



Language
Technologies
Institute

Carnegie
Mellon
University

Multimodal Machine Learning

Lecture 4.1: 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

Primary TAs

- Each team has a primary TA
 - Your primary TA should have contacted you to schedule a first meeting together
 - Groups will be created in Piazza for each team
- Some projects may have a secondary TA, with complementary expertise

Schedule a meeting with your Primary TA this week!

First Project Assignment

Due date: Sunday 9/20 at 8pm

Four main sections:

- Introduction
 - Related work
 - Experimental setup
 - Research ideas
- 

Follows ICML paper format



The two main sections are related work and research ideas



teammates = # research ideas



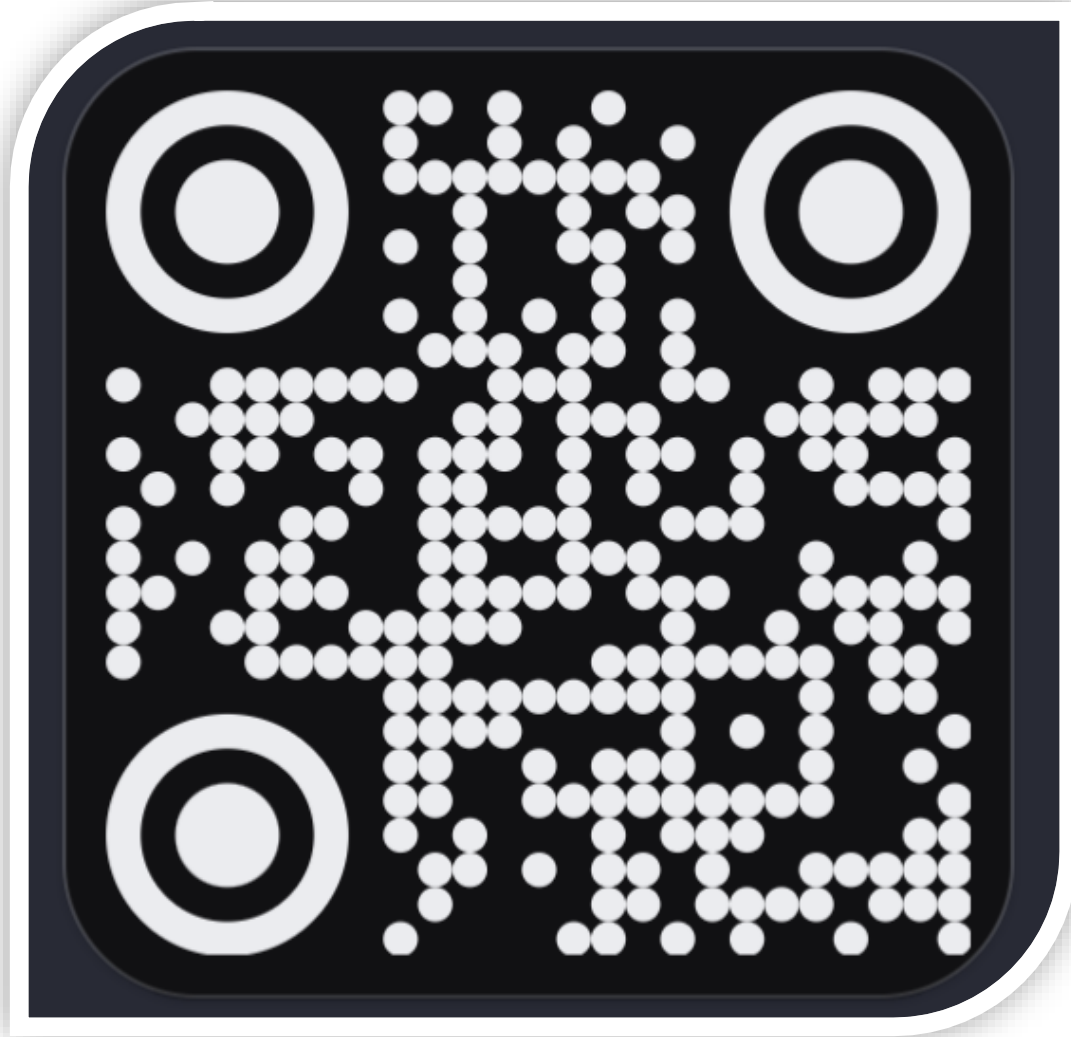
Page limit depends on team size:

- 4 students : 5 pages + references
- 5 students : 5.5 pages + references
- 6 students : 6 pages + references

Proposal Presentations: Team Meetings with Instructor

- No lectures on Tuesday 9/22 and Thursday 9/24
- 15-mins meeting with instructor
 - 5-mins presentations (instruction details will be shared soon)
 - All teammates are required to attend
 - Meetings next week
- Signup form will be shared via Piazza

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 4.1: Multimodal alignment

Louis-Philippe Morency

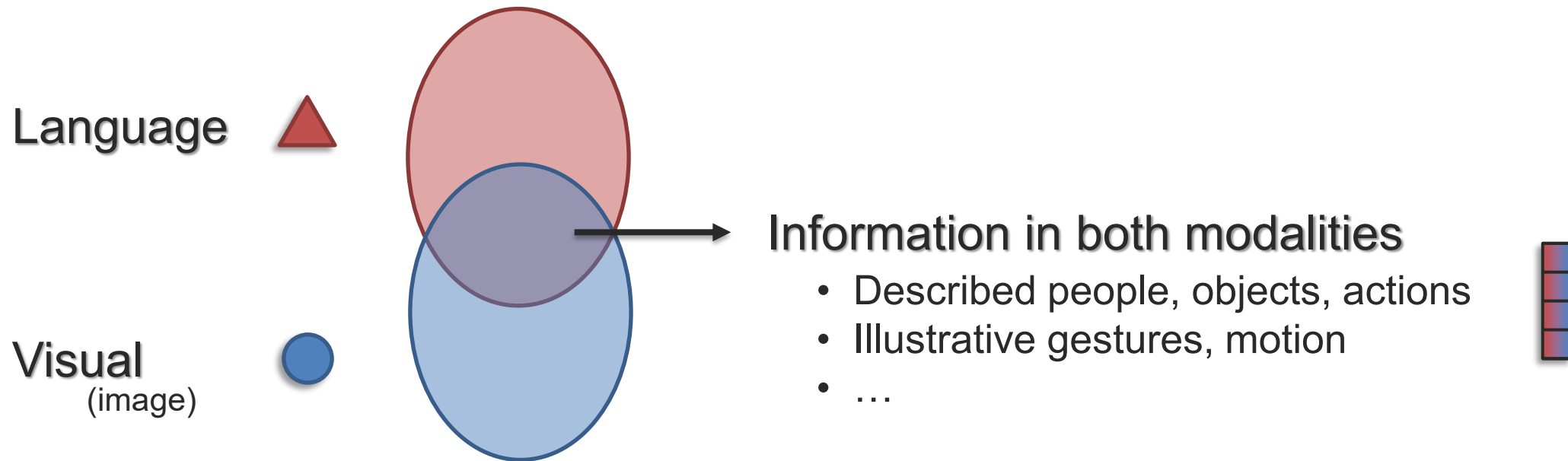
** Original course co-developed with Tadas Baltrusaitis.
Major revision co-developed with Paul Liang. Co-lecturers
included Yonatan Bisk, Daniel Fried and Carlos Busso.*

Lecture objectives

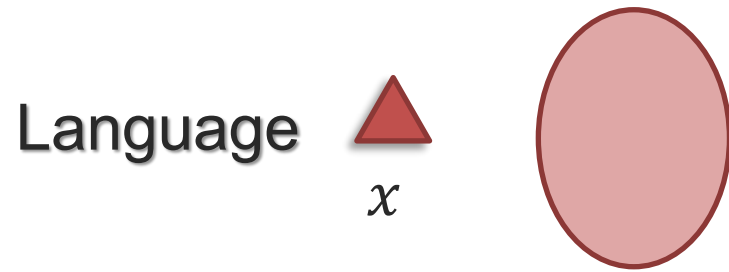
- Representation coordination
 - Information, entropy and mutual information
- Representation fission
 - Factorized multimodal representations
- Discrete alignment
 - Local alignment - Coordinated representations; soft attention
 - Global alignment - Assignment problem and optimal transport
- Continuous alignment
 - Continuous warping- Dynamic time warping
 - Discretization and segmentation

Representation Coordination

Multimodal Coordination – Information Theory



Information and Entropy – Information Theory



How much information in the modality?

Information Theory (Shannon, 1948)

Main intuition: “Information value” of a communicated message x depends on how random its content is

x : “1,1,1,1,1,1,1,1,1,1,1,1,1”

➡ Not very random... So, low information

x : “0,1,0,1,0,0,1,1,1,0,0,1”

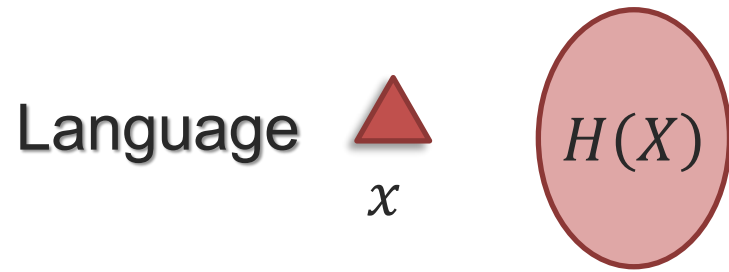
➡ More random... So, higher information

Information content $I(x)$

$$I(x) \sim \frac{1}{p(x)}$$

$$I(x) = \log \left(\frac{1}{p(x)} \right) = -\log(p(x))$$

Information and Entropy – Information Theory



How much information in the modality?

Information Theory (Shannon, 1948)

Information content $I(X) = -\log(p(X))$

➡ For discrete alphabet $\mathcal{X}$, then X is discrete random variable

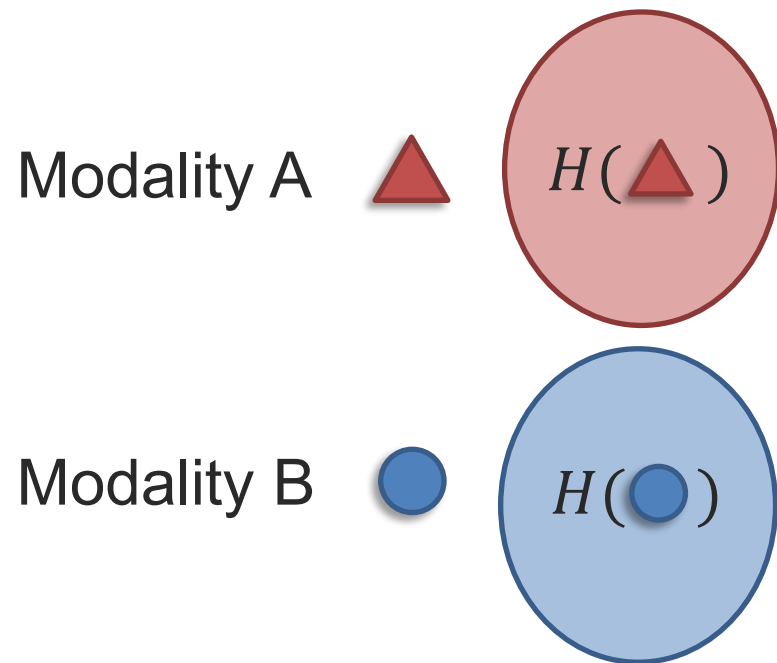
Entropy: weighted average of all possible outcomes from $\mathcal{X}$

$$H(X) = \mathbb{E}[I(X)] = \mathbb{E}[-\log(p(X))] = - \sum_{x \in \mathcal{X}} p(X) \log(p(X))$$

➡ Entropy can also be defined for continuous random variables

Entropy with Two Modalities

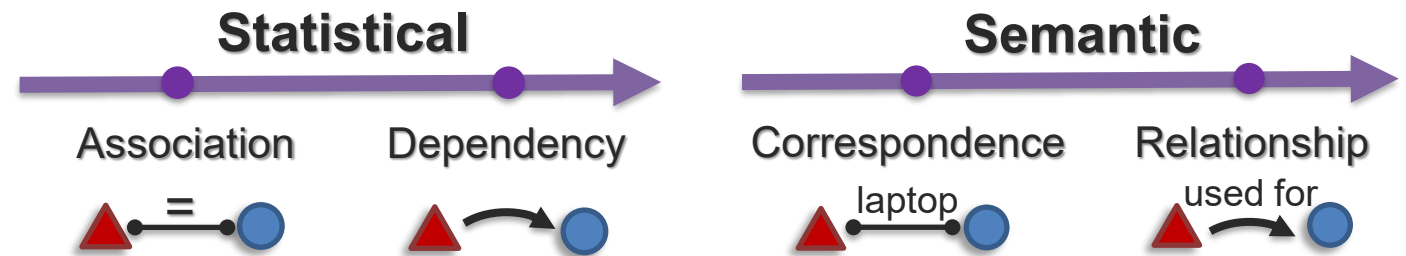
If no overlapping information



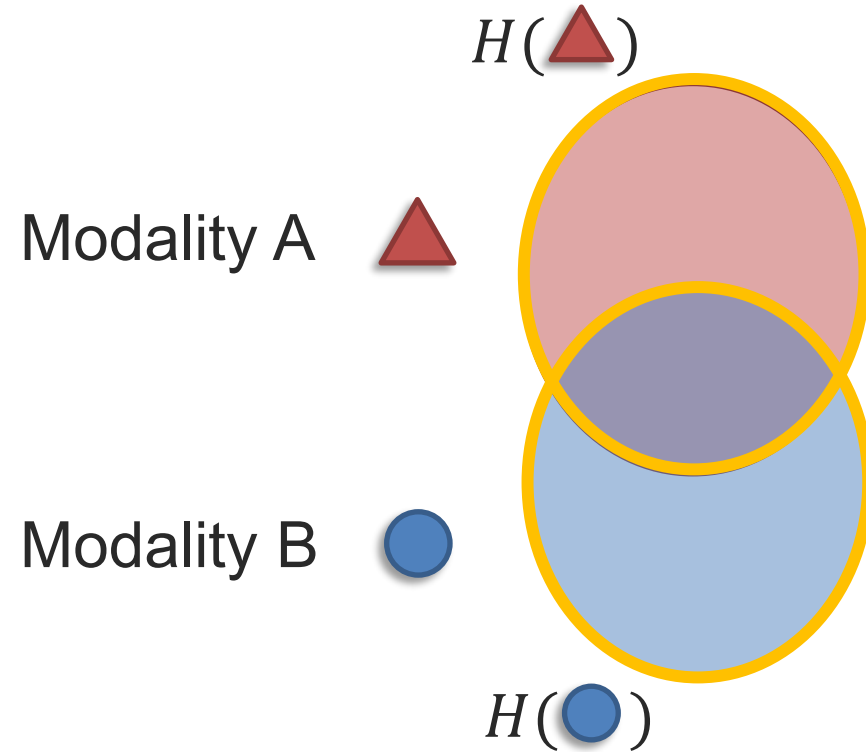
➡ But in most real-world scenarios, modalities are *inter-connected*



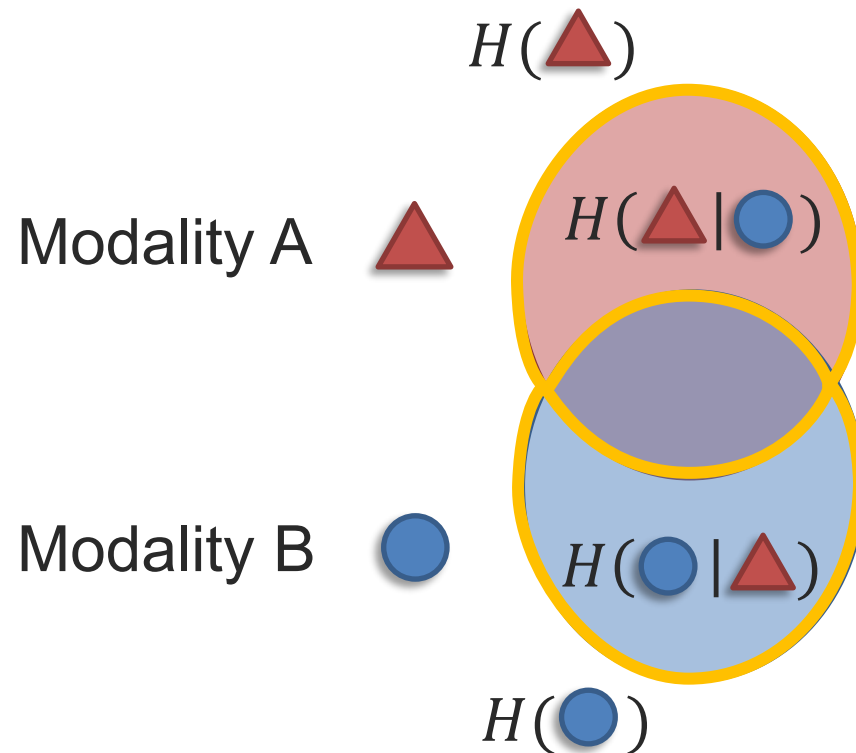
A **teacup** on the right of a **laptop** in a clean room.



Entropy with Two Modalities



Entropy with Two Modalities



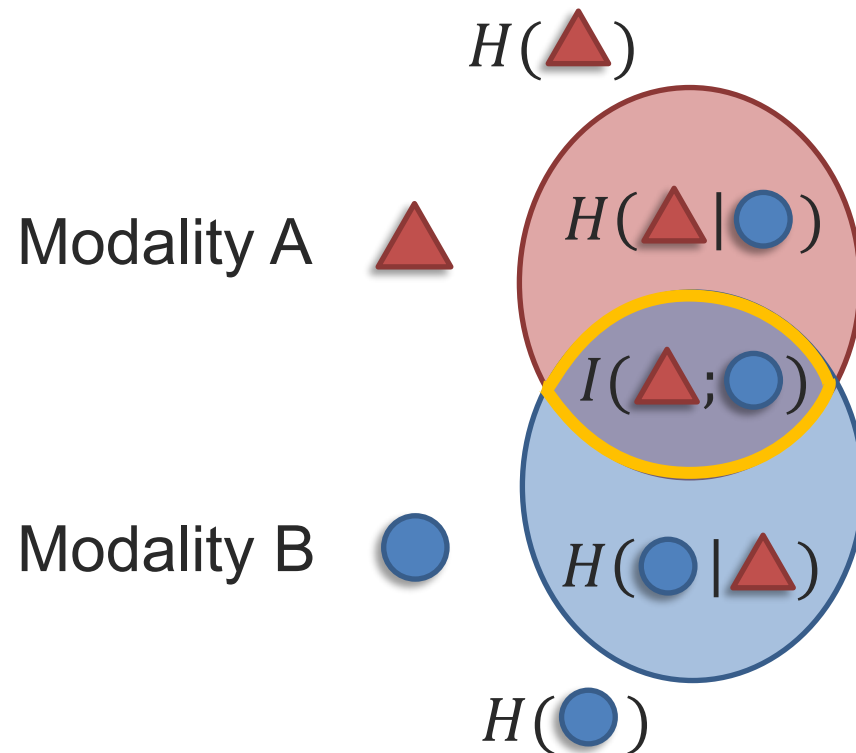
Conditional entropy $H(Y|X)$

$$H(Y|X) = -\mathbb{E}_{X,Y}[\log p(y|x)]$$

$$= -\mathbb{E}_{X,Y} \left[\log \frac{p(x,y)}{p(x)} \right]$$

If X and Y independent, $H(Y|X) = H(Y)$.
If X fully determines Y , then $H(Y|X) = 0$.

Entropy with Two Modalities



Mutual information $I(X; Y)$

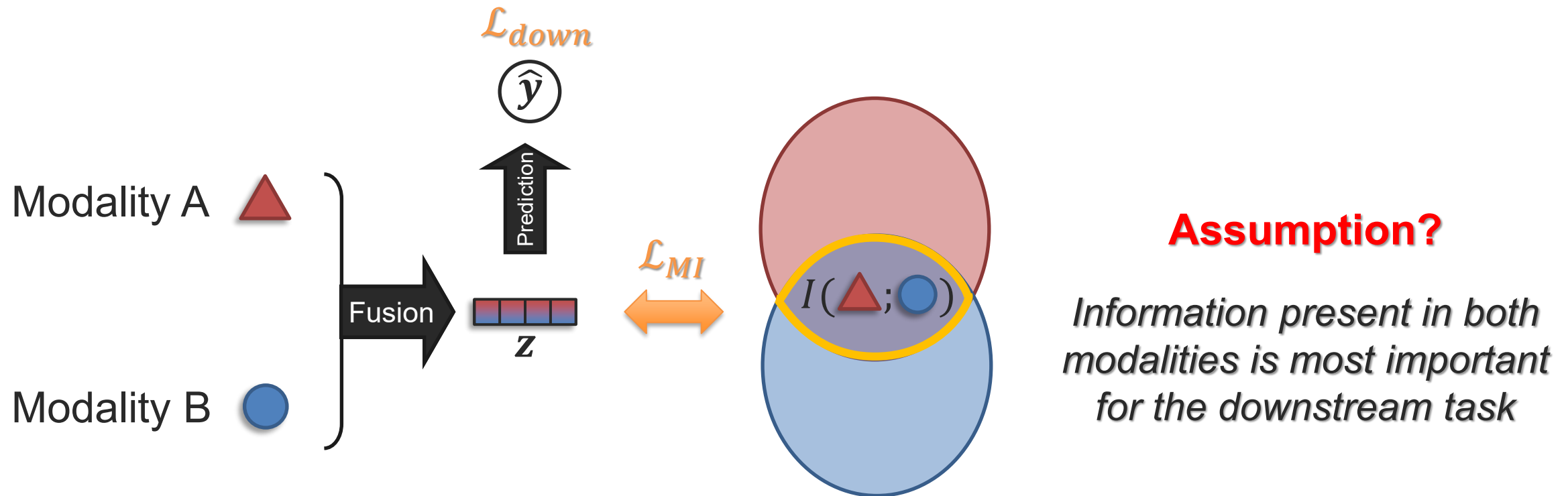
$$I(X; Y) = H(X) - H(X|Y)$$

$$= \mathbb{E}_{X,Y} \left[\log \frac{1}{P_X(x)} + \log \frac{P_{XY}(x, y)}{P_Y(y)} \right]$$

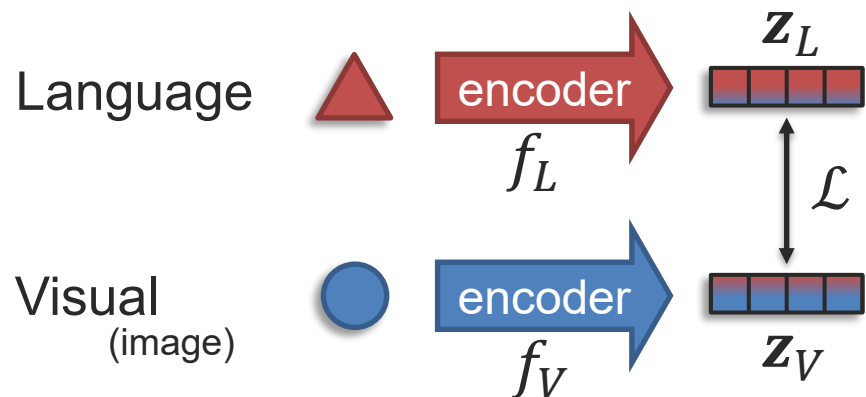
$$I(X; Y) = \mathbb{E}_{X,Y} \left[\log \frac{P_{XY}(x, y)}{P_X(x)P_Y(y)} \right]$$

using KL-divergence $\leftarrow I(X; Y) = D_{KL}(P_{XY}(x, y) \parallel P_X(x)P_Y(y))$

Multimodal Fusion with Mutual Information



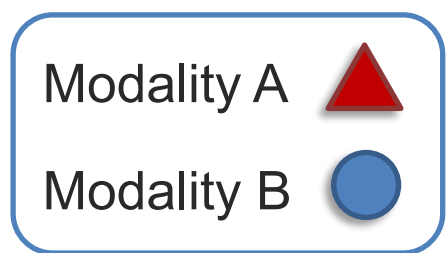
Contrastive Learning and Connected Modalities



Popular contrastive loss: InfoNCE

$$\mathcal{L} = -\frac{1}{N} \sum_{i=1}^N \log \frac{\text{sim}(\mathbf{z}_A^i, \mathbf{z}_B^i)}{\sum_{j=1}^N \text{sim}(\mathbf{z}_A^i, \mathbf{z}_B^j)}$$

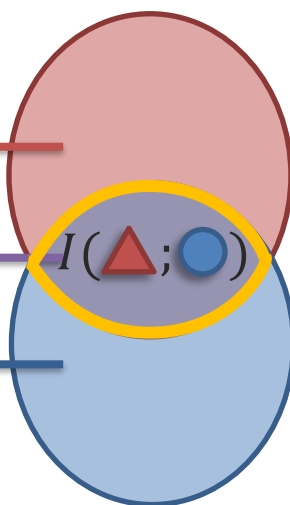
Connected modalities:



unique

shared

unique



Mutual information $I(X; Y)$

$$\mathbb{E}_{X,Y} \left[\log \frac{P_{XY}(x, y)}{P_X(x)P_Y(y)} \right]$$

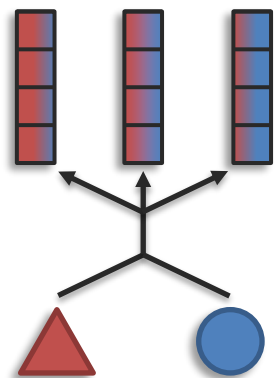


CLIP focuses on
shared connections

Representation

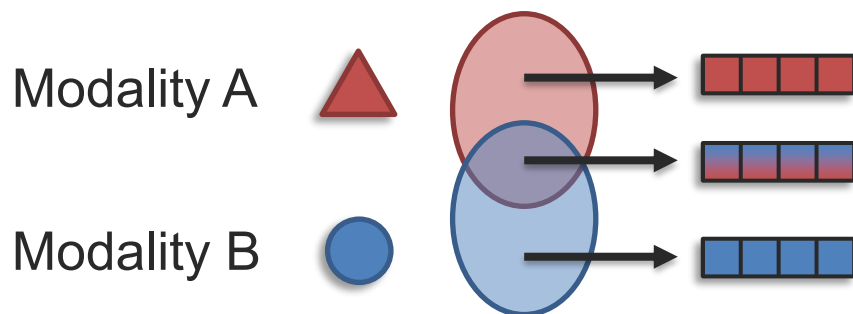
Fission

Sub-Challenge 1c: Representation Fission

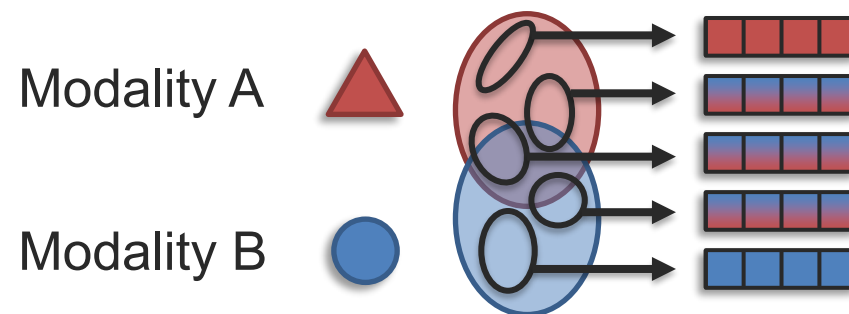


Definition: Learning a new set of representations that reflects multimodal internal structure such as data factorization or clustering

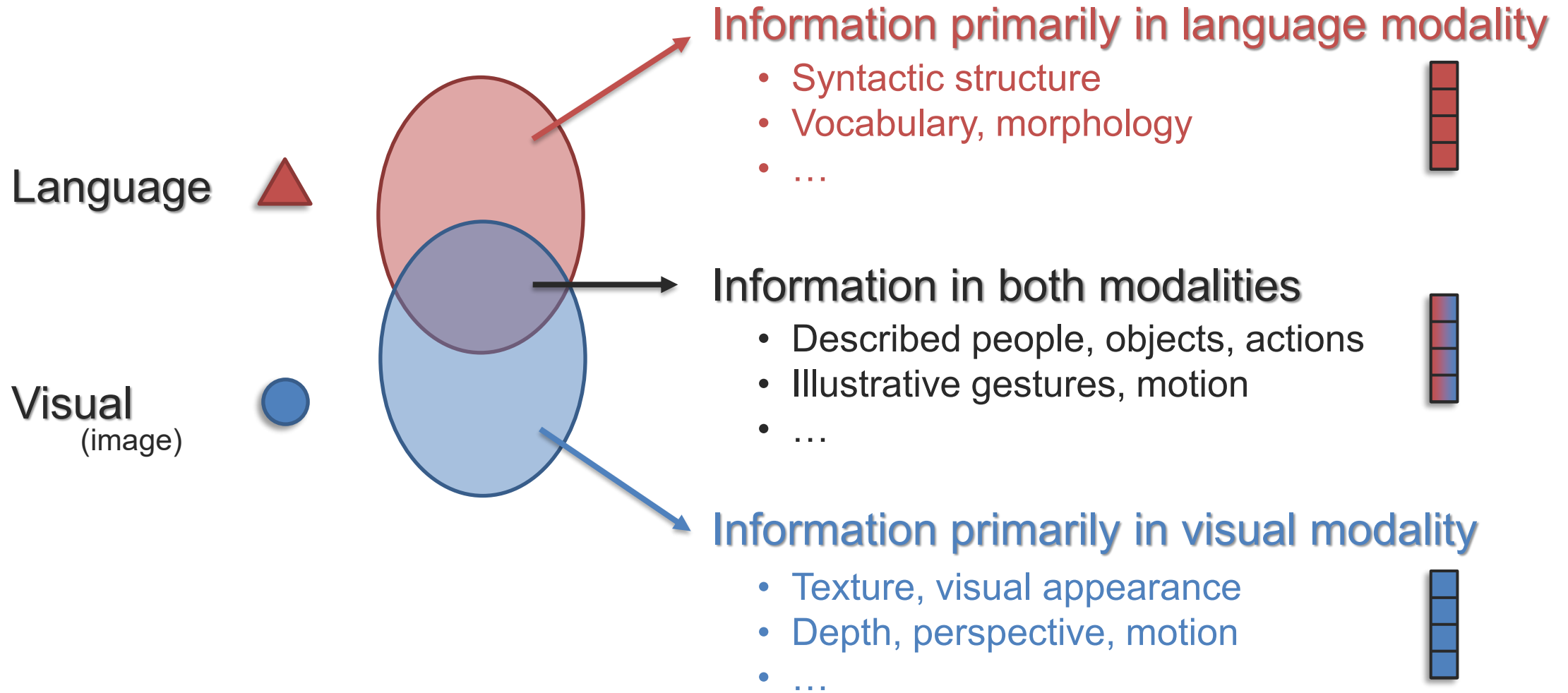
Modality-level fission:



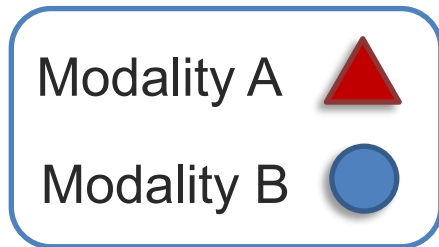
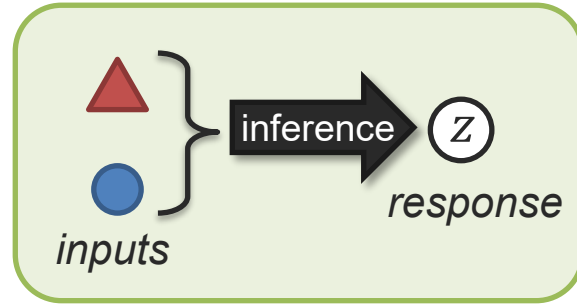
Fine-grained fission:



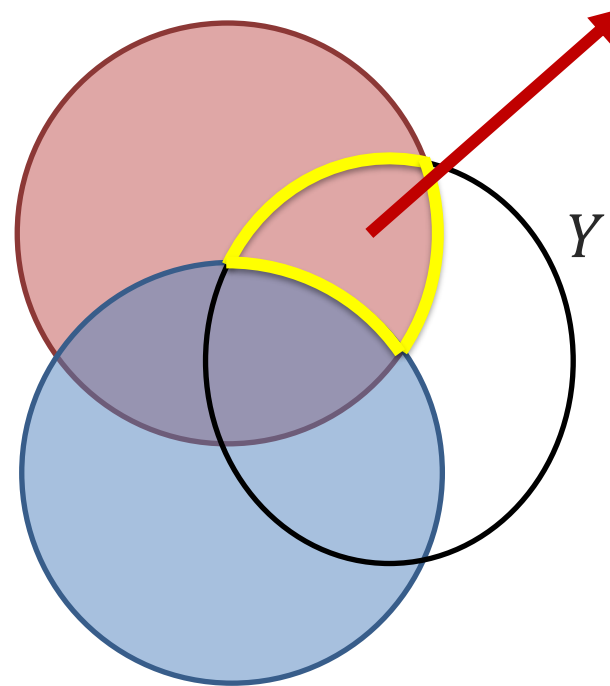
Modality-Level Fission



Modeling Multimodal Interactions

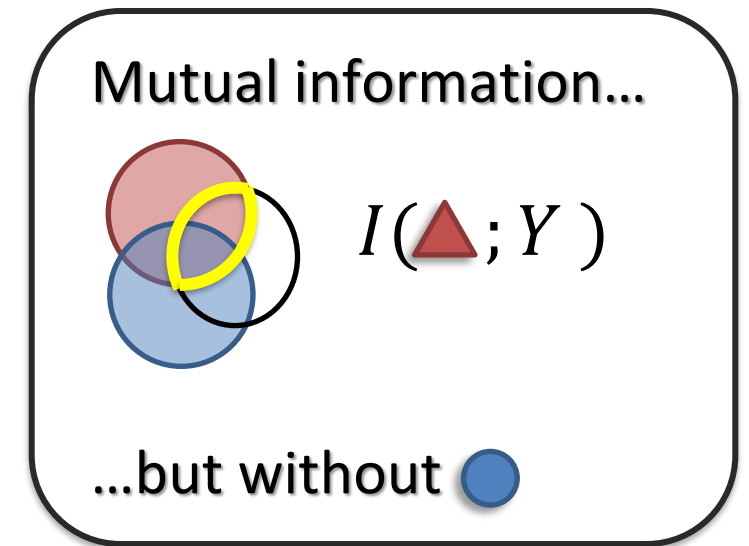


Information theory as a framework for modeling multimodal interactions between input modalities and responses

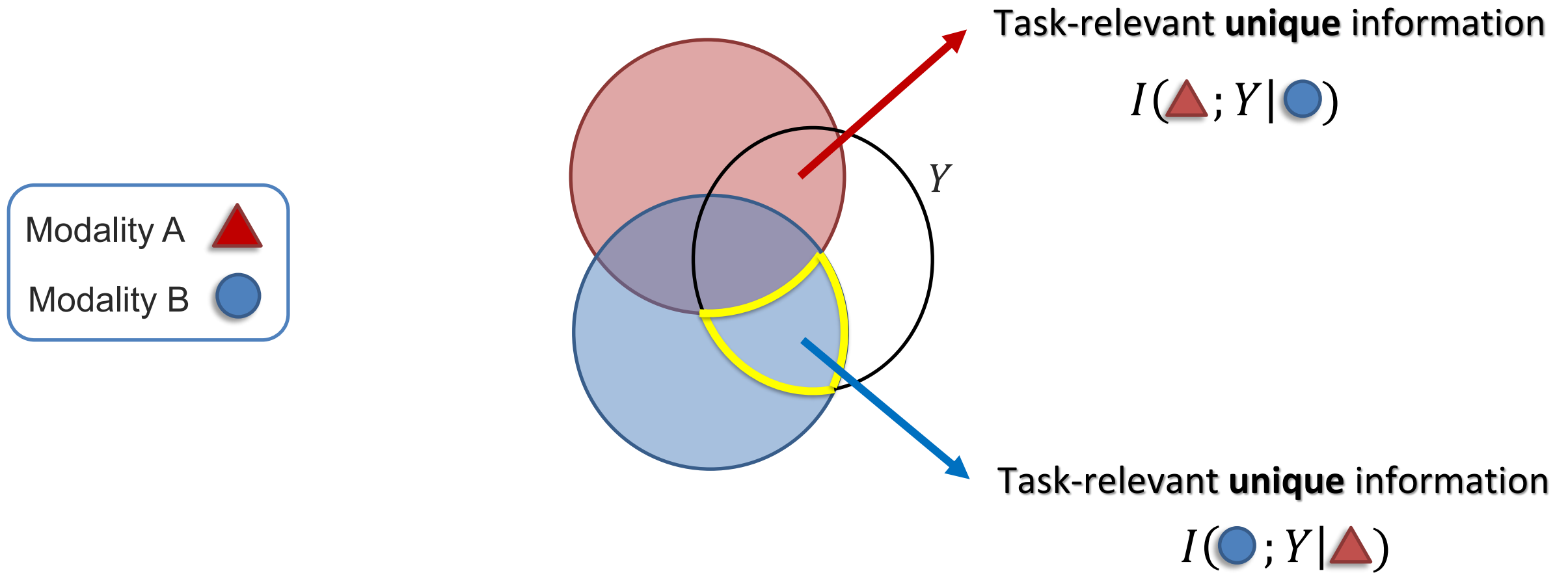


Task-relevant **unique** information

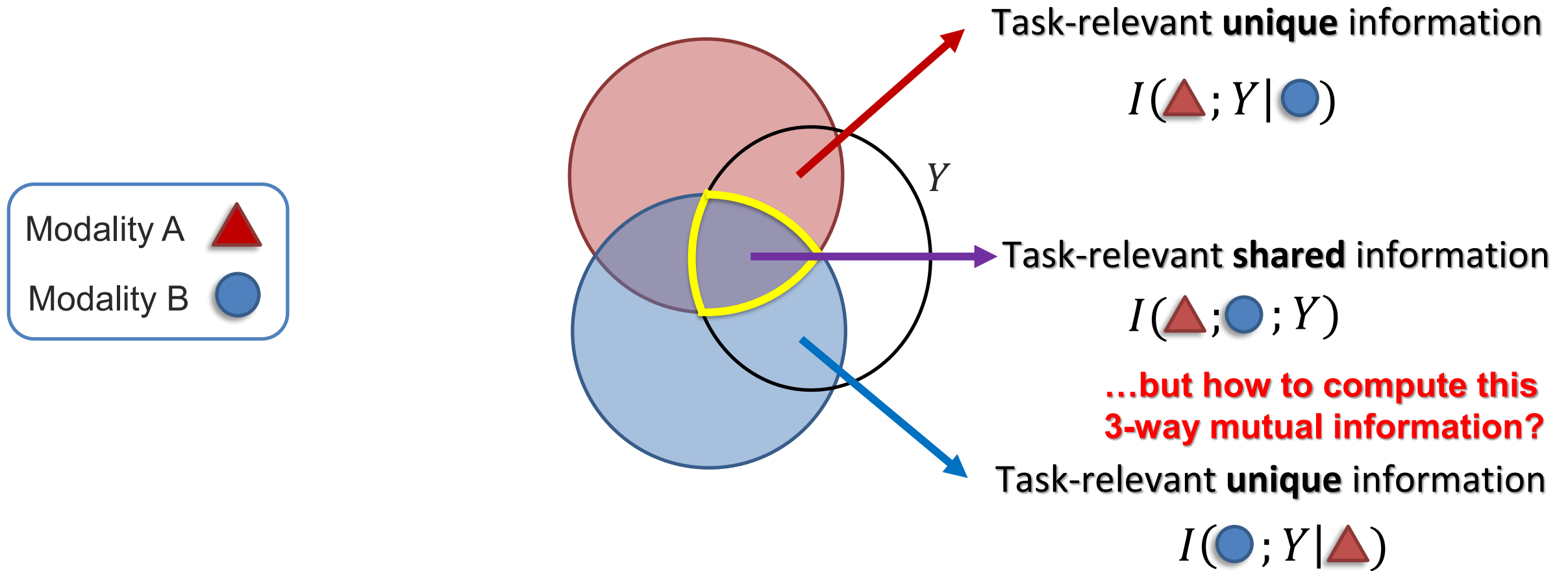
$$I(\triangle; Y | \bullet)$$



Modeling Multimodal Interactions



Modeling Multimodal Interactions



Modeling Multimodal Interactions

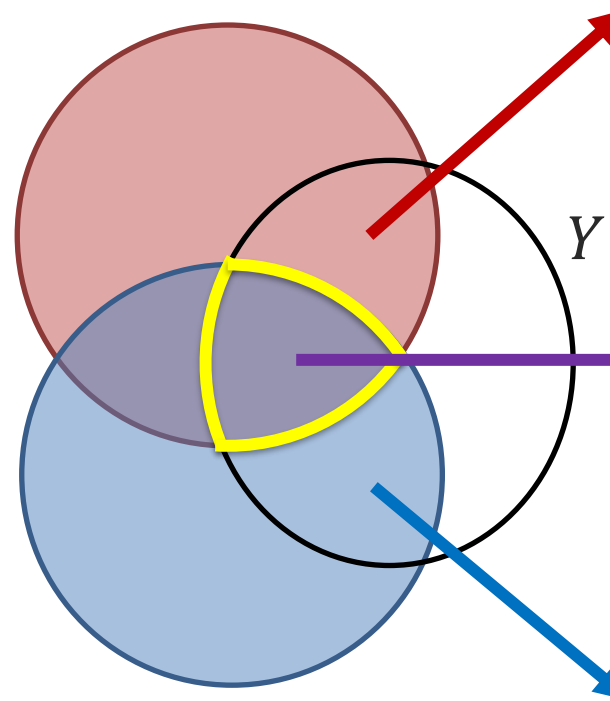
Partial Information Decomposition (PID)

$$I(\triangle; \bullet; Y) = R - S$$

factorizes 3-way mutual information into:

R: redundancy

S: Synergy



Task-relevant **unique** information

$$I(\triangle; Y | \bullet)$$

Task-relevant **shared** information

$$I(\triangle; \bullet; Y)$$

...but how to compute
3-way mutual information?

Task-relevant **unique** information

$$I(\bullet; Y | \triangle)$$

More details:

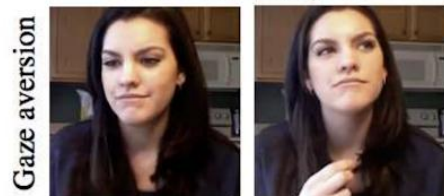
Liang et al., Quantifying & Modeling Feature Interactions: An Information Decomposition Framework. Neurips 2023

Quantifying Multimodal Interactions

These interactions can be efficiently estimated – gives a path towards understanding interactions

Language: *And he I don't think he got mad when hah
I don't know maybe.*

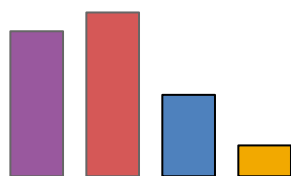
Vision:



Acoustic:

(frustrated voice)

Sentiment



$R \ U_\ell \ U_{av} \ S$

Language/Agreement

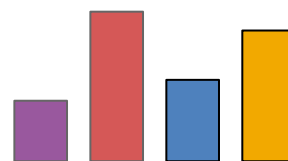
Sheldon :

Its just a *privilege* to watch your mind at work.

- **Text** : suggests a compliment.
- **Audio** : neutral tone.
- **Video** : straight face.

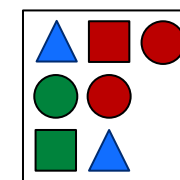


Sarcasm



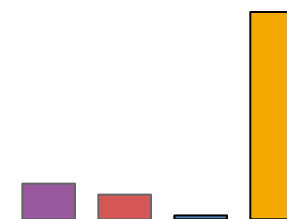
$R \ U_\ell \ U_{av} \ S$

Multimodal Transformer



Is there a red shape above a circle?

VQA



$R \ U_\ell \ U_i \ S$

Multiplicative/Transformer

Also matches human judgment of interactions, and other sanity checks on synthetic datasets
They can also be used to choose most appropriate models.

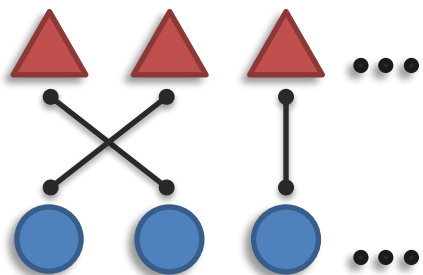
Challenge 2: Alignment

Challenge 2: Alignment

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

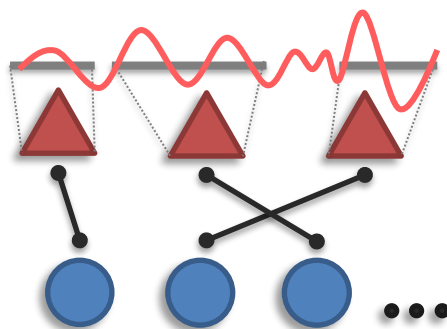
Sub-challenges:

Discrete Alignment



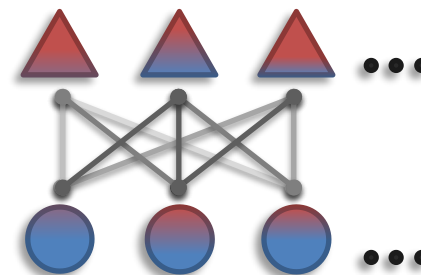
Discrete elements
and connections

Continuous Alignment



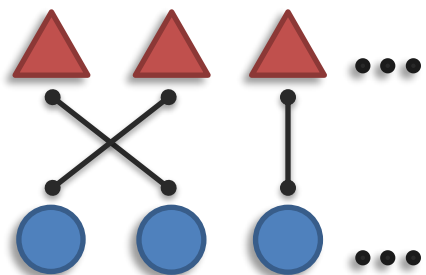
Segmentation and
continuous warping

Contextualized Representation

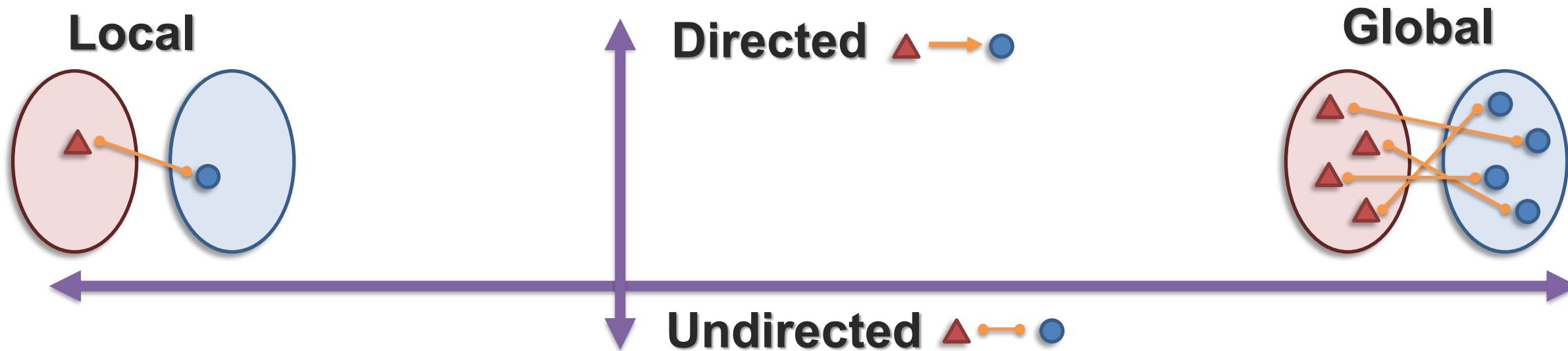


Alignment + representation

Sub-Challenge 2a: Discrete Alignment



Definition: Identify and model discrete connections between elements of multiple modalities



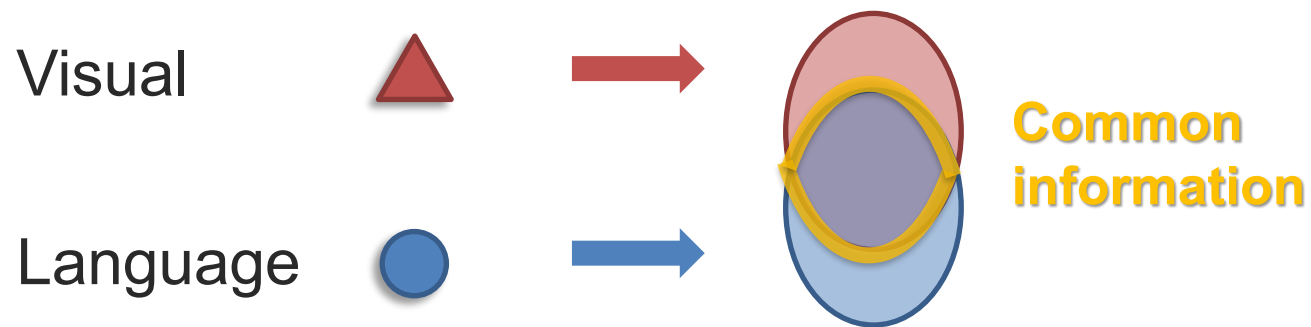
Language Grounding

Definition: Tying language (words, phrases,...) to non-linguistic elements, such as the visual world (objects, people, ...)



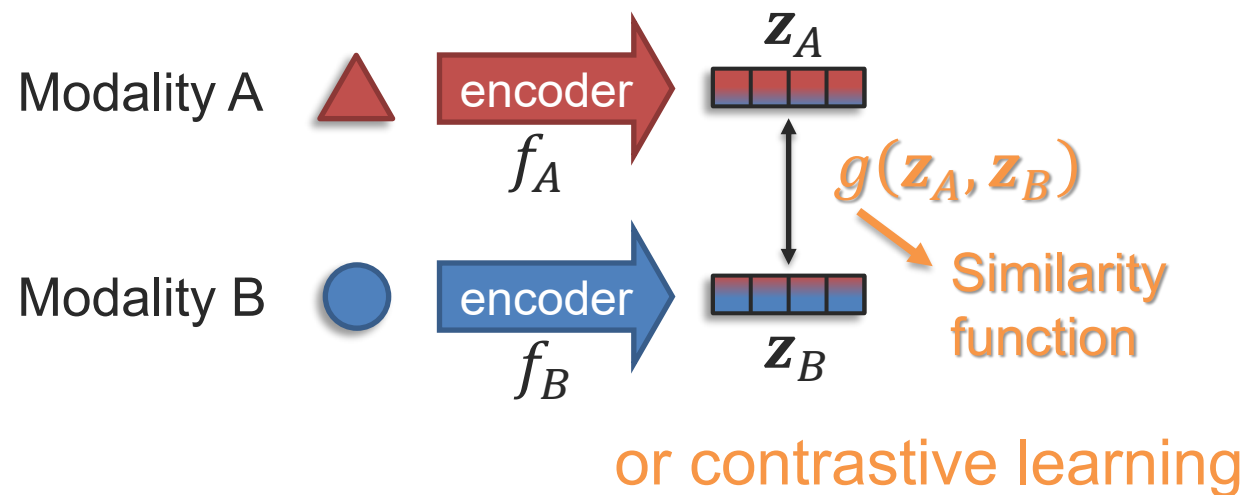
A **woman** reading **newspaper**

Local Alignment – Coordinated Representations

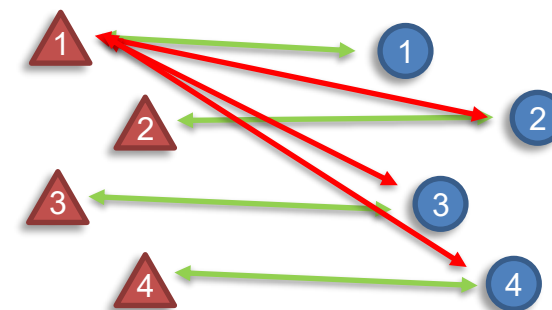


A **woman** reading **newspaper**

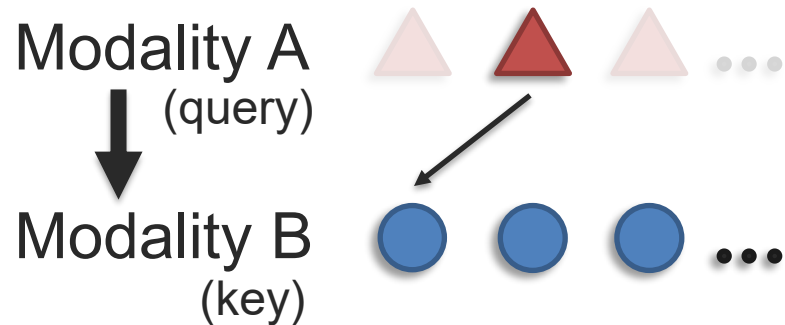
Learning coordinated representations:



Supervision: Paired data



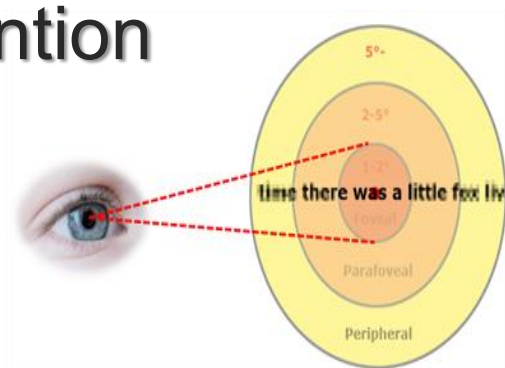
Directed Alignment



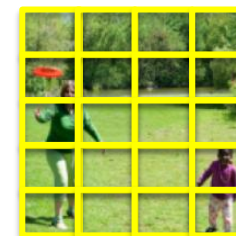
A woman is throwing a frisbee

Which
object?

Attention



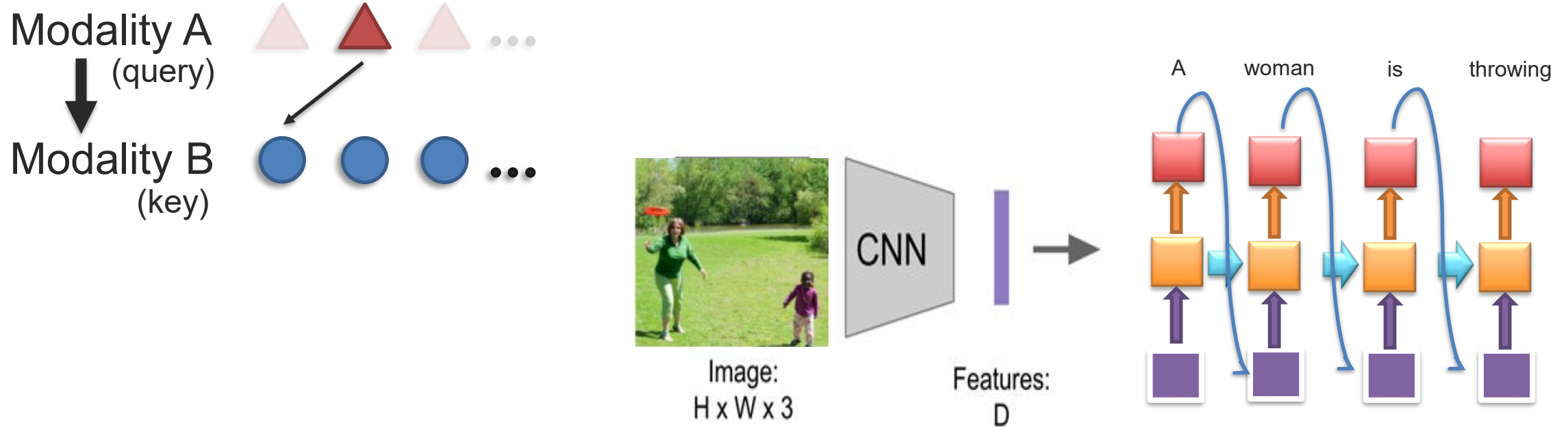
1 Soft attention



2 Hard attention

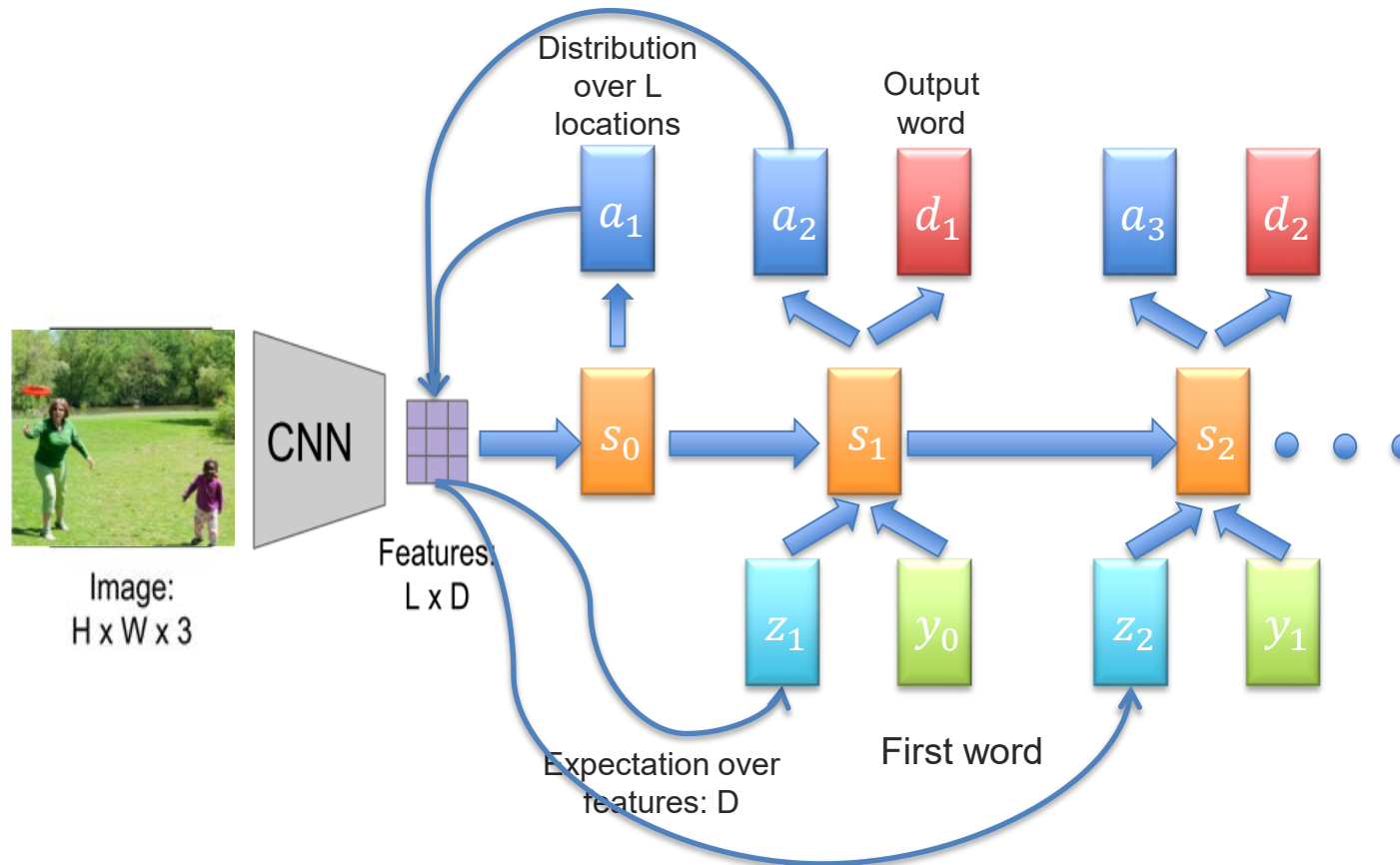


Directed Alignment – Image Captioning



Should we always use the final layer of the CNN for all generated words?

Directed Alignment – Image Captioning



Attention Gates

Before:

$$p(y_i | y_1, \dots, y_{i-1}, \mathbf{x}) = g(y_{i-1}, \mathbf{s}_i, \mathbf{z}),$$

where $\mathbf{z} = \mathbf{h}_T$, last encoder state and $\mathbf{s}_i$ is the current state of the decoder

Now:

$$p(y_i | y_1, \dots, y_{i-1}, \mathbf{x}) = g(y_{i-1}, \mathbf{s}_i, \mathbf{z}_i)$$

Have an attention “gate”

- A different context $\mathbf{z}_i$ used at each time step!

- $\mathbf{z}_i = \sum_{j=1}^{T_x} \alpha_{ij} \mathbf{h}_j$

α_{ij} is the (scalar) attention for word j at generation step i

Attention Gates

So how do we determine α_{ij} ?

$$\alpha_{i,j} = \frac{\exp(e_{ij})}{\sum_{k=1}^{T_x} \exp(e_{ik})} \quad \Rightarrow \text{softmax, making sure they sum to 1}$$

where:

$$e_{ij} = \mathbf{v}^T \sigma(W \mathbf{s}_{i-1} + U \mathbf{h}_j)$$

a feedforward network that can tell us how important the current encoding is

$\mathbf{v}, W, U$ — learnable weights

$$\mathbf{z}_i = \sum_{j=1}^{T_x} \alpha_{ij} \mathbf{h}_j$$

← expectation of the context (a fancy way to say it's a weighted average)

Example – Image Captioning



[Show, Attend and Tell: Neural Image Caption Generation with Visual Attention, Xu et al., 2015]

Global Alignment

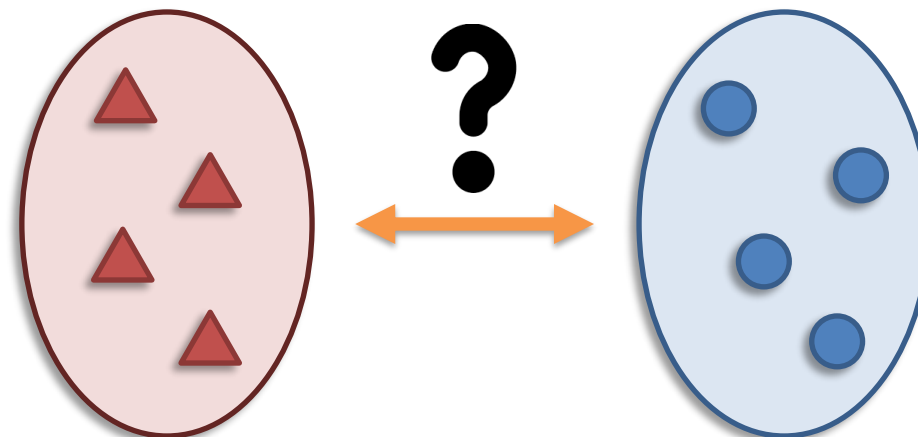
Visual



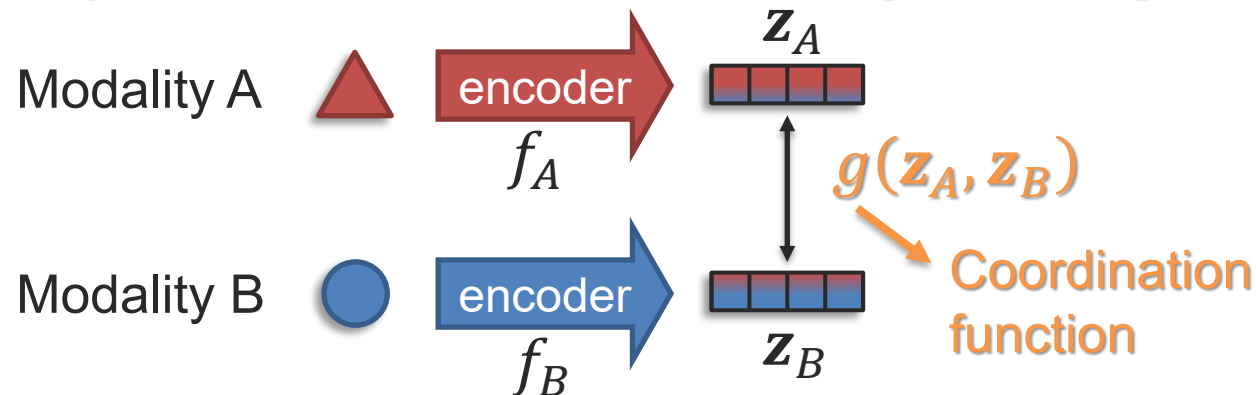
Language



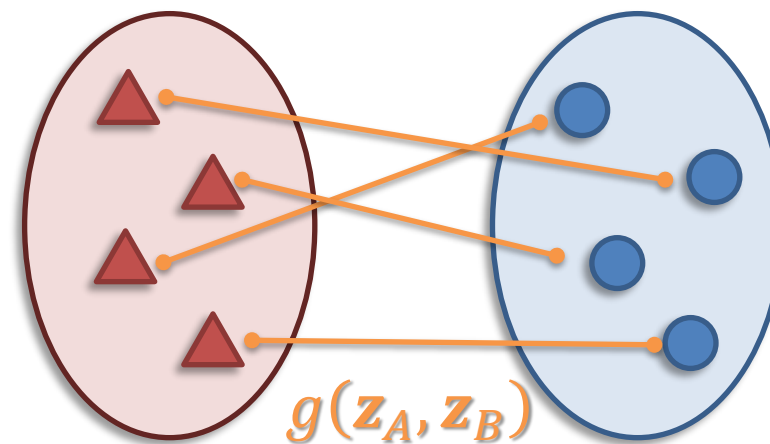
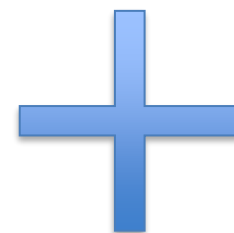
Latent pairing information



Jointly optimize representation + global alignment:

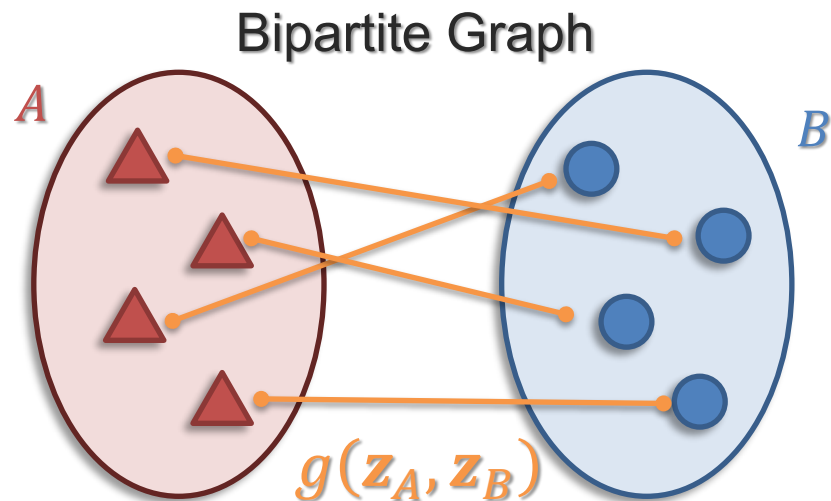


(representation)



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

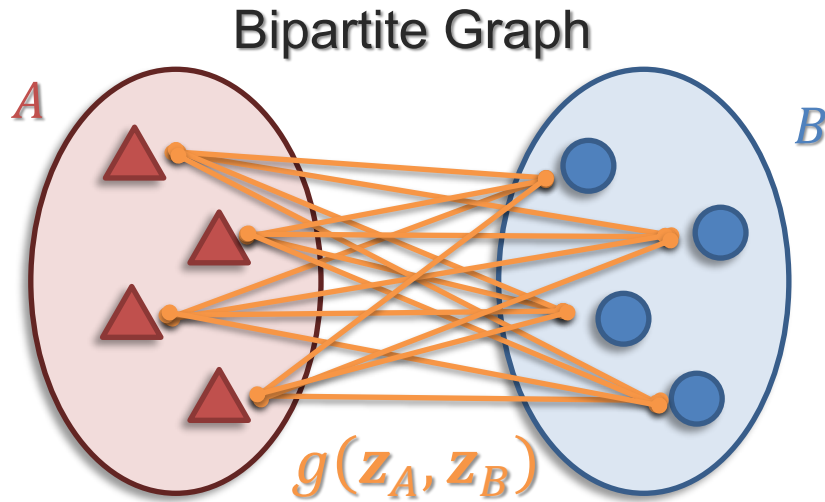
$w_{(i,j)} = g(\mathbf{z}_A^i, \mathbf{z}_B^j)$

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 $\mathbf{z}_B^j$

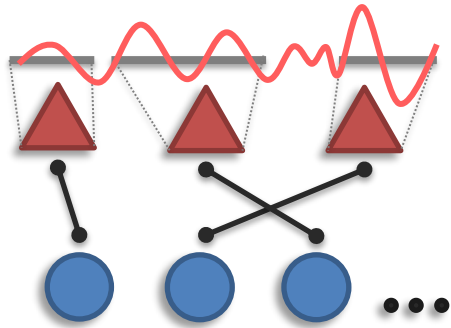
Similarity weights: $w_{(i,j)} = g(\mathbf{z}_A^i, \mathbf{z}_B^j)$

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

➡ Wassertein distance gives optimal transport

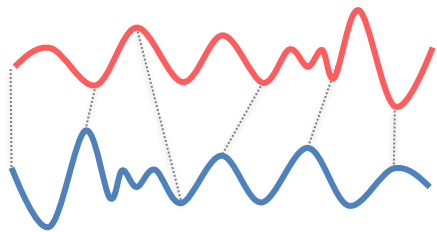
Continuous Alignment

Challenge 2b: Continuous Alignment

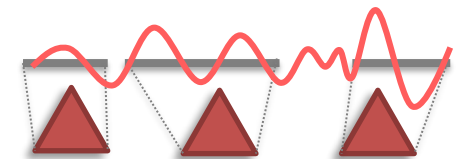


Definition: Model alignment between modalities with continuous signals and no explicit elements

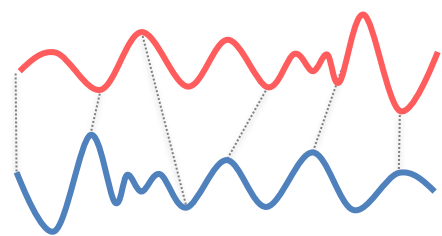
Continuous
warping



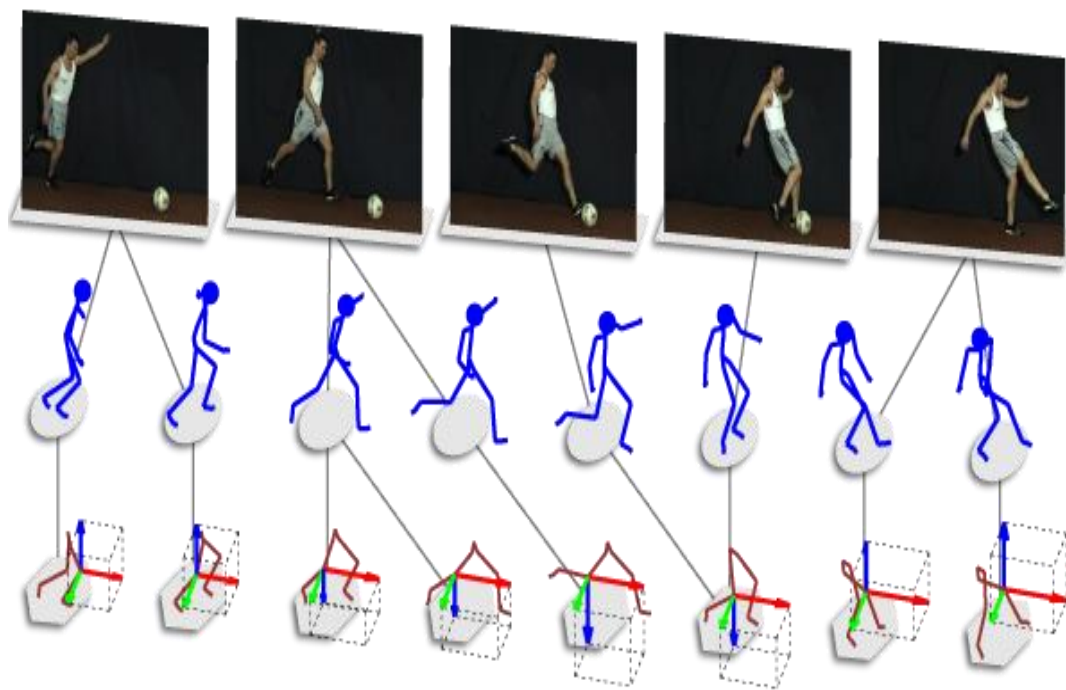
Discretization
(segmentation)



Continuous Warping – Example



➔ Aligning video sequences



Dynamic Time Warping (DTW)

We have two unaligned temporal unimodal signals

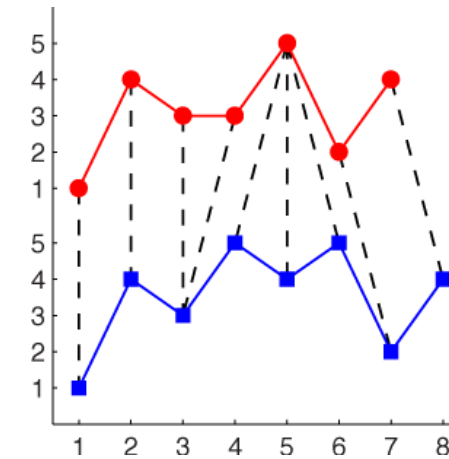
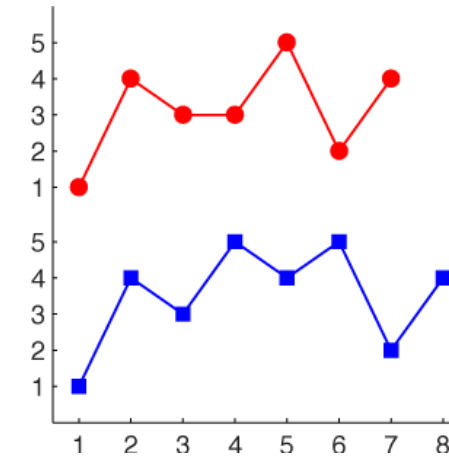
- $\mathbf{X} = [x_1, x_2, \dots, x_{n_x}] \in \mathbb{R}^{d \times n_x}$
- $\mathbf{Y} = [y_1, y_2, \dots, y_{n_y}] \in \mathbb{R}^{d \times n_y}$

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!



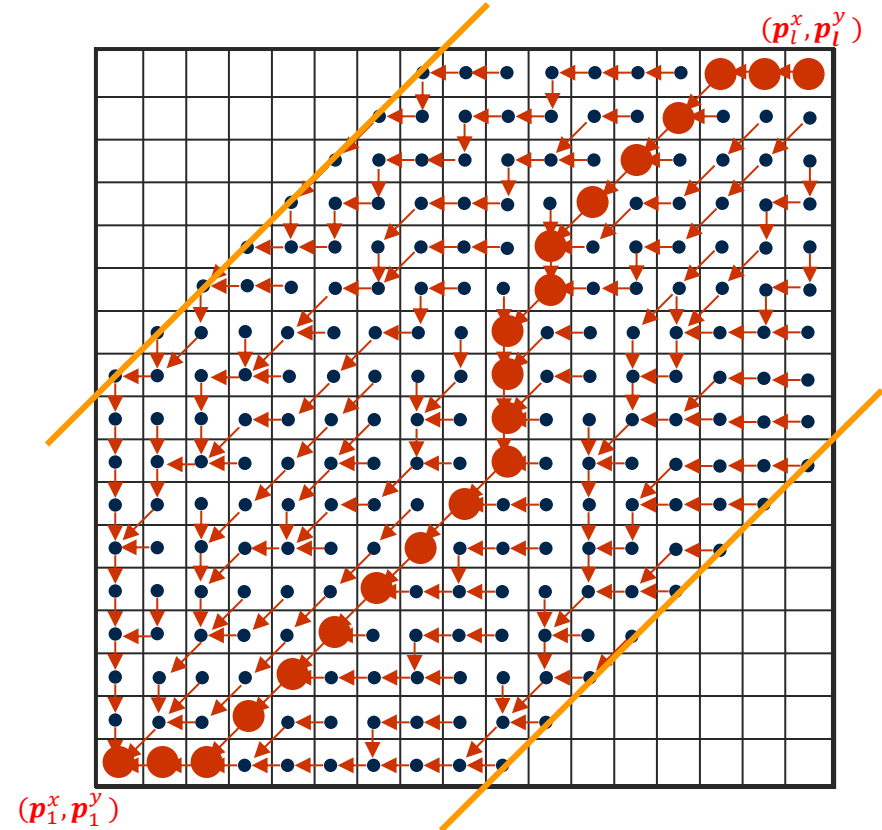
Dynamic Time Warping (DTW)

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

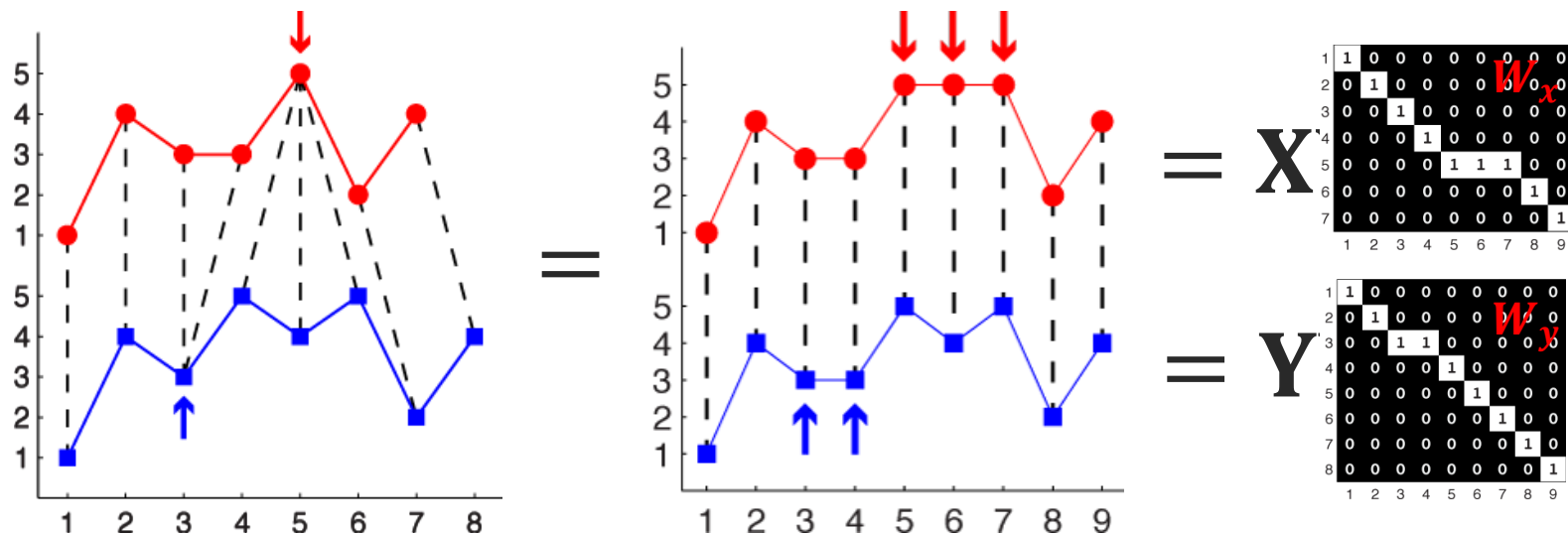
Solved using dynamic programming while respecting the restrictions



DTW alternative formulation

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

Replication doesn't change the objective!



Alternative objective:

$$L(W_x, W_y) = \|\mathbf{X}W_x - \mathbf{Y}W_y\|_F^2$$

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)

W_x, W_y – alignment matrices

A differentiable version of DTW also exists...

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

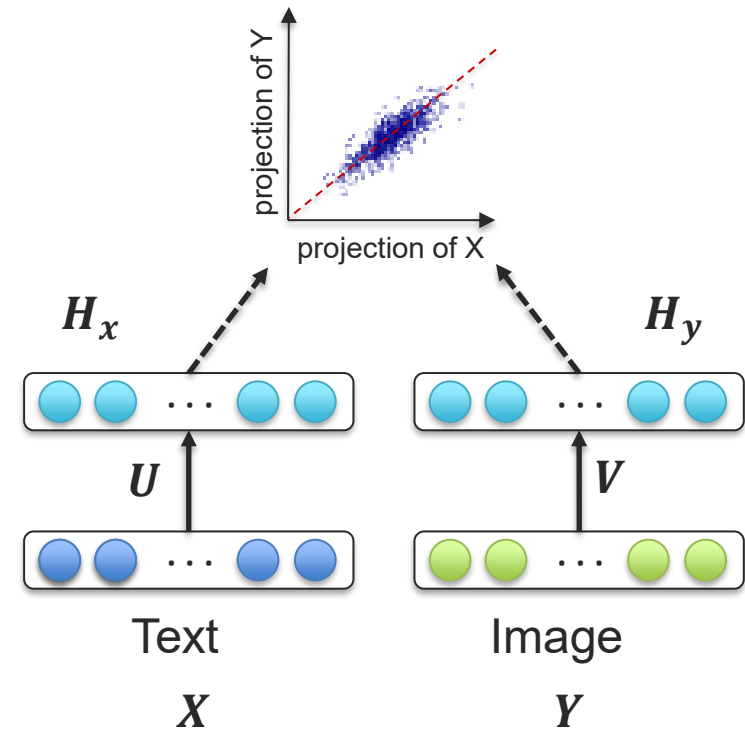
Canonical Correlation Analysis – Reminder

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

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

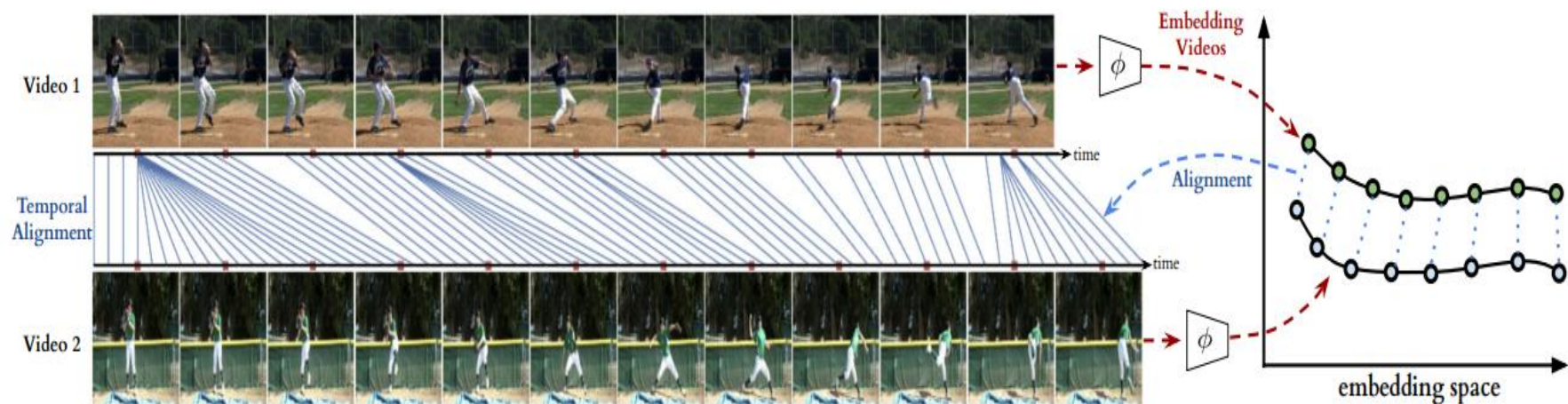
Allows to align multi-modal or multi-view (same modality but from a different point of view)

- $\mathbf{W}_x, \mathbf{W}_y$ – 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

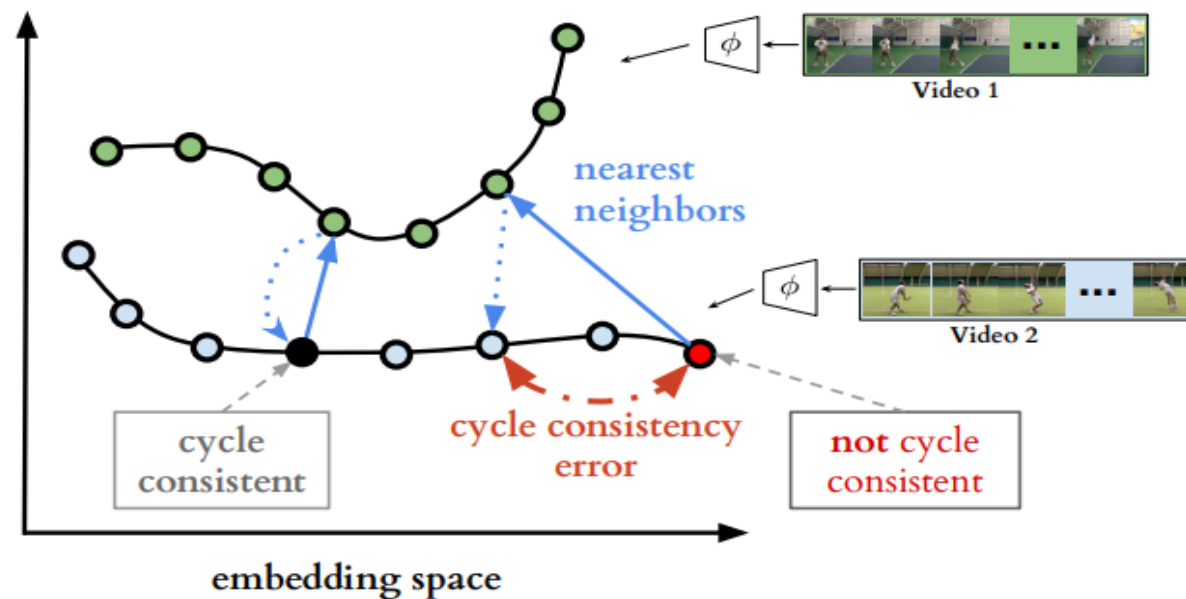
Premise: we have paired video sequences that can be temporally aligned



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

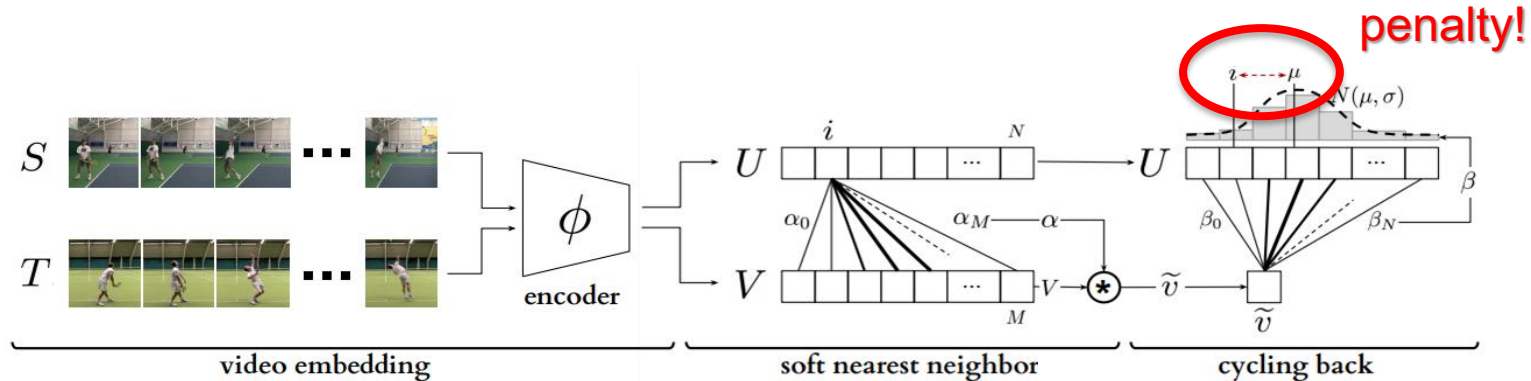
Temporal Cycle-Consistency Learning

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



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

Temporal Cycle-Consistency Learning



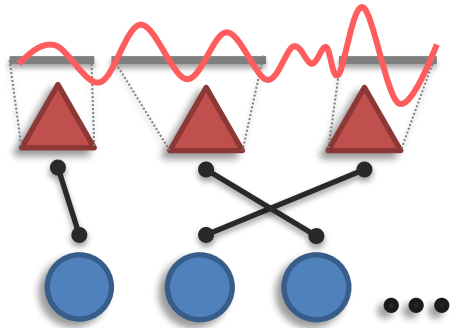
Compute “soft” / “weighted” nearest neighbour:

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:

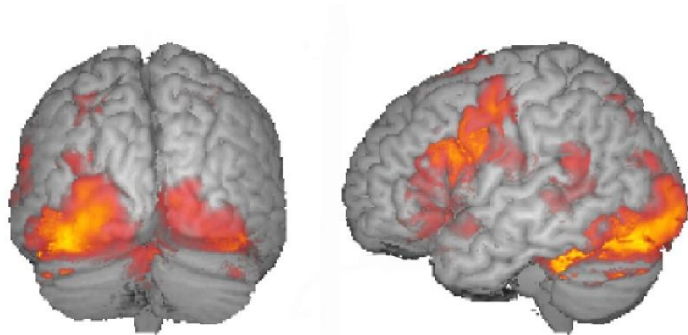
$$\beta_k = \frac{e^{-||\tilde{v} - u_k||^2}}{\sum_j^N e^{-||\tilde{v} - u_j||^2}} \quad L_{cbr} = \frac{|i - \mu|^2}{\sigma^2} + \lambda \log(\sigma)$$

Discretization (aka Segmentation)

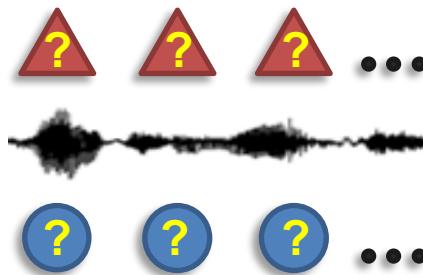


Common assumptions: ① Segmented elements

Examples:



Medical imaging



Signals

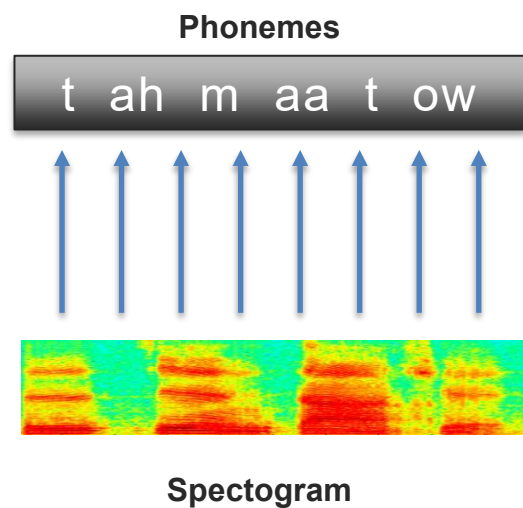


Images

objects

Discretization – Example

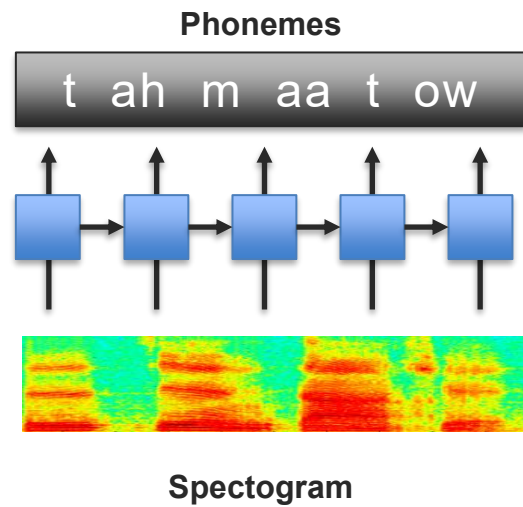
Sequence Labeling and Alignment



How can we predict the sequence
of phoneme labels?

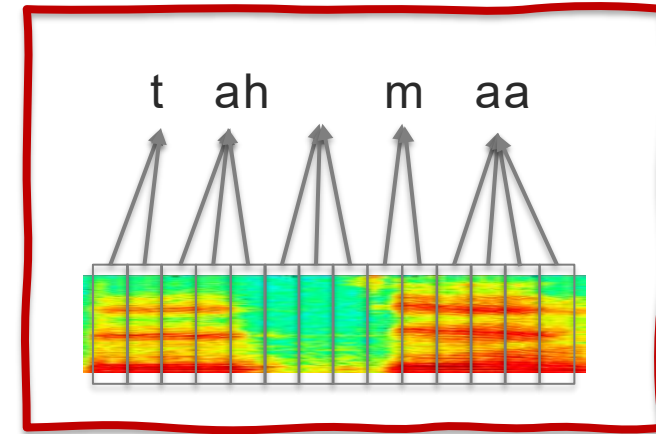
Discretization – Example

Sequence Labeling and Alignment



How can we predict the sequence of phoneme labels?

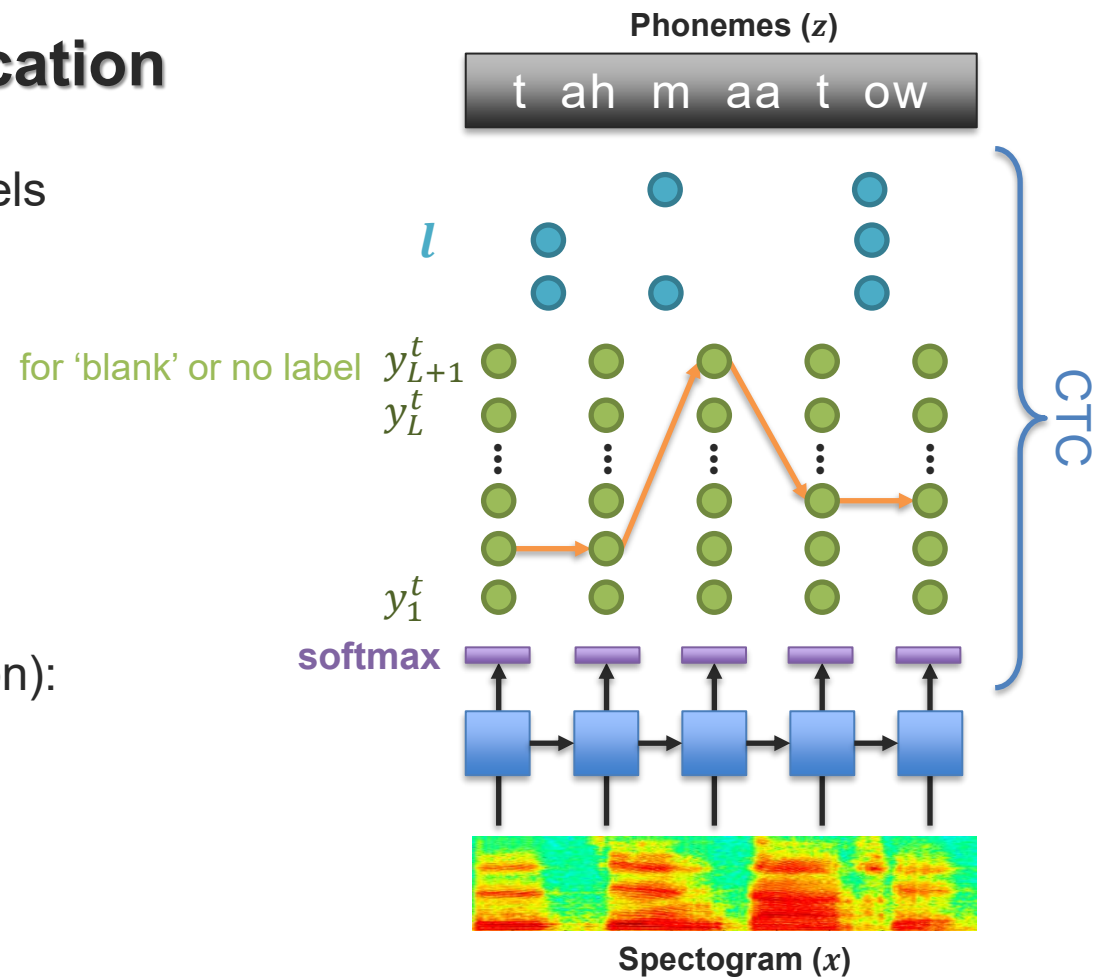
Challenge: many-to-1 alignment



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

