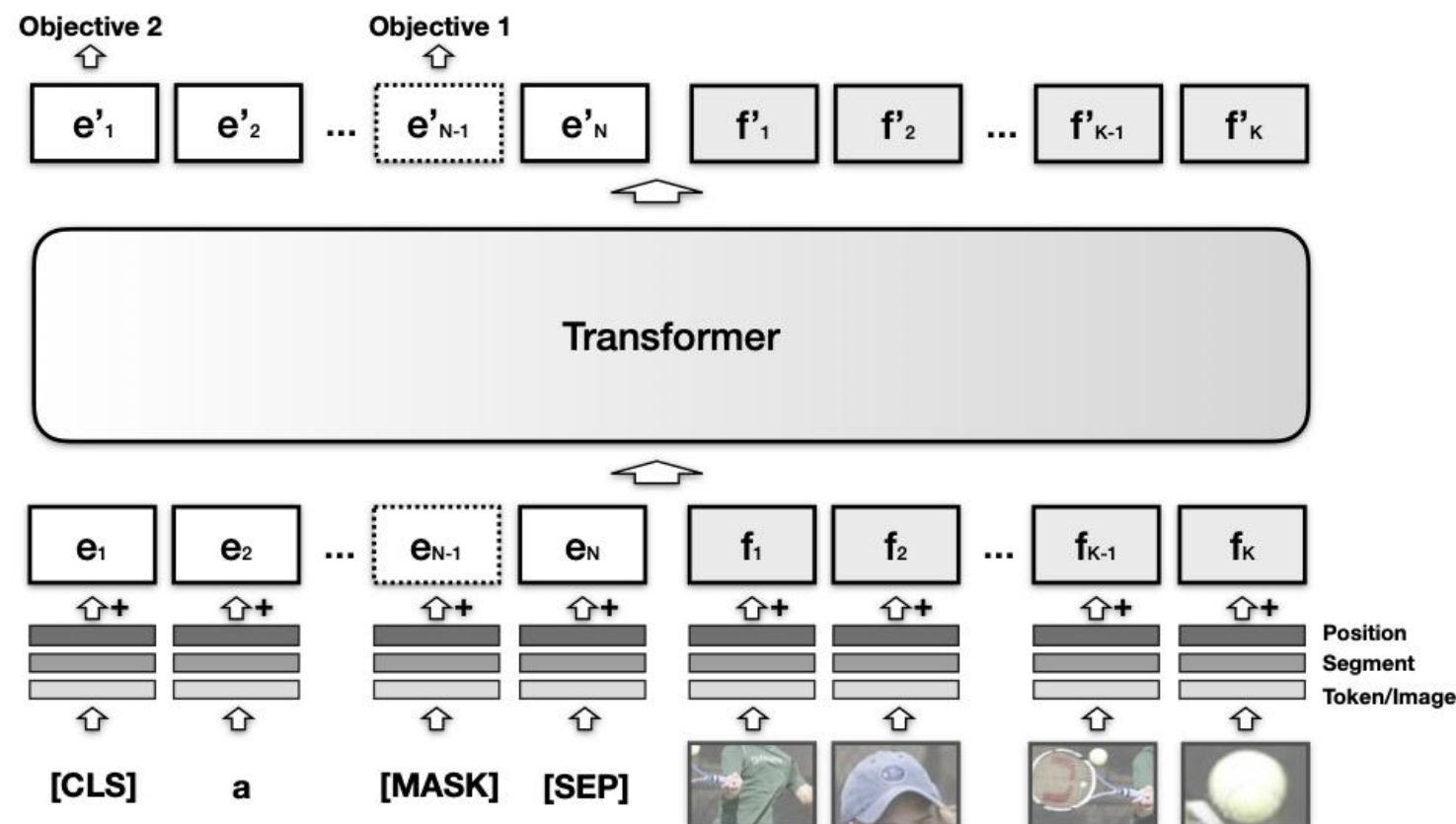


Option 1: Concatenate modalities and learn BERT transformer

VisualBERT

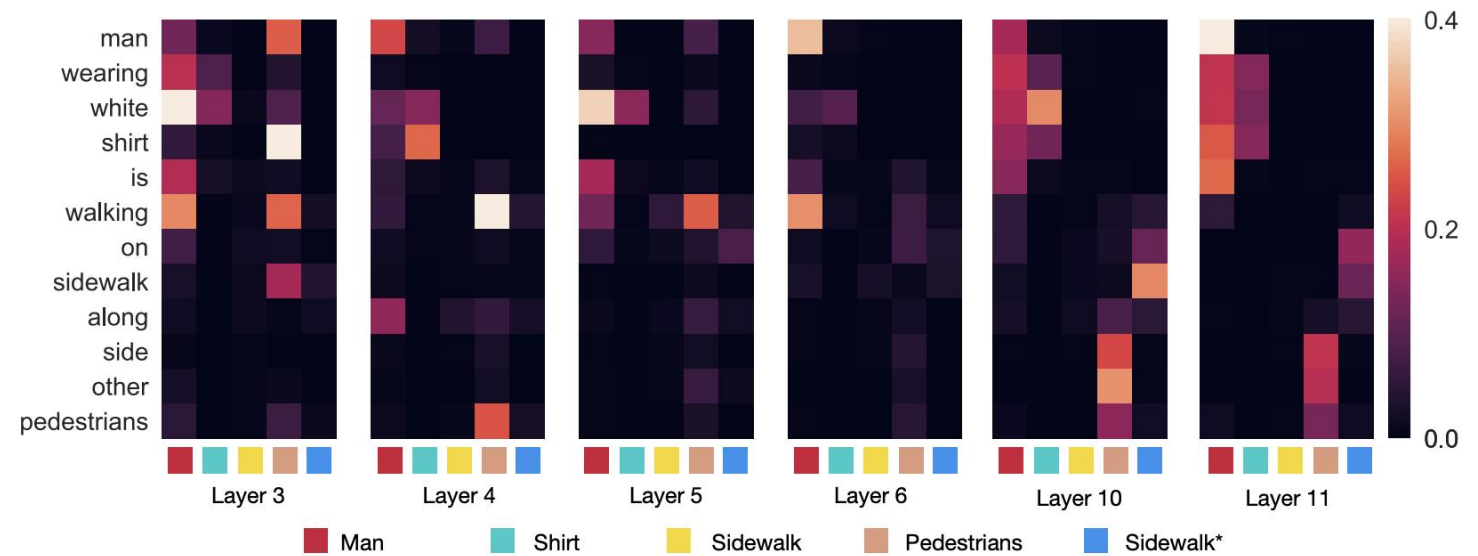
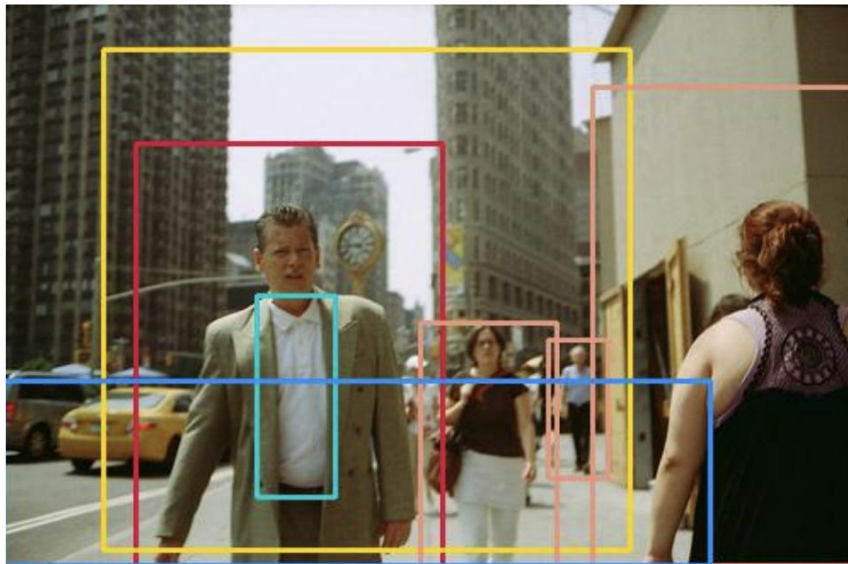


A person hits a ball with a tennis racket



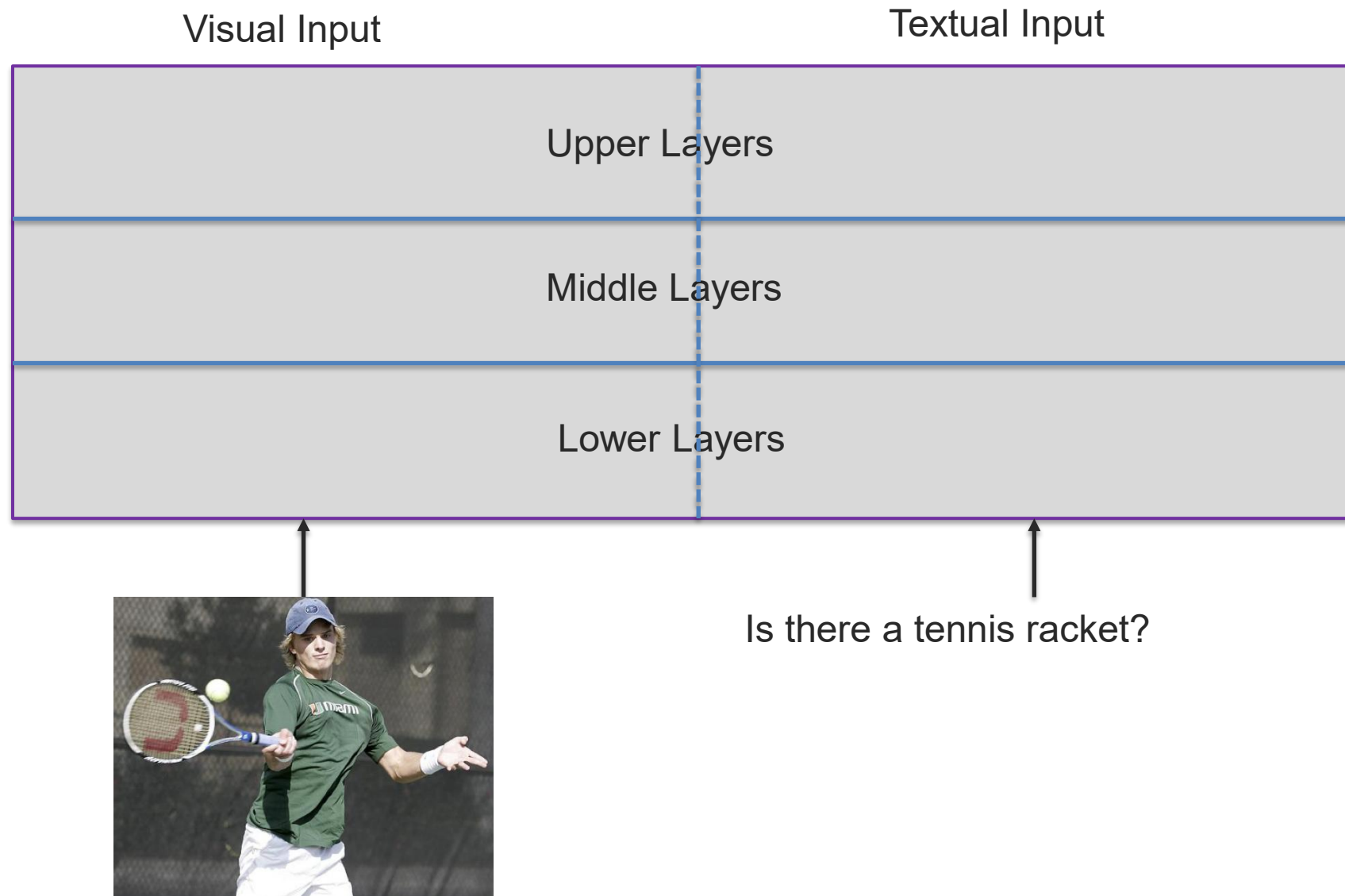
Li, Liunian Harold, et al. "Visualbert: A simple and performant baseline for vision and language." *ACL* (2020).

VisualBERT

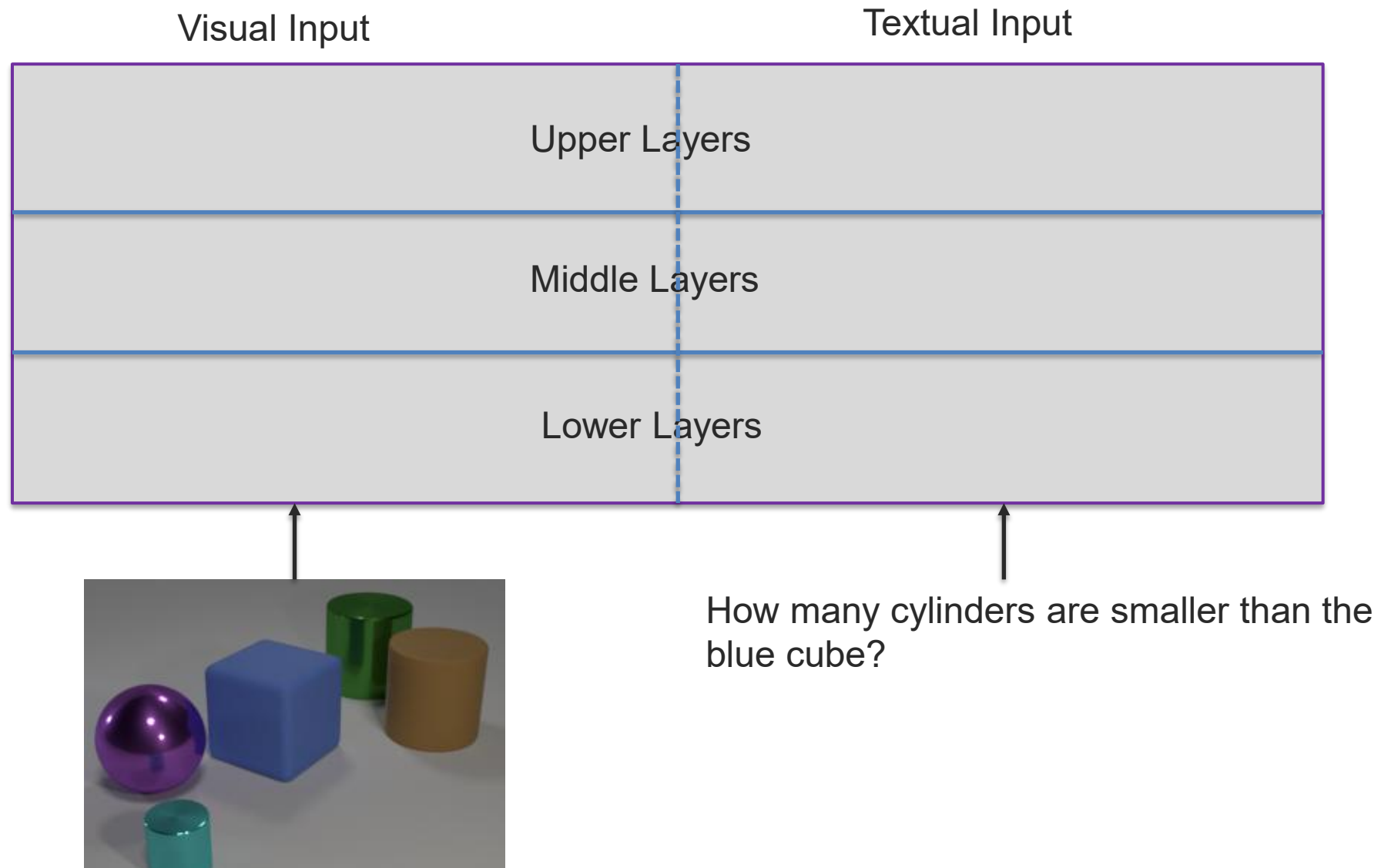


Li, Liunian Harold, et al. "Visualbert: A simple and performant baseline for vision and language." *ACL* (2020).

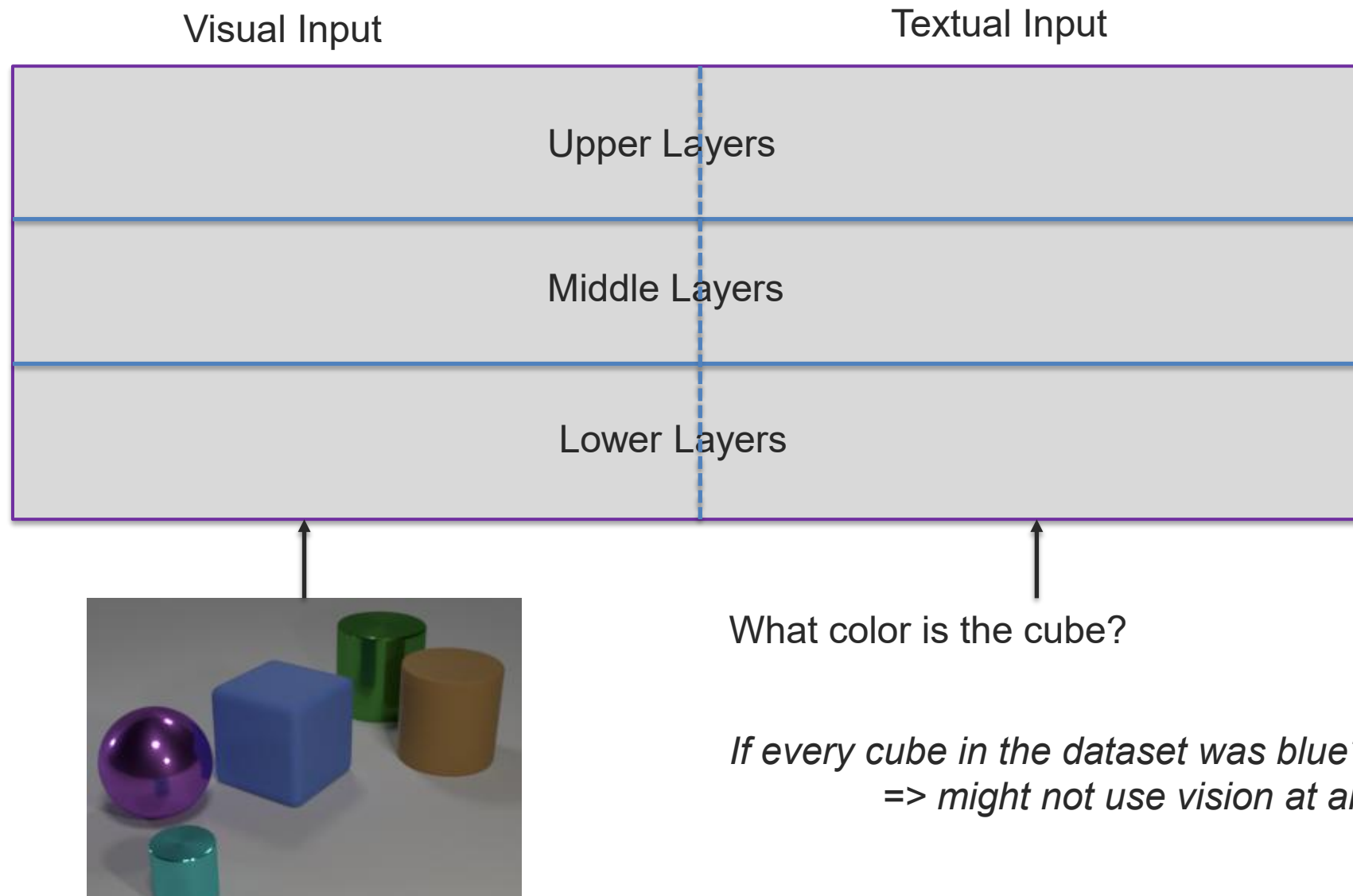
Where Do Cross-Modal Interactions Happen?



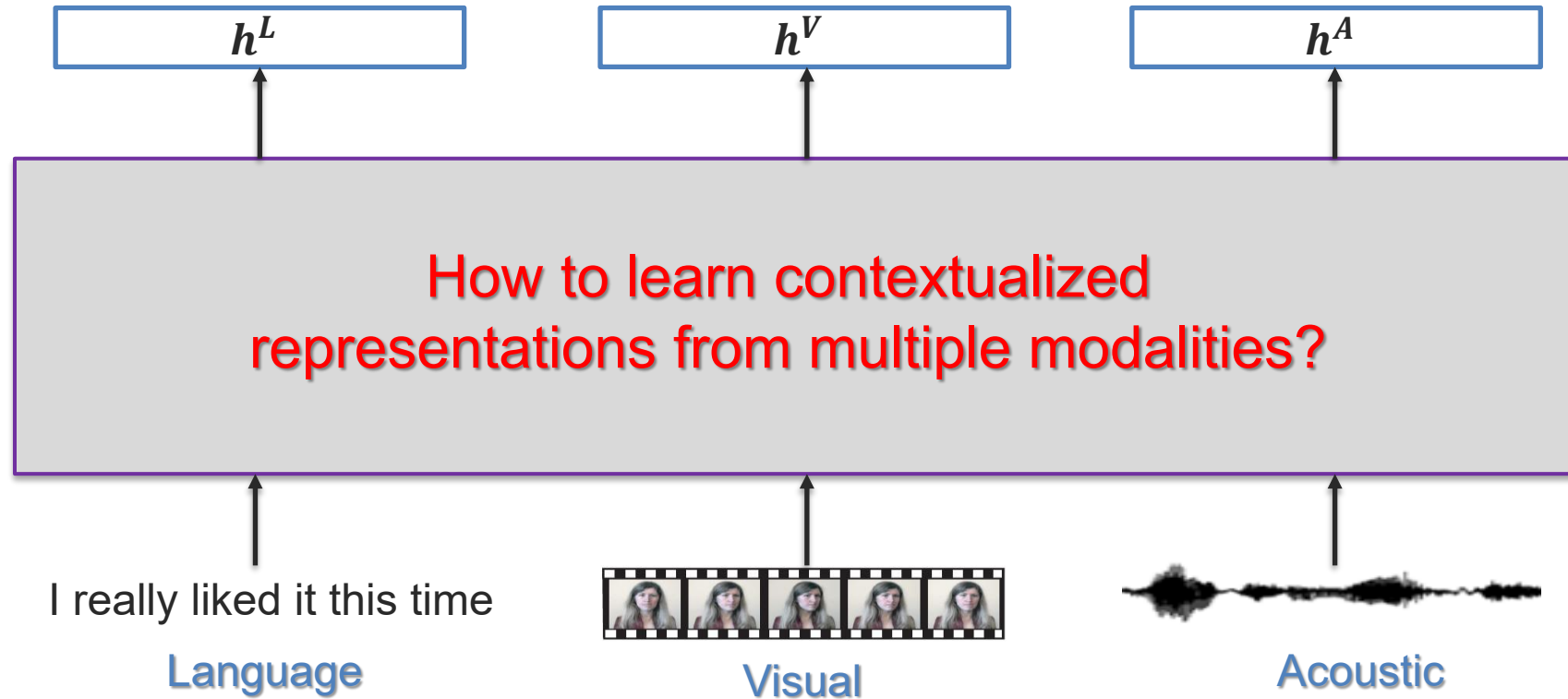
Where Do Cross-Modal Interactions Happen?



Do Cross-Modal Interactions Happen?

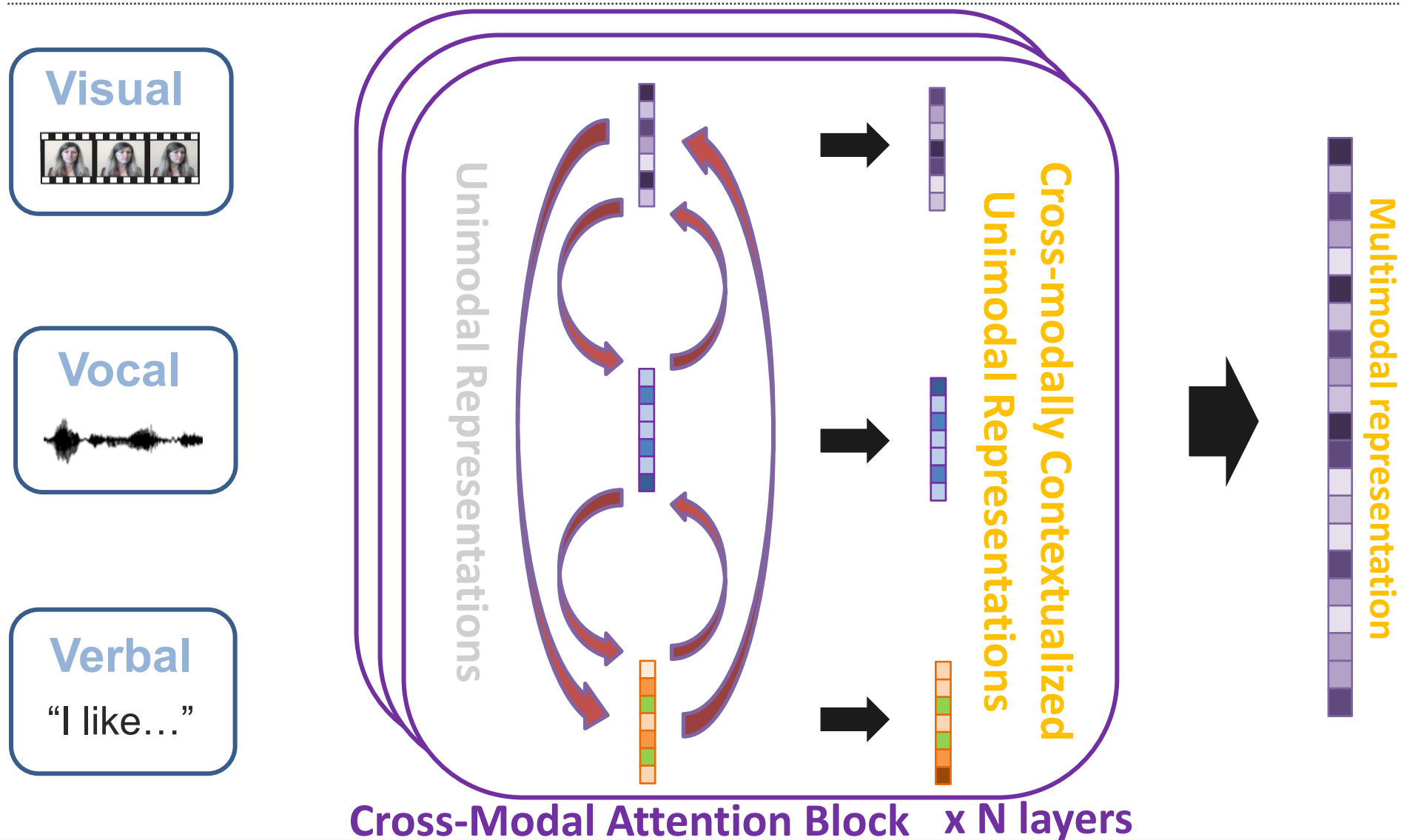


Multimodal Embeddings

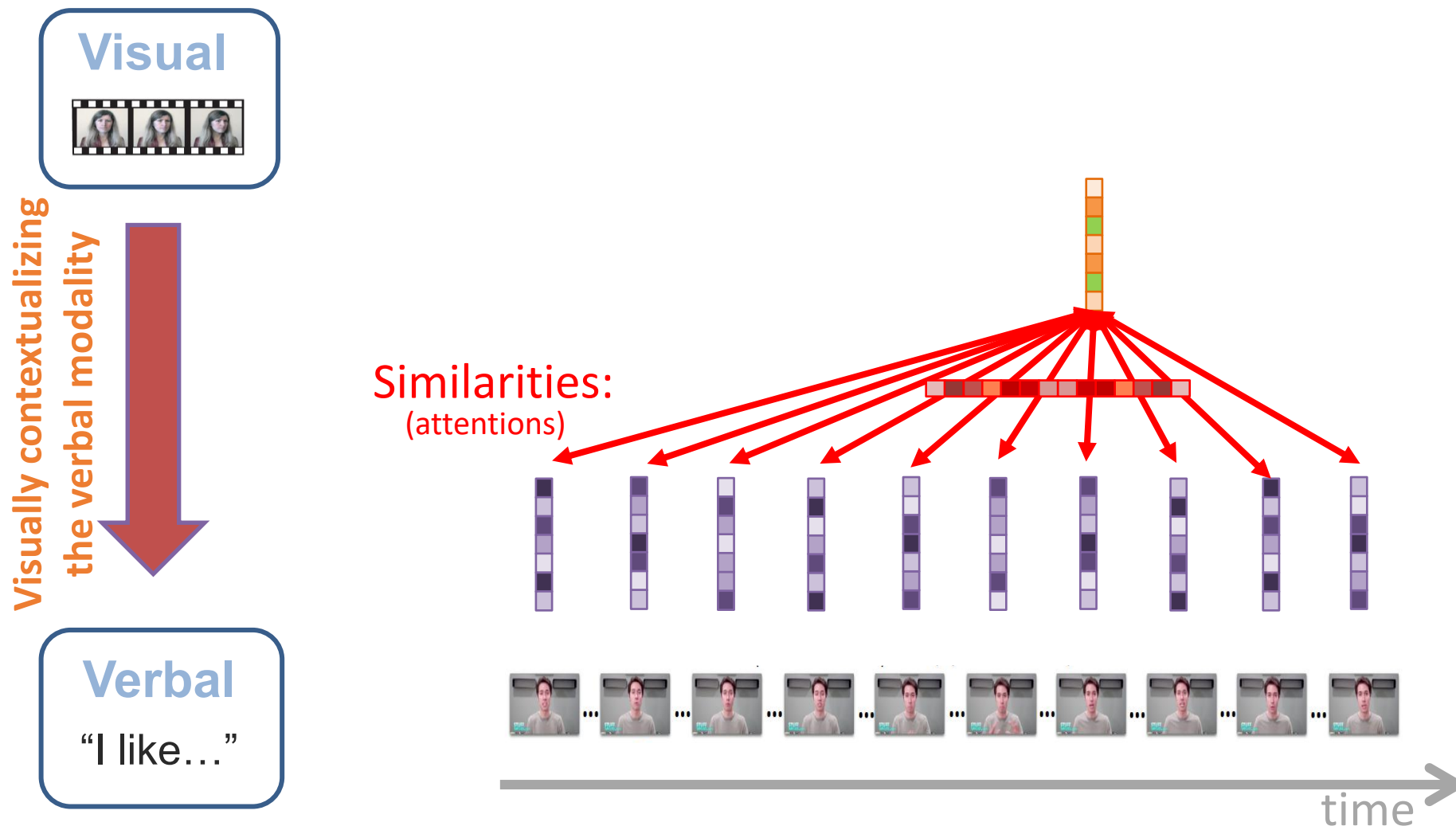


Option 2: Look at pairwise interactions between modalities

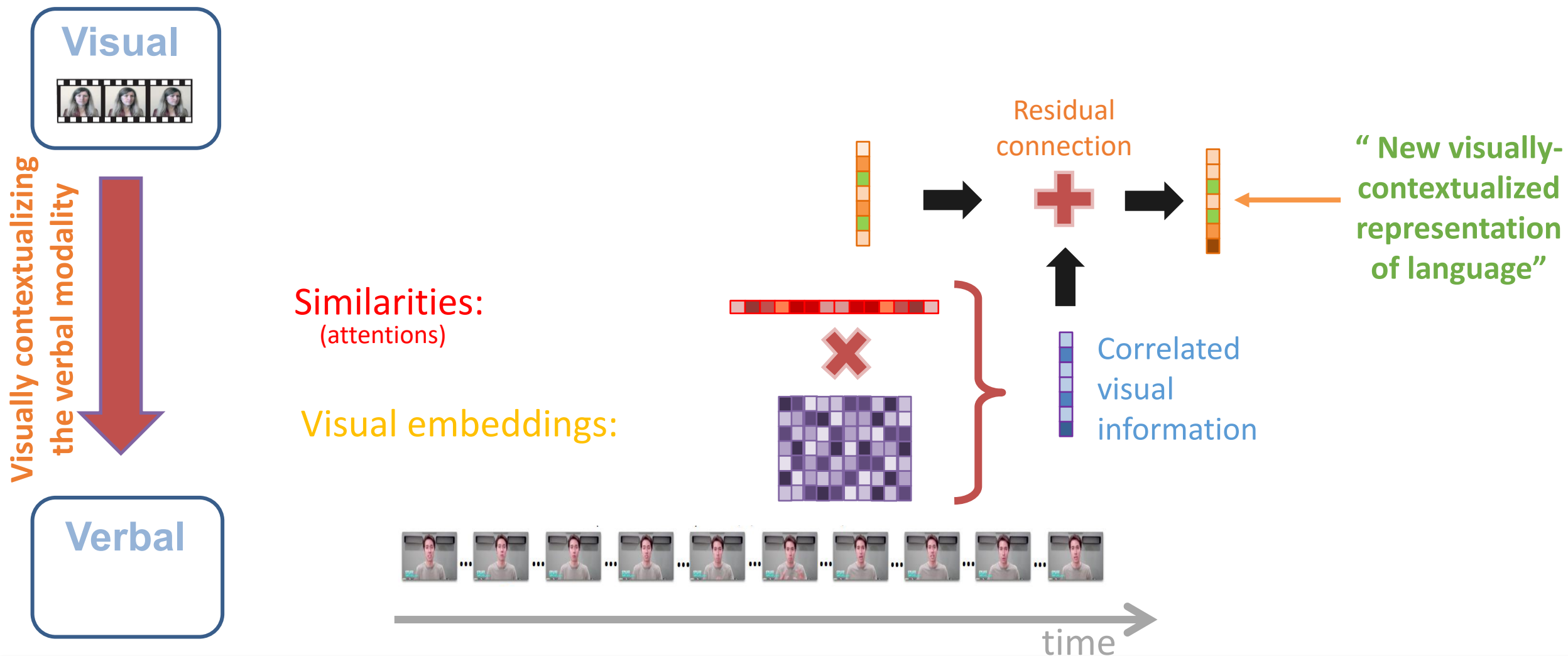
Multimodal Transformer – Pairwise Cross-Modal



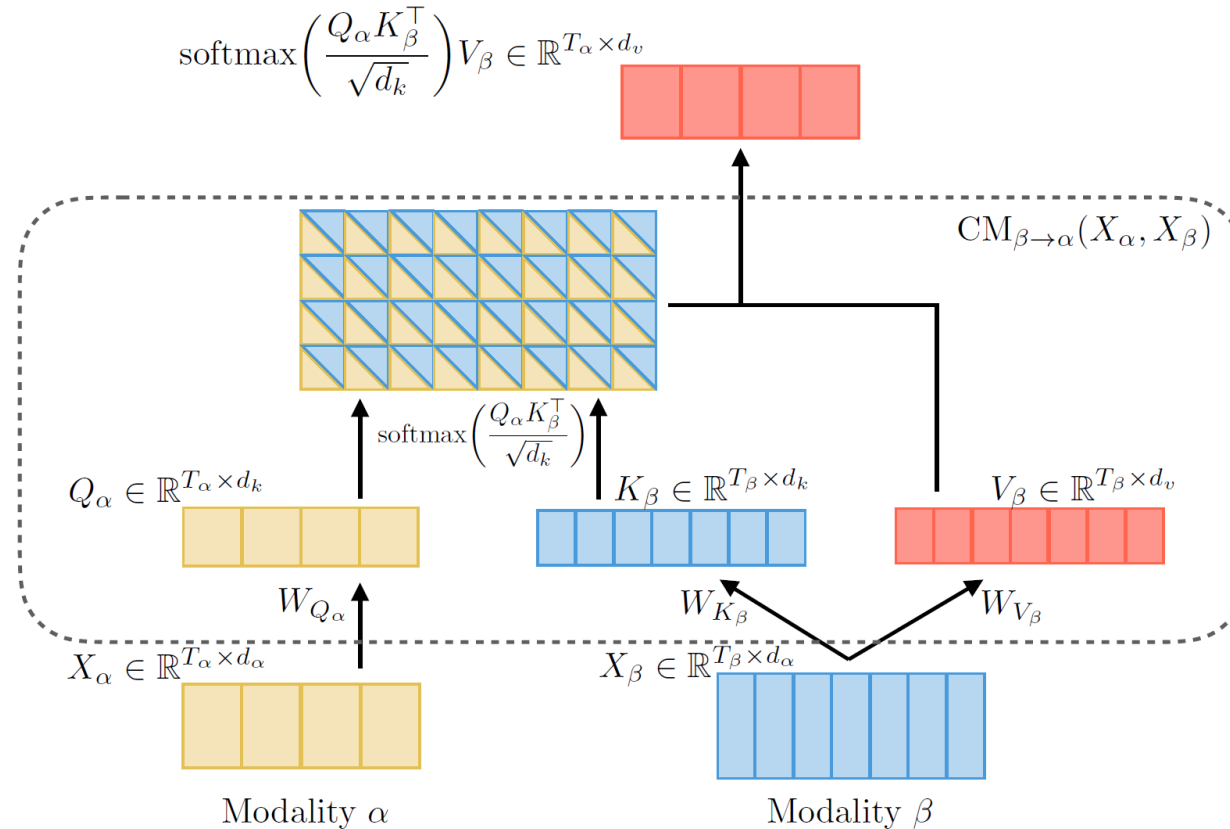
Cross-Modal Transformer Module ($V \rightarrow L$)



Cross-Modal Transformer Module ($V \rightarrow L$)

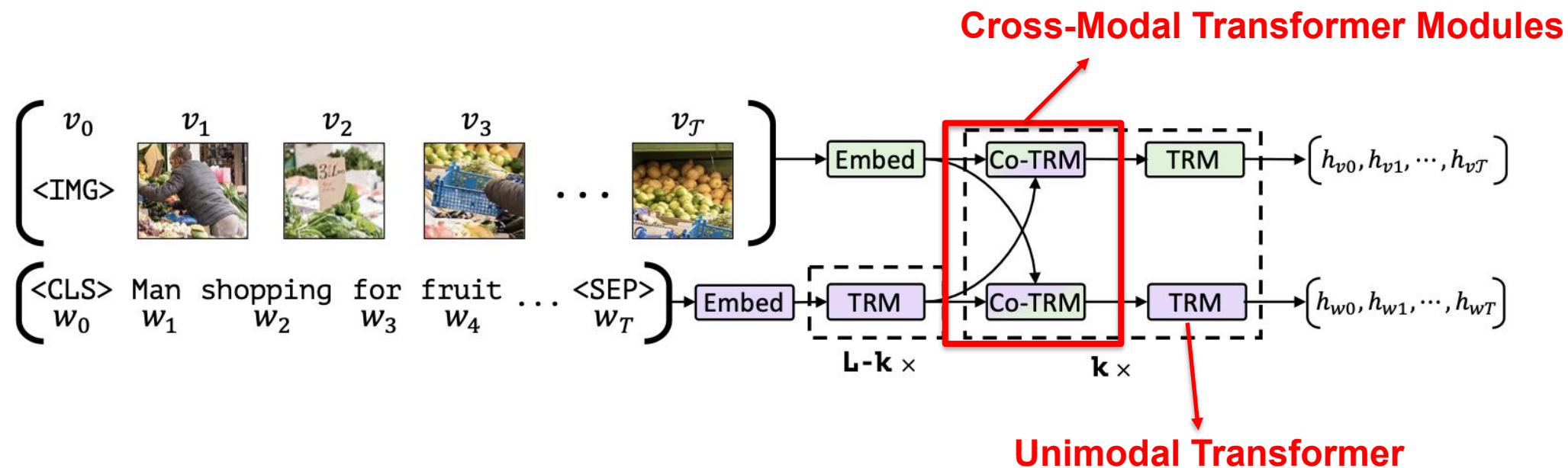


Cross-Modal Transformer Module ($\beta \rightarrow \alpha$)



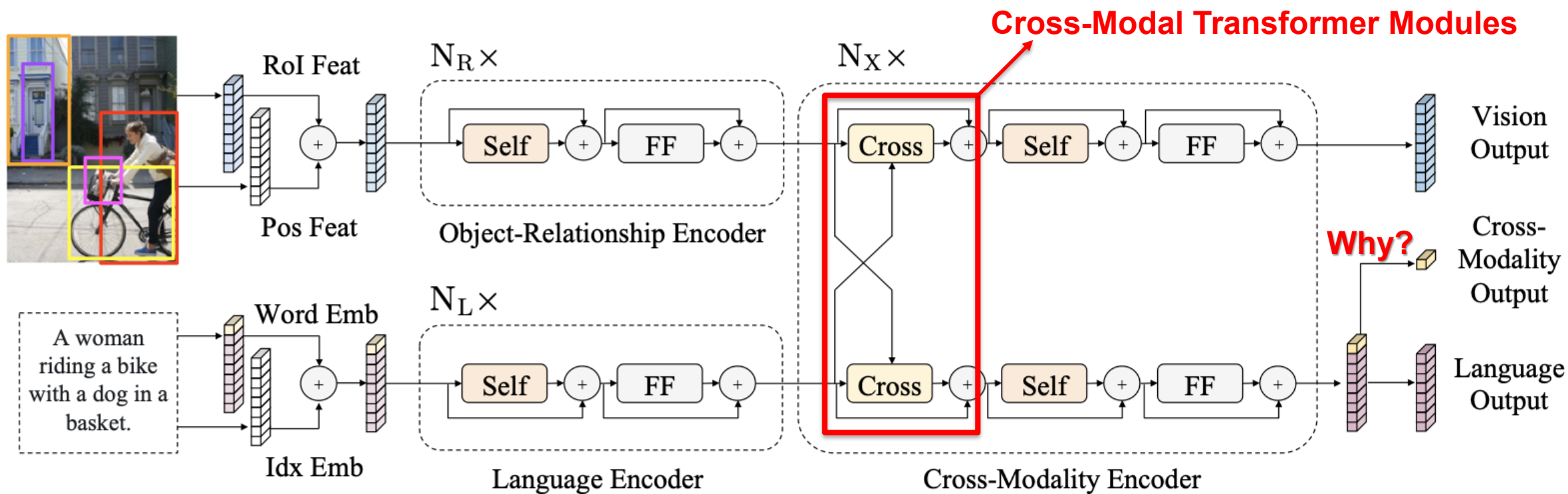
Tsai et al., Multimodal Transformer for Unaligned Multimodal Language Sequences, ACL 2019

ViLBERT



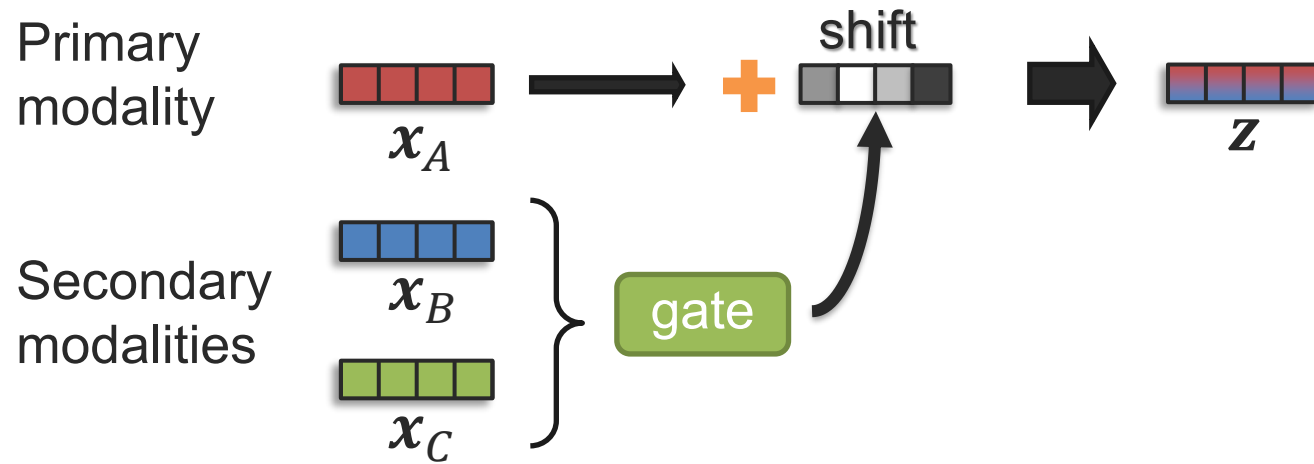
Lu, Jiasen, et al. "Vilbert: Pretraining task-agnostic visiolinguistic representations for vision-and-language tasks." *arXiv* (August 6, 2019).

LXMERT



Tan, Hao, and Mohit Bansal. "Lxmert: Learning cross-modality encoder representations from transformers." *arXiv* (August 20, 2019).

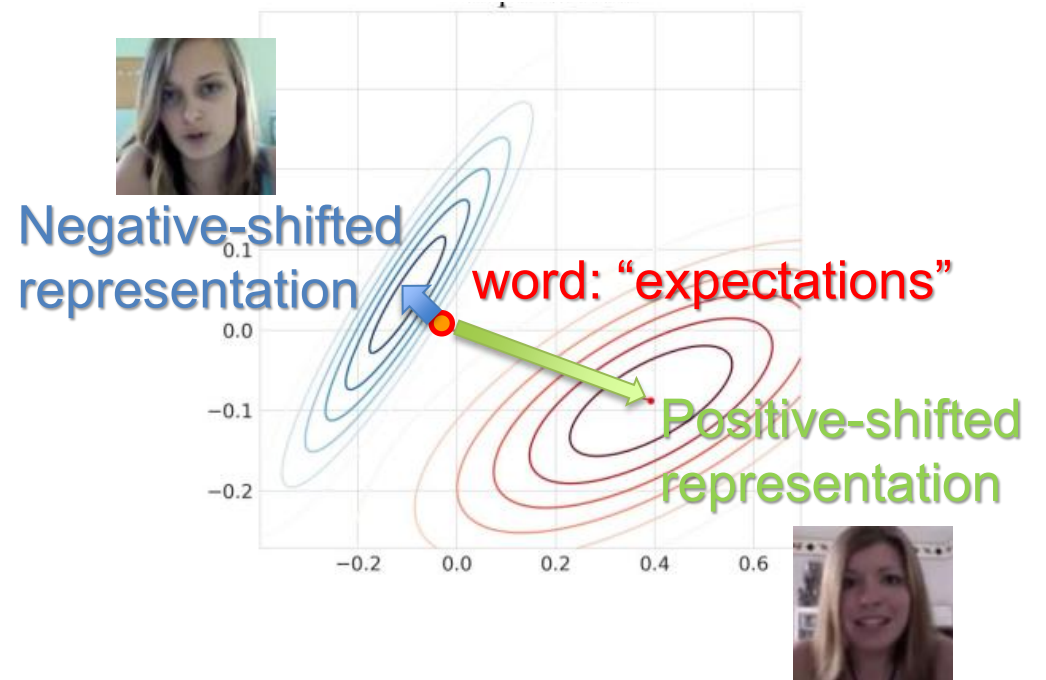
Reminder: Modality-Shifting Fusion



Example with language modality:

Primary modality: language

Secondary modalities: acoustic and visual



Modality-Shifting with Transformers

Multimodal Adaptation Gate (MAG) + BERT

