



Language
Technologies
Institute

Carnegie
Mellon
University

Multimodal Machine Learning

Lecture 3.2: Multimodal Coordination and Fission

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.*

Upcoming Course Assignments

Week 3 Reading Assignment (due Friday 9/11 at 8pm ET)

- Two papers: a first about unimodal and a second about multimodal

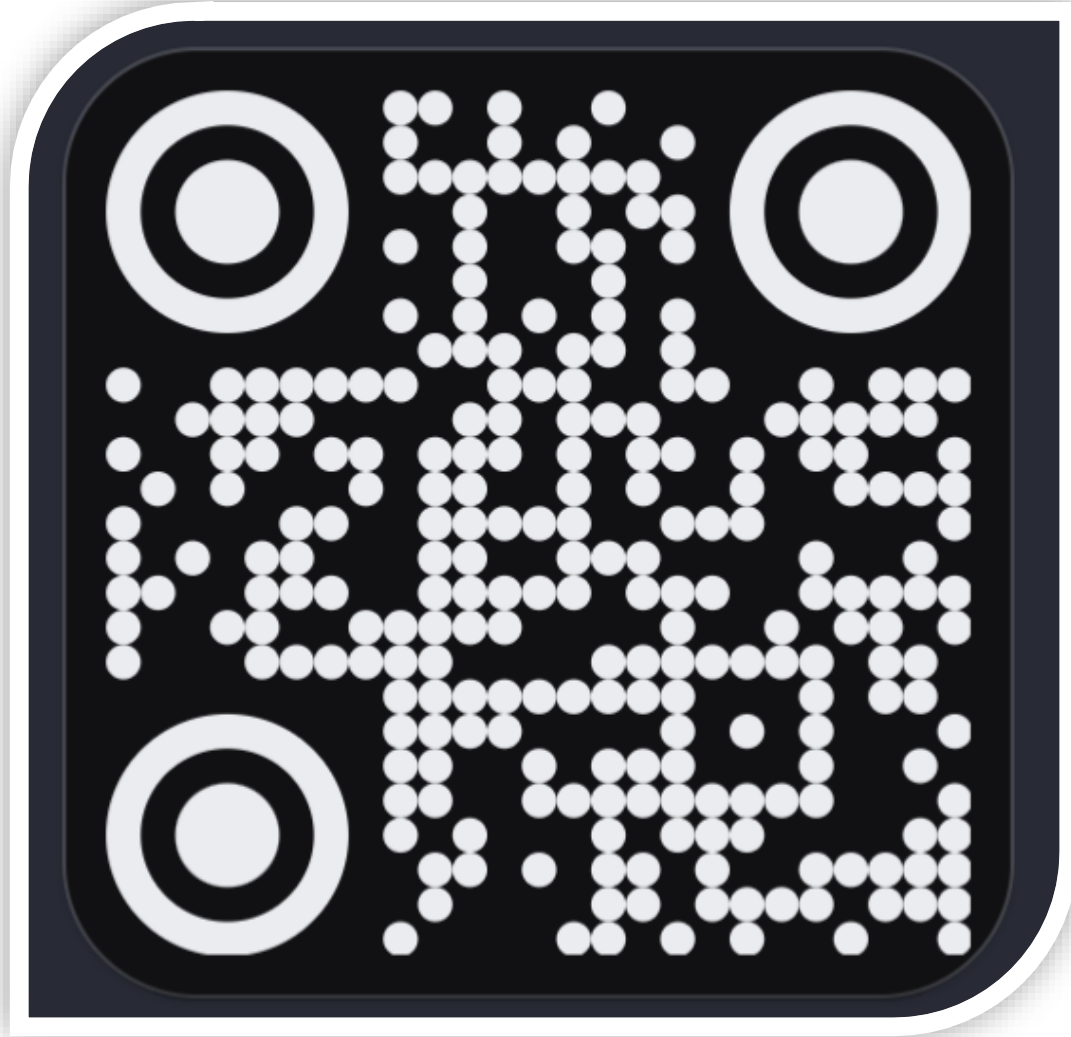
Proposal, aka First Assignment (due Sunday 9/20 at 8pm ET)

- Literature review of your research problem
- Second revision of your research ideas (i.e., top-down approach)

Coming soon: Proposal presentations

- No lectures on Tuesday 9/22 and Thursday 9/24
- Each team will schedule a 15mins timeslot
- 5 mins presentations + 10 mins discussion

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 3.2: Multimodal Coordination and Fission

Louis-Philippe Morency

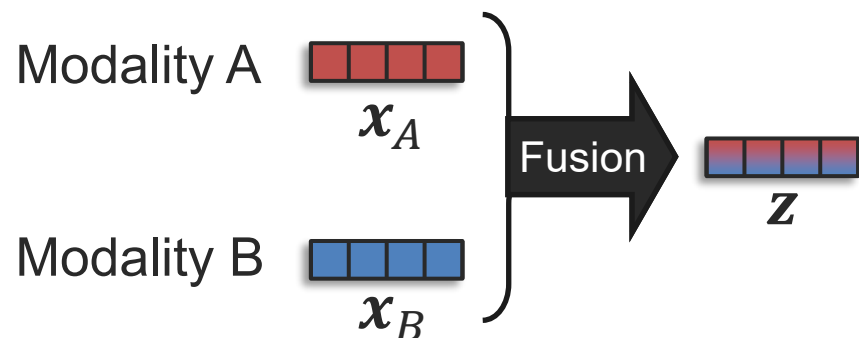
** Original course co-developed with Tadas Baltrusaitis.
Major revision co-developed with Paul Liang. Co-lecturers
included Yonatan Bisk, Daniel Fried and Carlos Busso.*

Objectives of today's class

- Representation fusion
 - Additive and multiplicative fusion
 - Gated fusion
- Representation coordination
 - Coordination functions
 - Kernel similarity functions
 - Canonical correlation analysis
 - Contrastive learning
 - Information, entropy and mutual information
- Representation fission
 - Factorized multimodal representations
 - Clustering and fine-grained fission

Representation Fusion

Additive Fusion

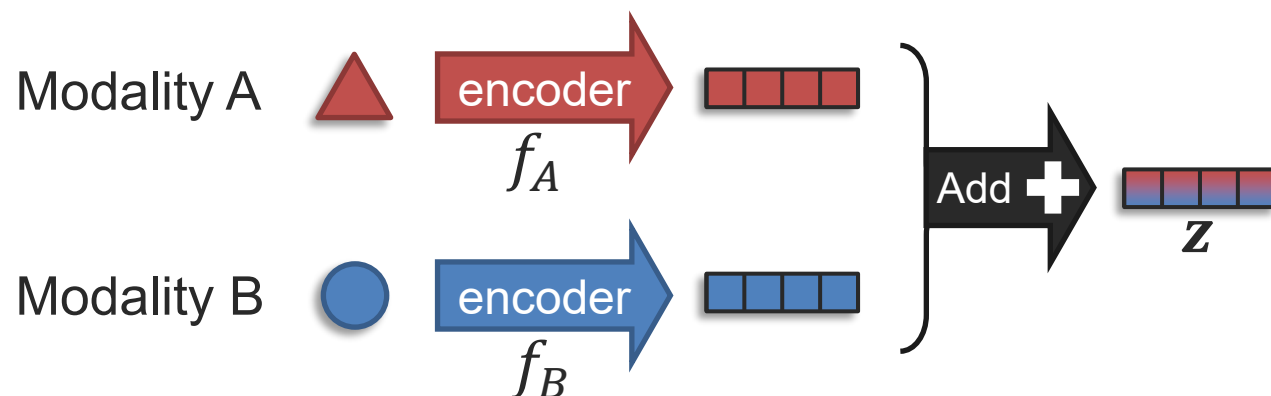


Additive fusion:

$$z = w_1 x_A + w_2 x_B$$

→ 1-layer neural network
can be seen as additive

With unimodal encoders:

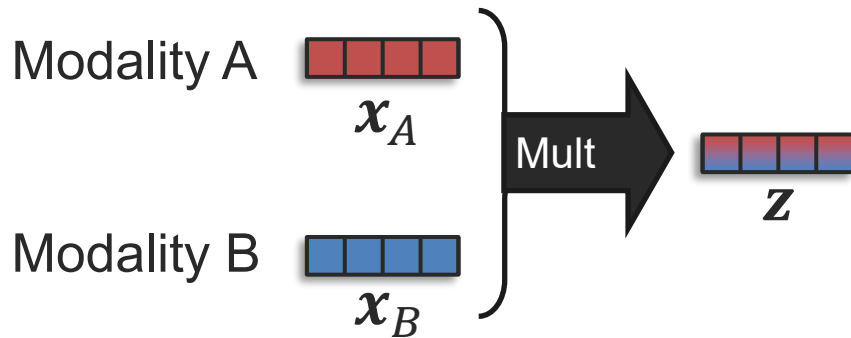


Additive fusion:

$$z = f_A(\triangle) + f_B(\circ)$$

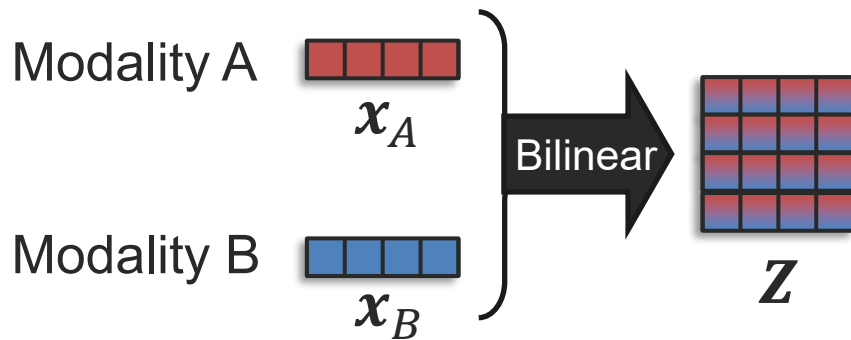
→ It could be seen as an
ensemble approach
(late fusion)

Multiplicative Fusion



Simple multiplicative fusion:

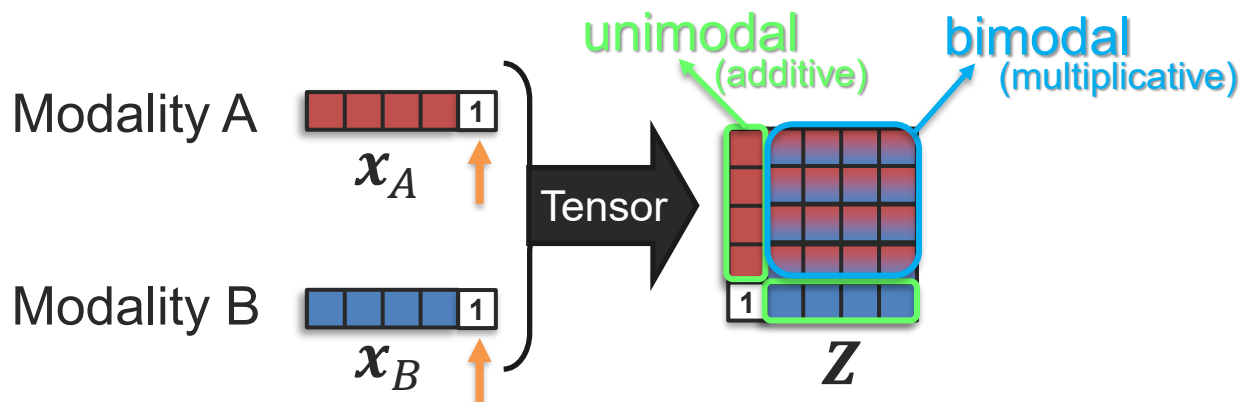
$$z = w(x_A \times x_B)$$



Bilinear Fusion:

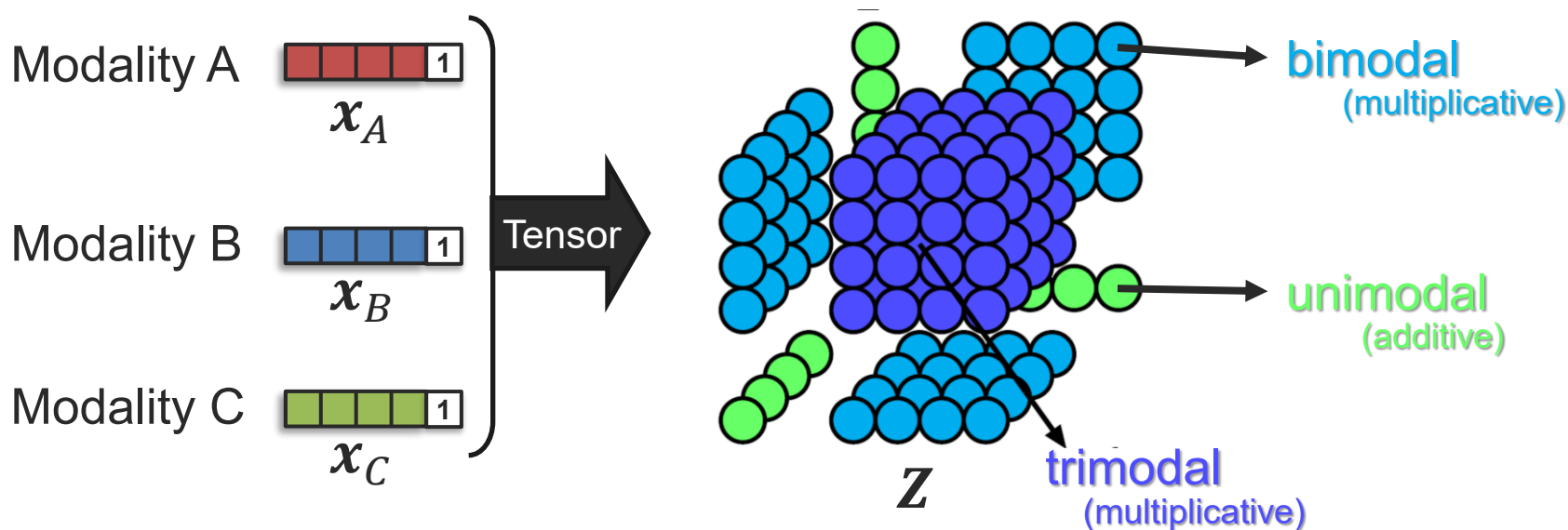
$$Z = W(x_A^T \cdot x_B)$$

Tensor Fusion



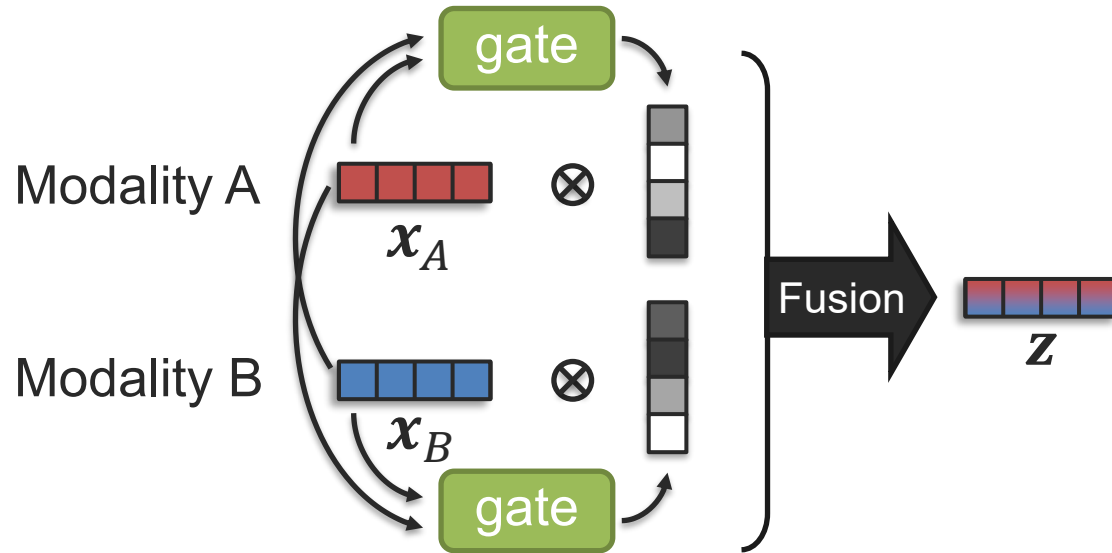
Tensor Fusion (bimodal):

$$\mathbf{Z} = \mathbf{w}([\mathbf{x}_A \ 1]^T \cdot [\mathbf{x}_B \ 1])$$



... but the weight matrix may end up quite large!

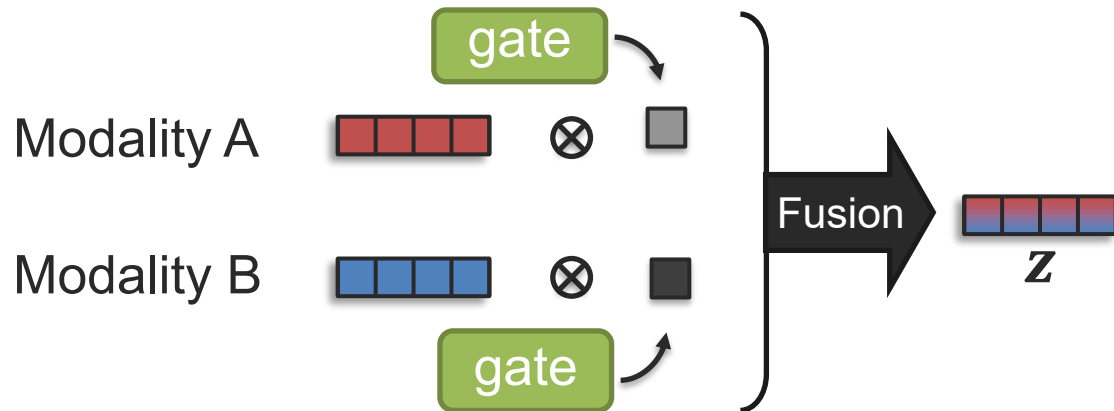
Gated Fusion



Example with additive fusion:

$$z = g_A(x_A, x_B) \cdot x_A + g_B(x_A, x_B) \cdot x_B$$

→ g_A and g_B can be seen as attention functions



→ Gating output can be one weight for the whole modality

Gating Module (aka, attention module)

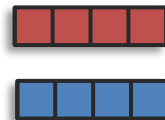


What should it be?

Target modality []

Other modality []

All modality []



“Neural network designed to mask unwanted signal from propagating forward” (gating)

...or with a more positive view:

“Neural network designed to select preferable signal to move forward” (attention)

Soft attention



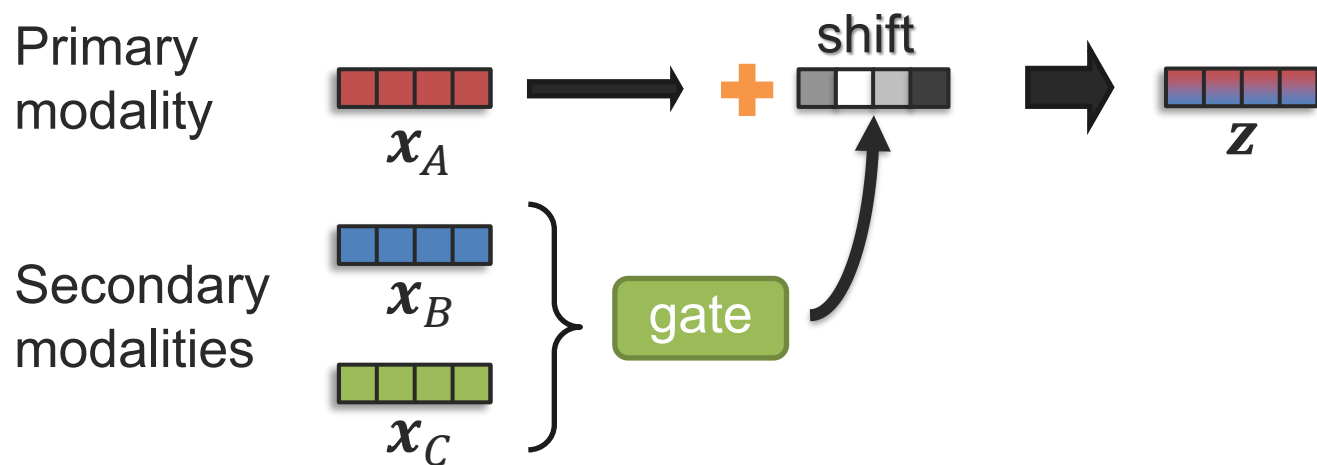
Easier to compute derivative (gradient)

Hard attention



Derivative is harder (e.g., use reinforcement learning)

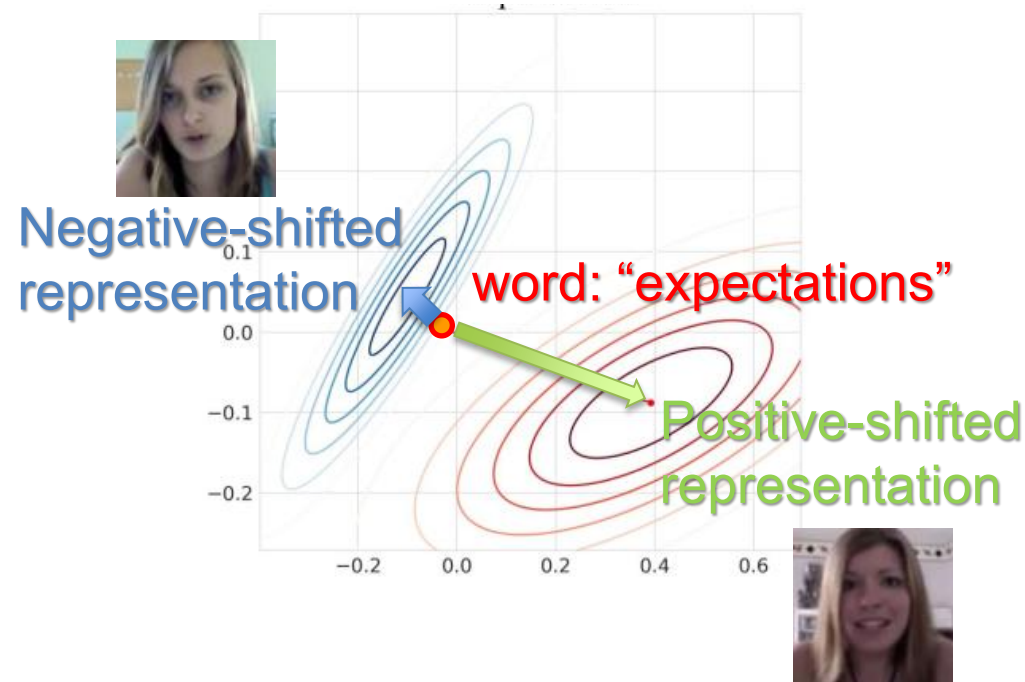
Modality-Shifting Fusion



Example with language modality:

Primary modality: language

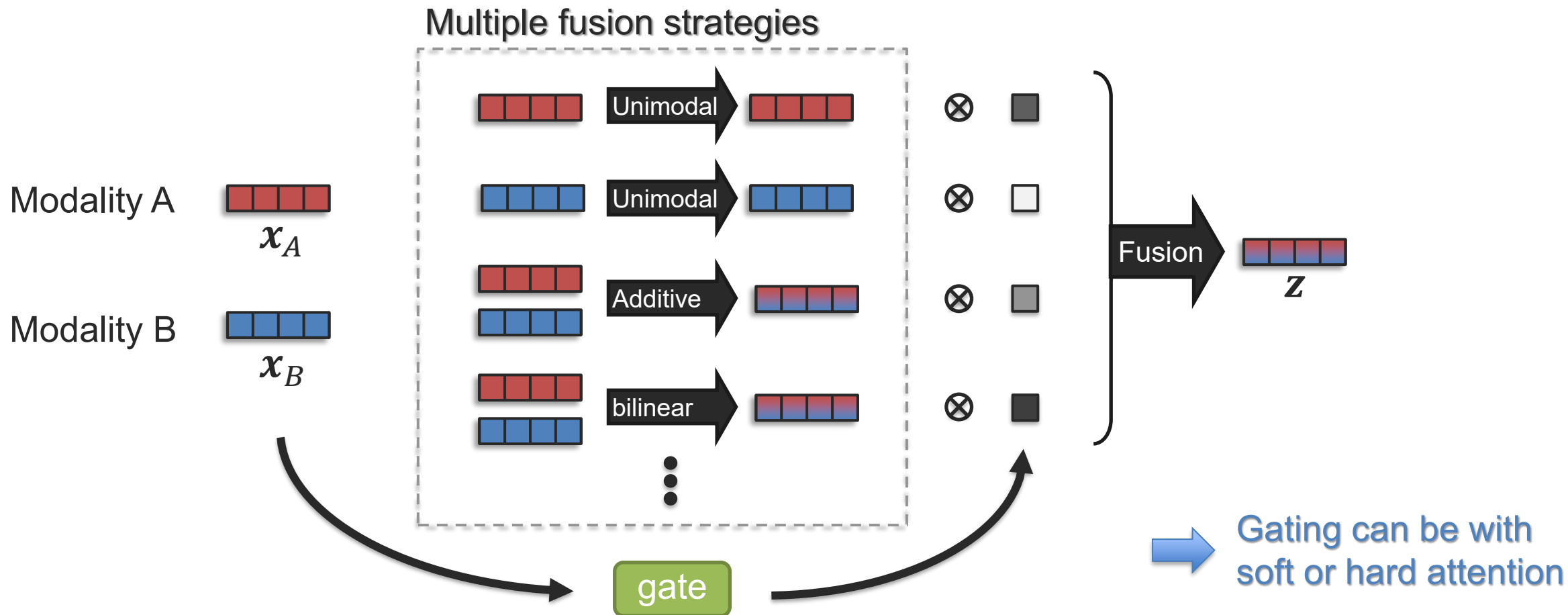
Secondary modalities: acoustic and visual



[Wang et al., Words Can Shift: Dynamically Adjusting Word Representations Using Nonverbal Behaviors, AAAI 2019]

[Rahman et al., Integrating Multimodal Information in Large Pretrained Transformers, ACL 2020]

Mixture of Fusions



[Zadeh et al., Multimodal Language Analysis in the Wild: CMU-MOSEI Dataset and Interpretable Dynamic Fusion Graph, ACL 2018]

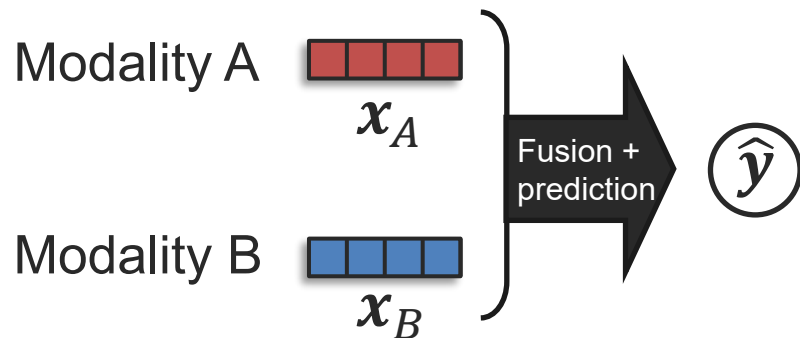
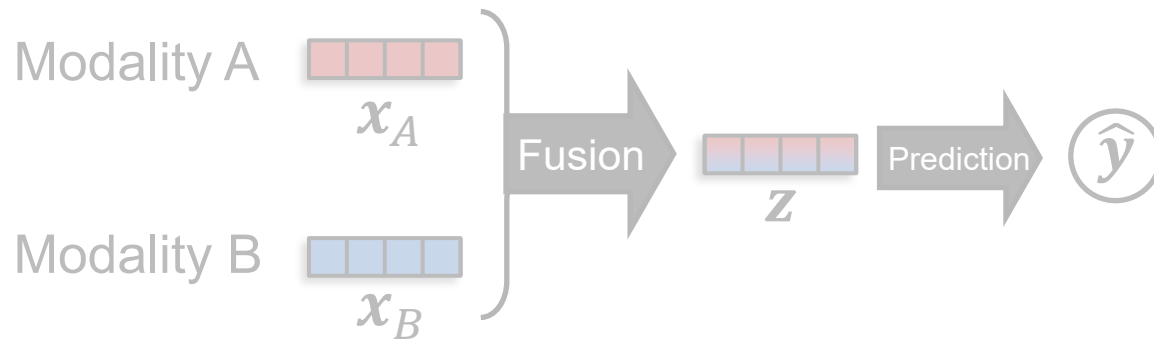
[Xu et al., MUFASA: Multimodal Fusion Architecture Search for Electronic Health Records, AAAI 2021]

Nonlinear Fusion

Nonlinear fusion:

$$\hat{y} = f(x_A, x_B) \in \mathbb{R}^d$$

where f could be a multi-layer perceptron or any nonlinear model

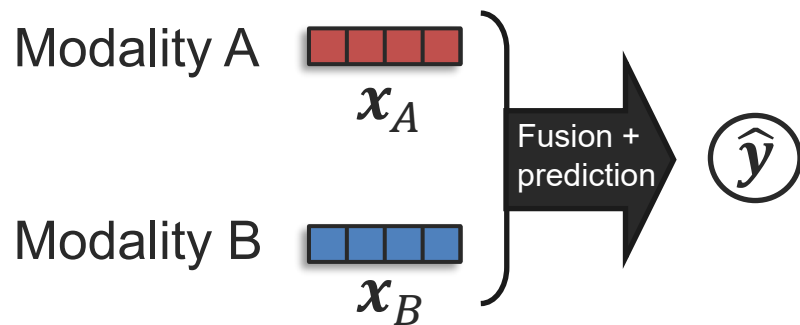


➡ This could be seen as *early fusion*:

$$\hat{y} = f([x_A, x_B])$$

... but will our neural network learn the nonlinear interactions?

Measuring Non-Additive Interactions



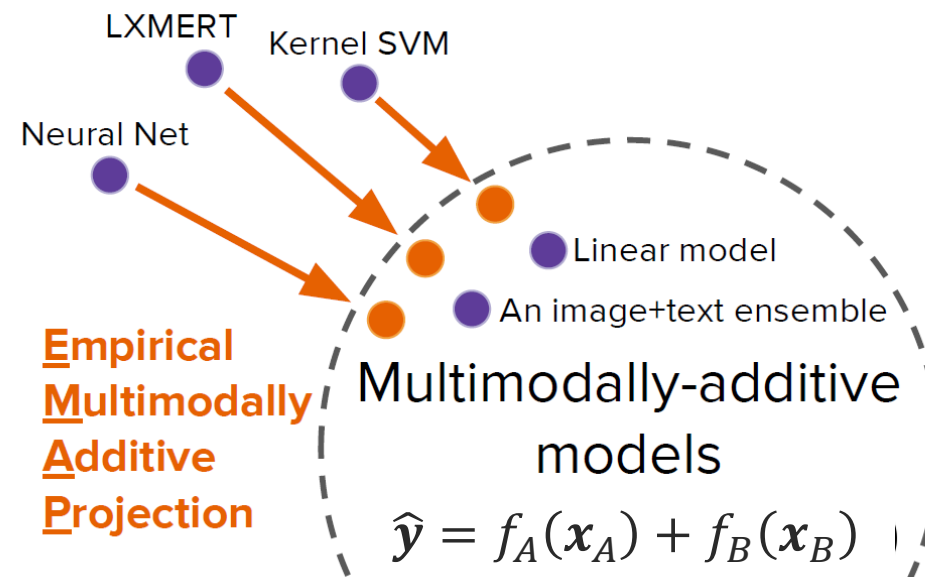
Nonlinear fusion:

$$\hat{y} = f(x_A, x_B)$$

Projection?

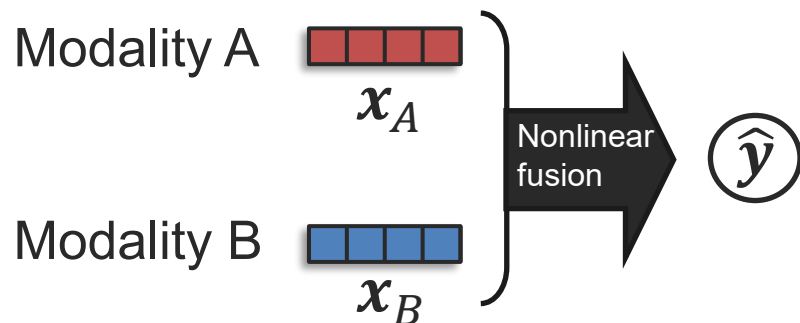
Additive fusion:

$$\hat{y} = f_A(x_A) + f_B(x_B)$$



Hessel and Lee, Does my multimodal model learn cross-modal interactions? It's harder to tell than you might think!, EMNLP 2020 → introduced the EMAP method

Measuring Non-Additive Interactions



Nonlinear fusion:

$$\hat{\mathbf{y}} = f(\mathbf{x}_A, \mathbf{x}_B)$$

Projection?

Additive fusion:

$$\hat{\mathbf{y}}' = f_A(\mathbf{x}_A) + f_B(\mathbf{x}_B)$$

Projection from nonlinear to additive (using EMAP):

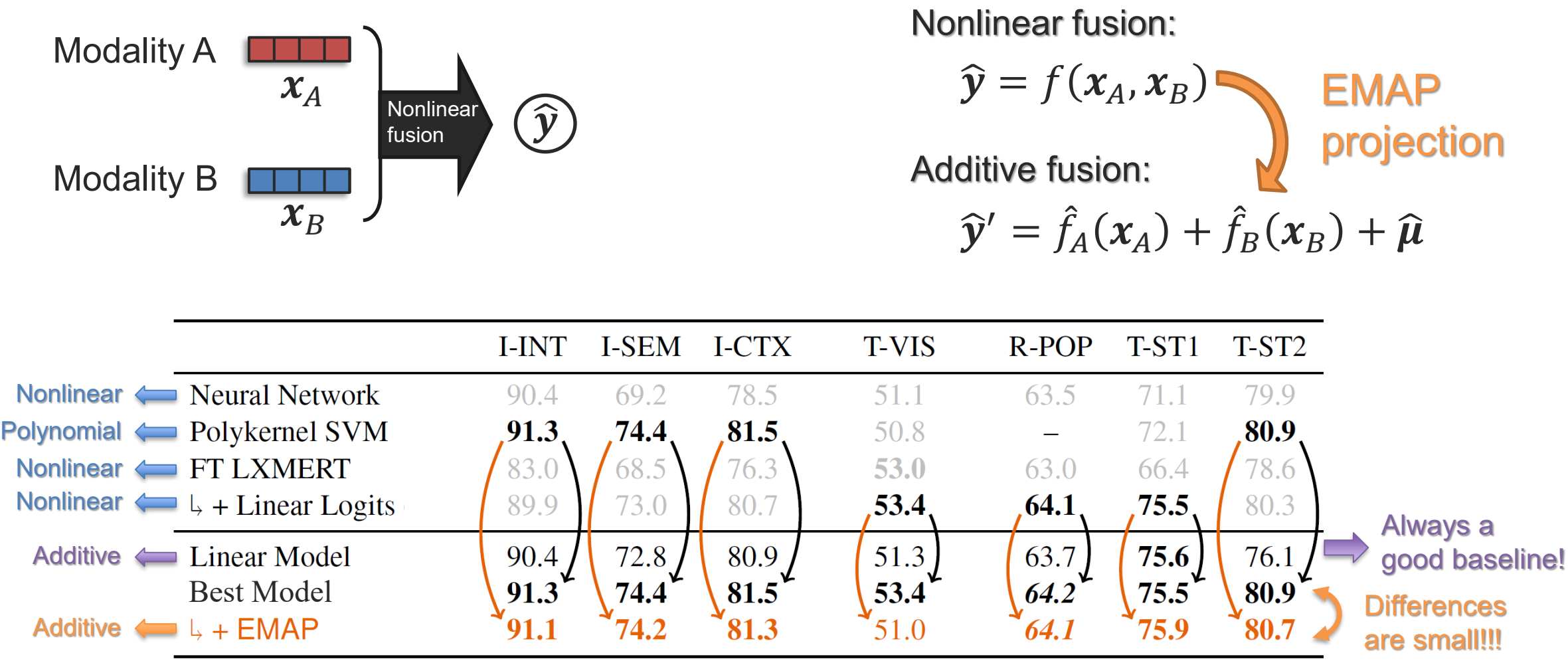
$$\tilde{f}(\mathbf{x}_A, \mathbf{x}_B) = \underbrace{\mathbb{E}_{\mathbf{x}_B} [f(\mathbf{x}_A, \mathbf{x}_B)]}_{f_A(\mathbf{x}_A)} + \underbrace{\mathbb{E}_{\mathbf{x}_A} [f(\mathbf{x}_A, \mathbf{x}_B)]}_{f_B(\mathbf{x}_B)}$$

Modality A Modality B

Additive fusion
(approximation)

[Hessel and Lee, Does my multimodal model learn cross-modal interactions? It's harder to tell than you might think!, EMNLP 2020]

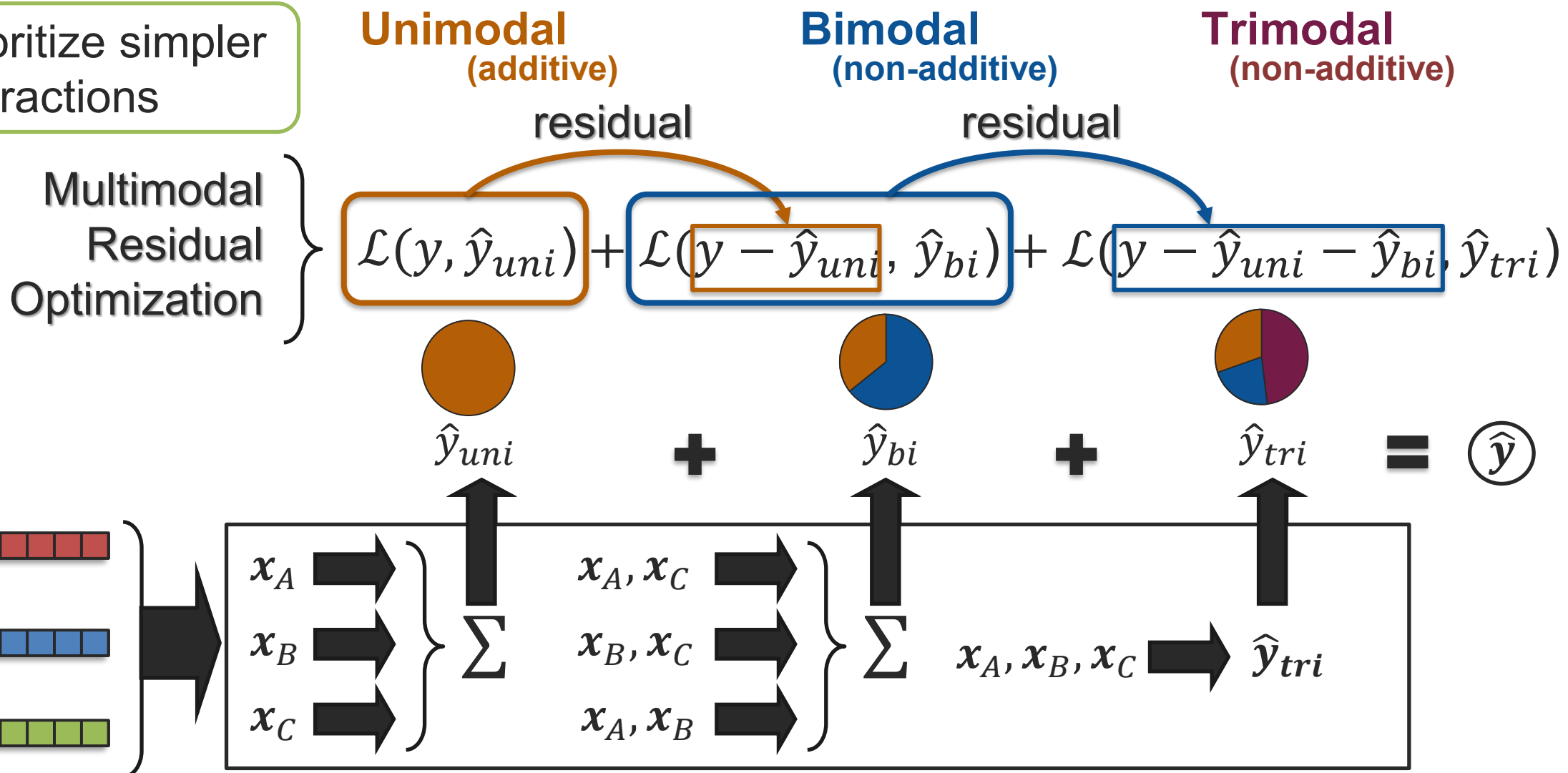
Measuring Non-Additive Interactions



[Hessel and Lee, Does my multimodal model learn cross-modal interactions? It's harder to tell than you might think!, EMNLP 2020]

Learning Non-additive Bimodal and Trimodal Interactions

Idea: prioritize simpler interactions



[Wortwein et al., Beyond Additive Fusion: Learning Non-Additive Multimodal Interactions, Findings-EMNLP 2022]

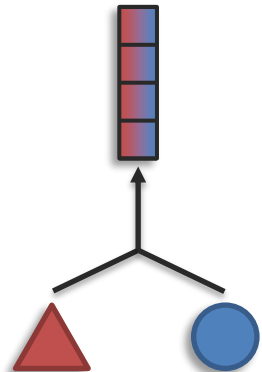
Representation Coordination

Challenge 1: Representation

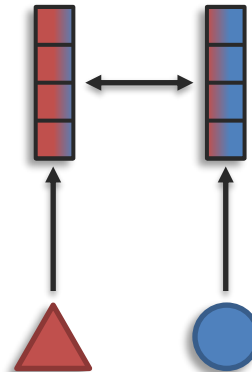
Definition: Learning representations that reflect cross-modal interactions between individual elements, across different modalities

Sub-challenges:

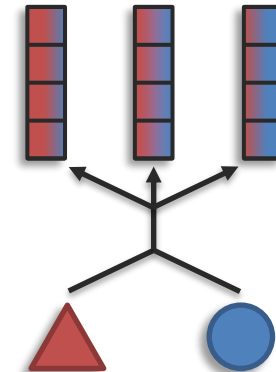
Fusion



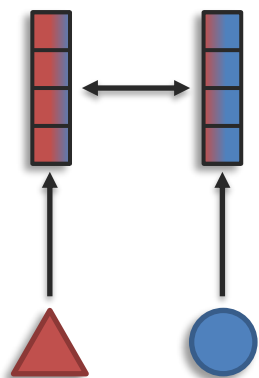
Coordination



Fission

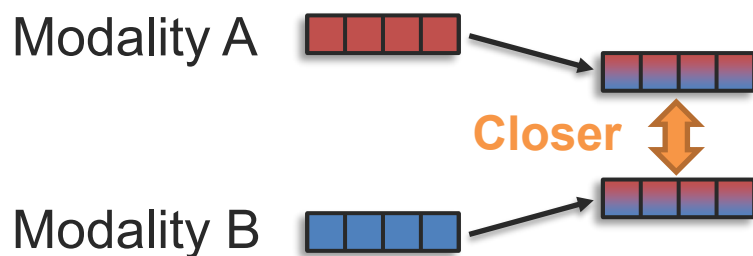


Sub-Challenge 1b: Representation Coordination

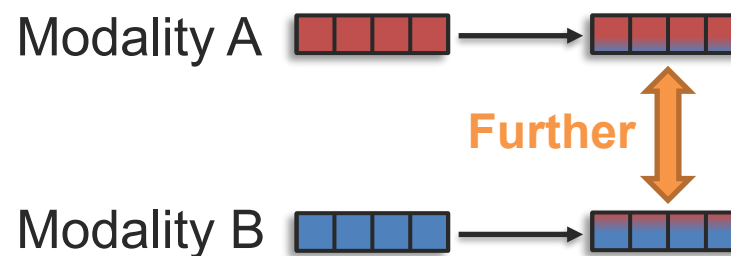


Definition: Learn multimodally-contextualized representations that are coordinated through their cross-modal interactions

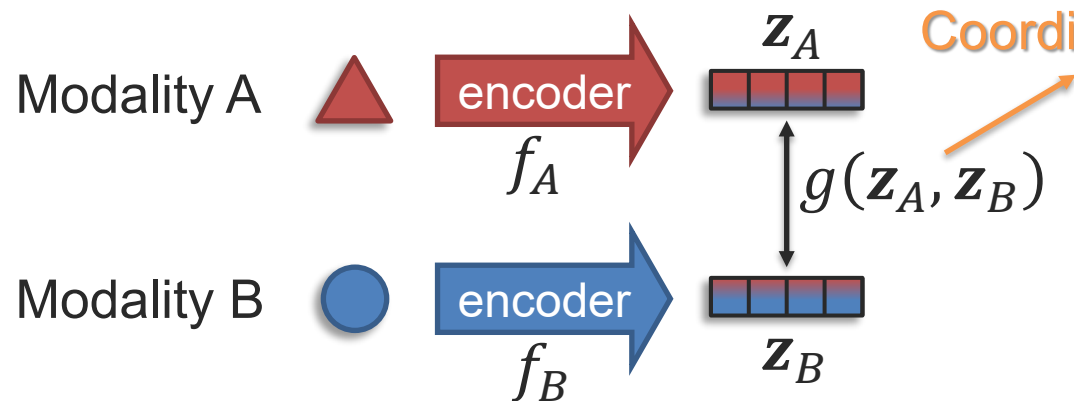
Strong Coordination:



Partial Coordination:



Coordination Function



Learning with coordination function:

$$\mathcal{L} = g(f_A(\triangle), f_B(\circ))$$

with model parameters θ_g , θ_{f_A} and θ_{f_B}

➡ Requires paired data

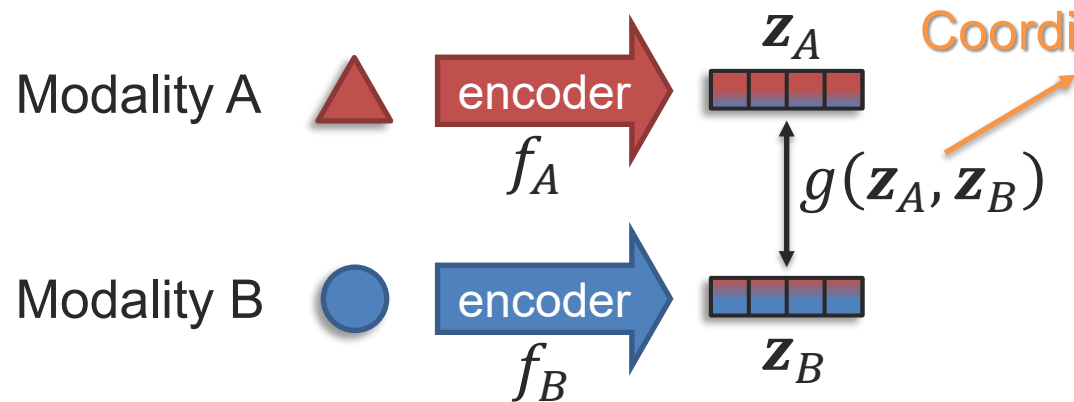
Examples of coordination function:

① Cosine similarity:
$$g(\mathbf{z}_A, \mathbf{z}_B) = \frac{\mathbf{z}_A \cdot \mathbf{z}_B}{\|\mathbf{z}_A\| \|\mathbf{z}_B\|}$$

Strong coordination!

➡ For normalized inputs (e.g., $\mathbf{z}_A - \overline{\mathbf{z}_A}$), equivalent to *Pearson correlation coefficient*

Coordination Function



Learning with coordination function:

$$\mathcal{L} = g(f_A(\triangle), f_B(\circ))$$

with model parameters θ_g , θ_{f_A} and θ_{f_B}

Examples of coordination function:

② Kernel similarity functions:

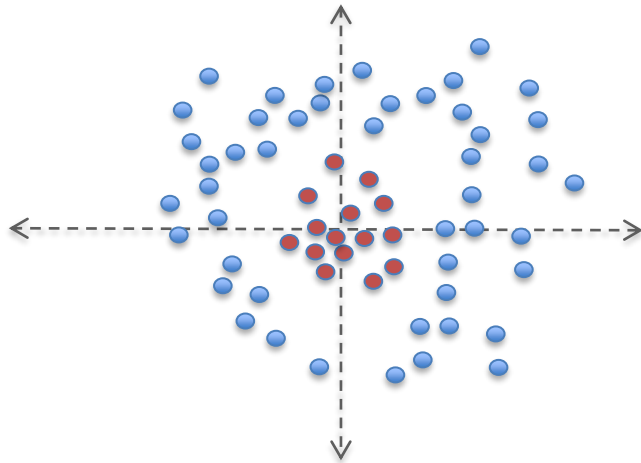
$$g(\mathbf{z}_A, \mathbf{z}_B) = k(\mathbf{z}_A, \mathbf{z}_B) \left\{ \begin{array}{l} \bullet \text{ Linear} \\ \bullet \text{ Polynomial} \\ \bullet \text{ Exponential} \\ \bullet \text{ RBF} \end{array} \right.$$

➡ All these examples bring relatively strong coordination between modalities

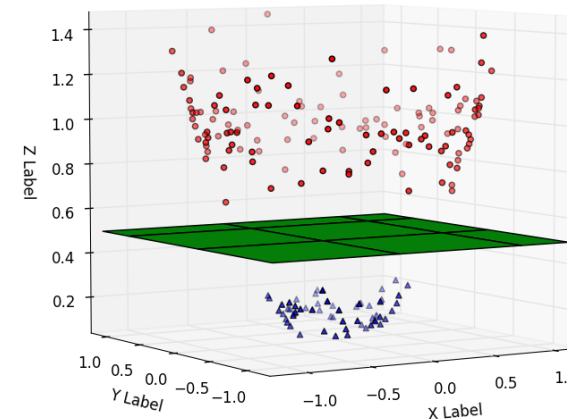
Kernel Function

A kernel function: Acts as a similarity metric between data points

$$K(x_i, x_j) = \phi(x_i)^T \phi(x_j) = \langle \phi(x_i), \phi(x_j) \rangle \quad \Rightarrow \phi(x) \text{ can be high-dimensional space!}$$



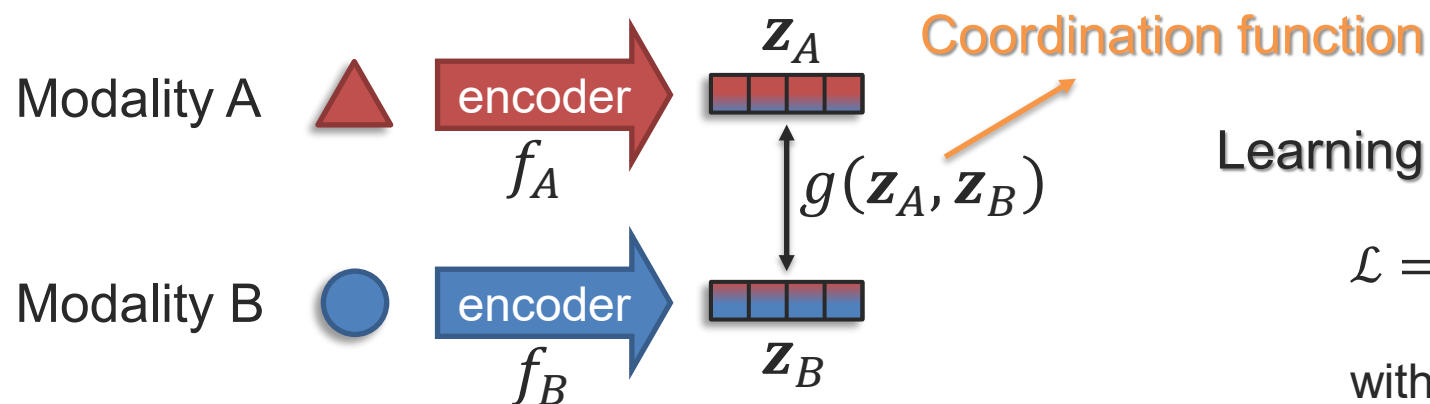
Not linearly separable in x space



Same data, but now linearly separable in $\phi(x)$ space

➡ **Radial Basis Function (RBF) Kernel :** $K(x_i, x_j) = \exp - \frac{1}{2\sigma^2} \|x_i - x_j\|^2$

Coordination Function



Learning with coordination function:

$$\mathcal{L} = g(f_A(\triangle), f_B(\circ))$$

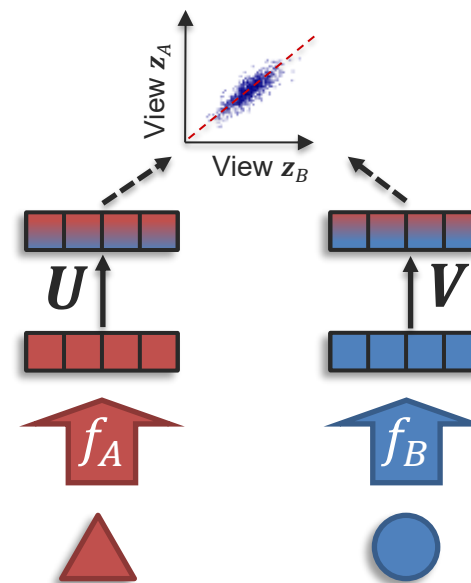
with model parameters θ_g , θ_{f_A} and θ_{f_B}

Examples of coordination function:

③ Canonical Correlation Analysis (CCA):

$$\operatorname{argmax}_{V, U, f_A, f_B} \operatorname{corr}(\mathbf{z}_A, \mathbf{z}_B)$$

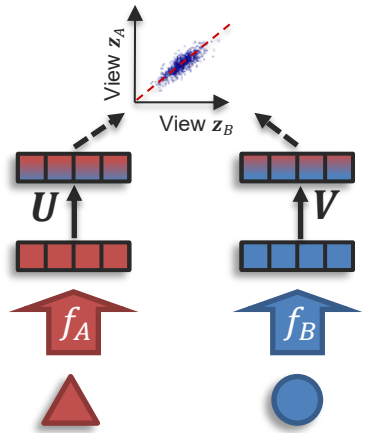
→ CCA includes multiple projections, all orthogonal with each others



Correlated Projection

- 1 Learn two linear projections, one for each view, that are maximally correlated:

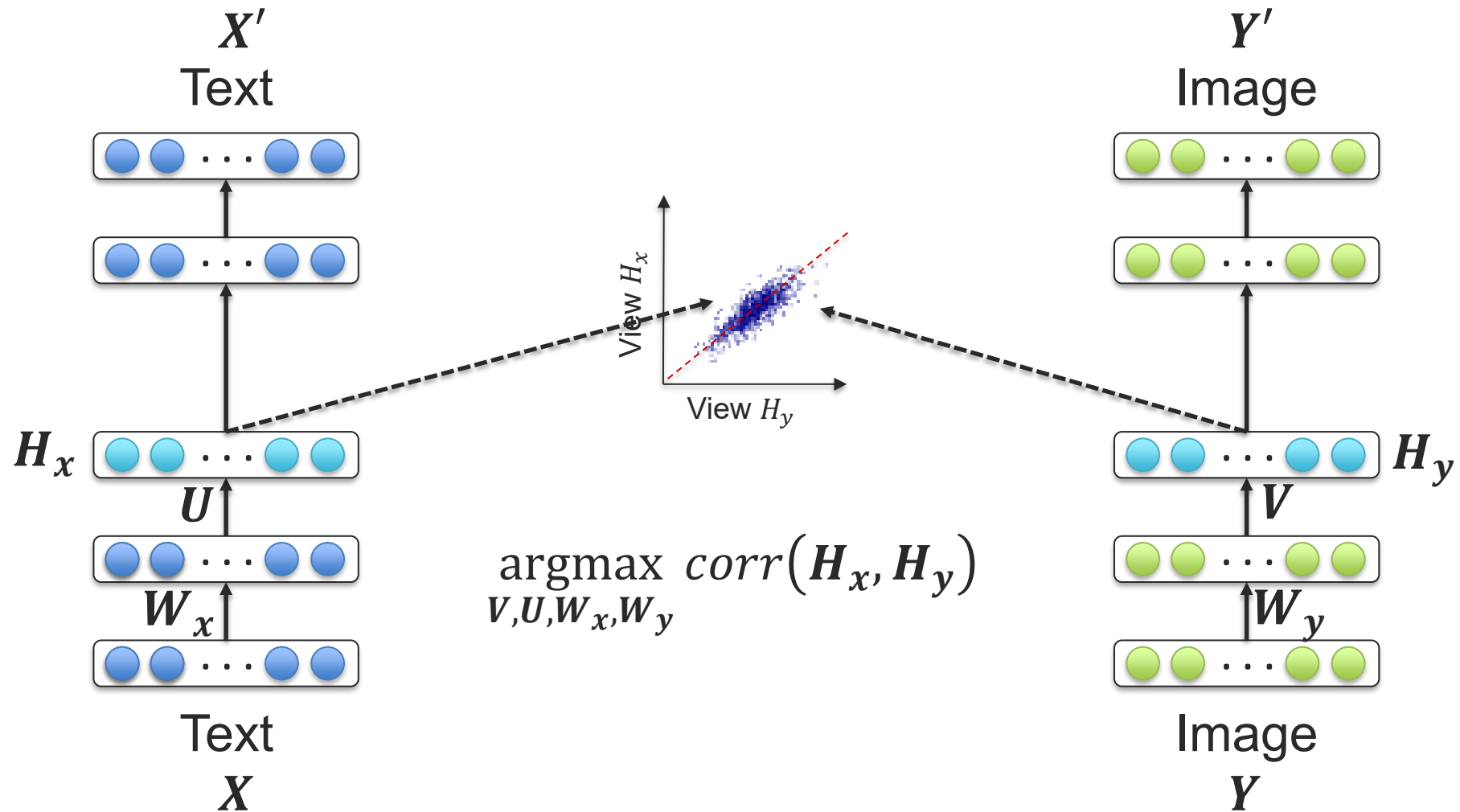
$$(u^*, v^*) = \operatorname{argmax}_{u, v} \operatorname{corr}(u^T X, v^T Y)$$



Two views X, Y where same instances have the same color

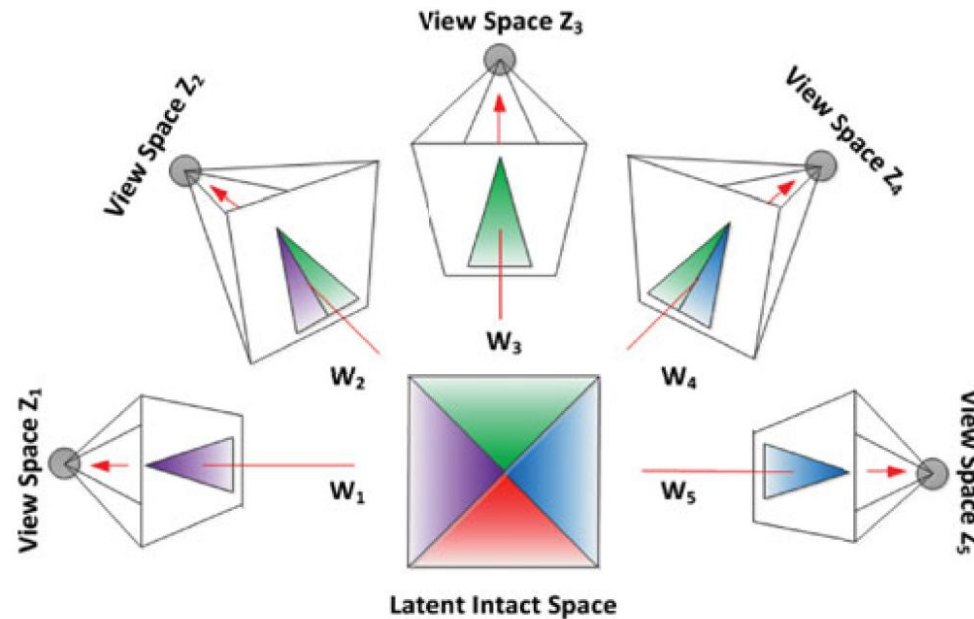
➡ Remember that X and Y consist of paired data

Deep Canonically Correlated Autoencoders (DCCAE)



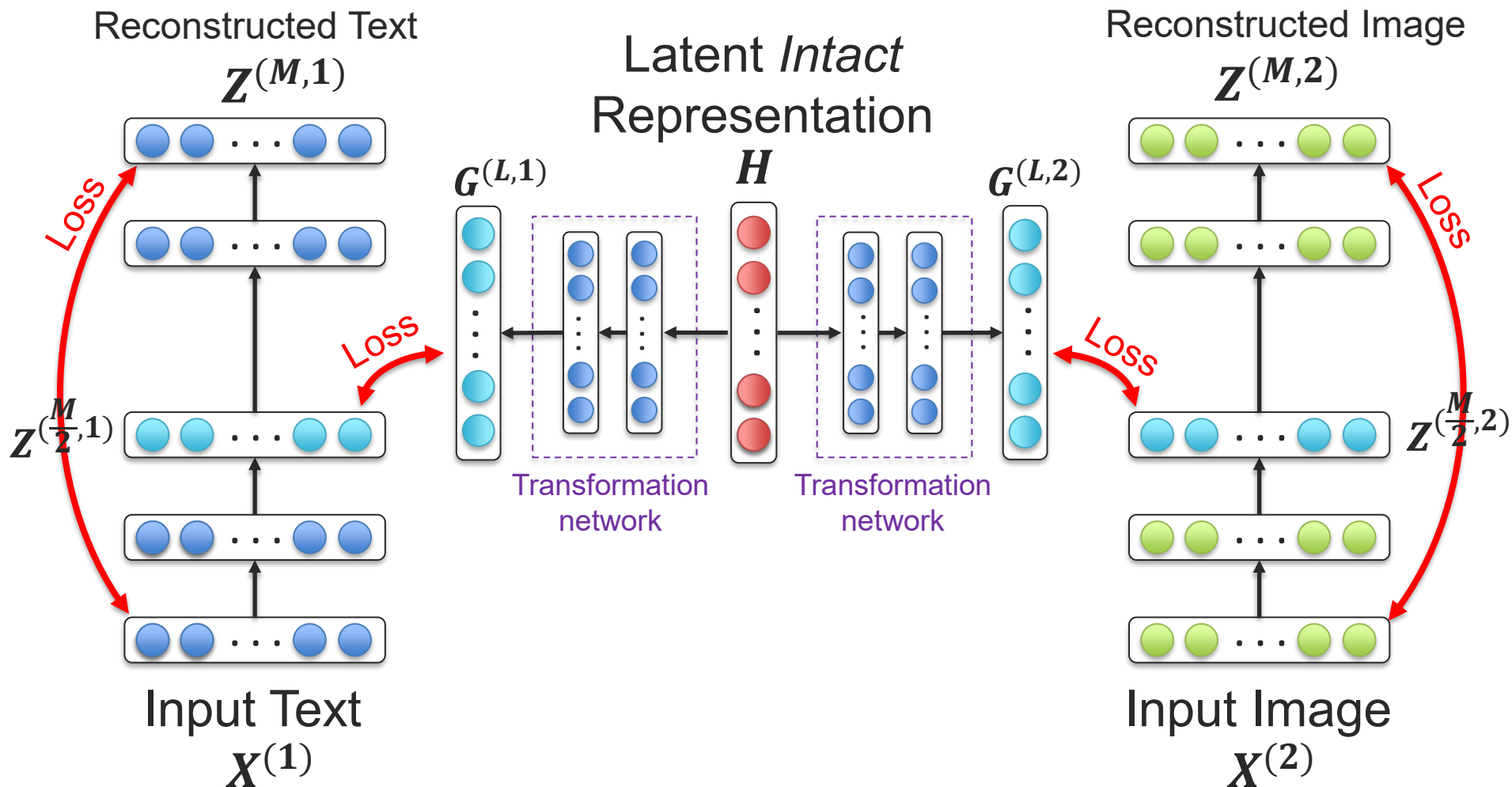
Multi-view Latent “Intact” Space

Given multiple views z_i from the same “object”:

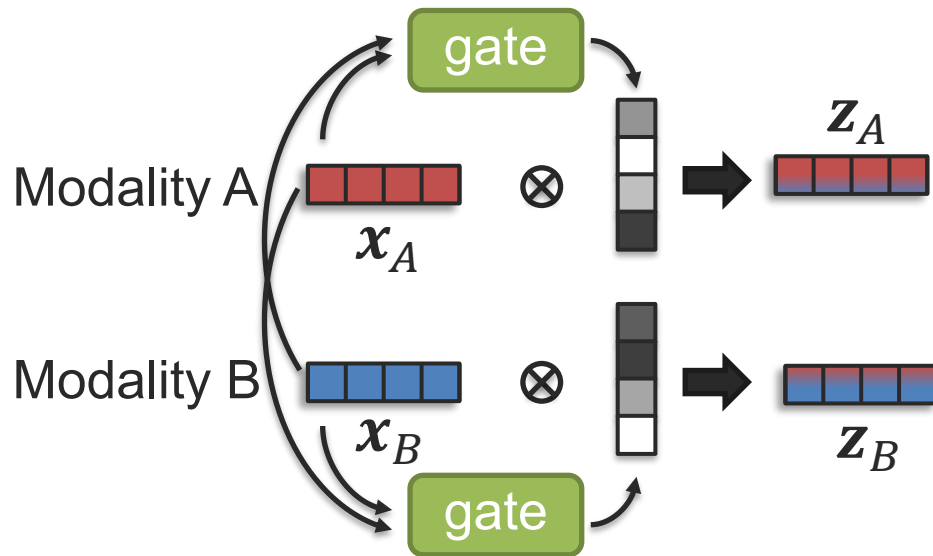


- 1) There is an “intact” representation which is *complete* and *not damaged*
- 2) The views z_i are partial (and possibly degenerated) representations of the intact representation

Auto-Encoder in Auto-Encoder Network



Gated Coordination



Gated coordination:

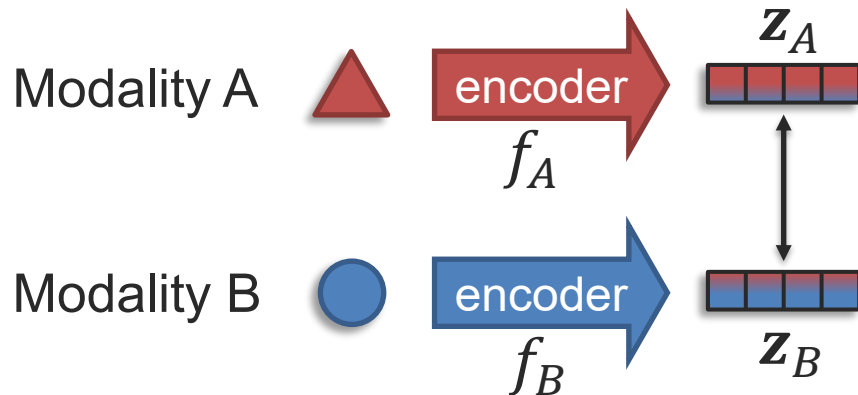
$$z_A = g_A(x_A, x_B) \cdot x_A$$

$$z_B = g_B(x_A, x_B) \cdot x_B$$

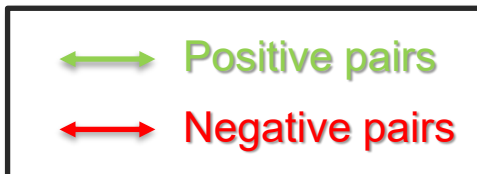
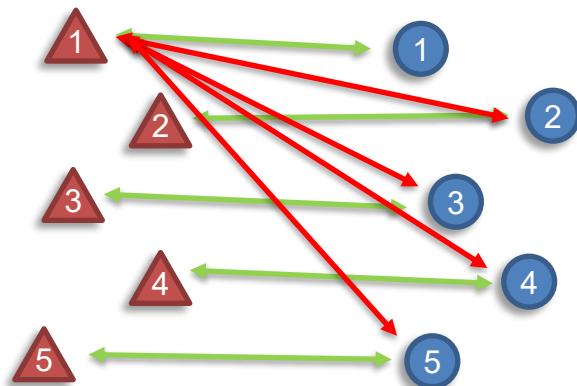
➡ Related to attention modules in transformers

More about it next week!

Coordination with Contrastive Learning



Paired data: $\{\triangle, \circ\}$
(e.g., images and text descriptions)



Contrastive loss:

→ brings **positive pairs** closer and pushes **negative pairs** apart

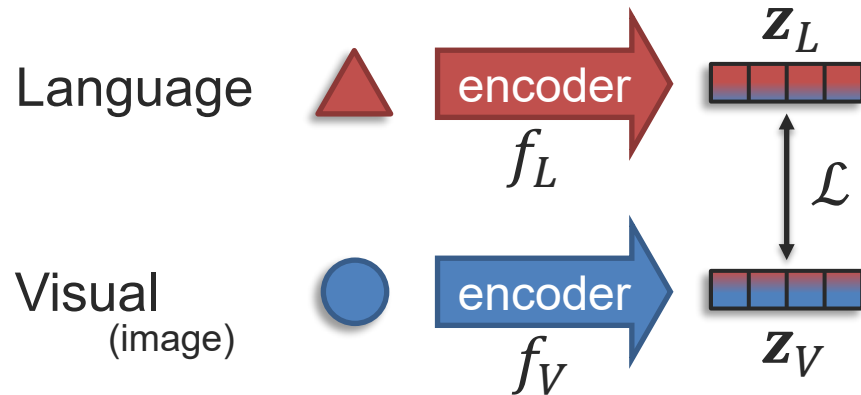
Simple contrastive loss:

$$\max\{0, \alpha + \underbrace{\text{sim}(z_A, z_B^+)}_{\text{positive pairs}} - \underbrace{\text{sim}(z_A, z_B^-)}_{\text{negative pair}}\}$$

Similarity functions are often cosine similarity

→ Similar to hinge loss

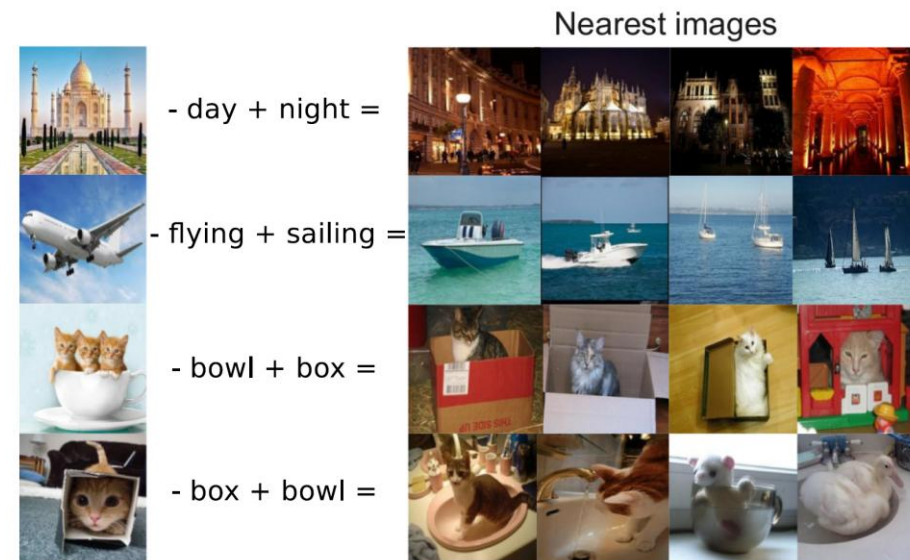
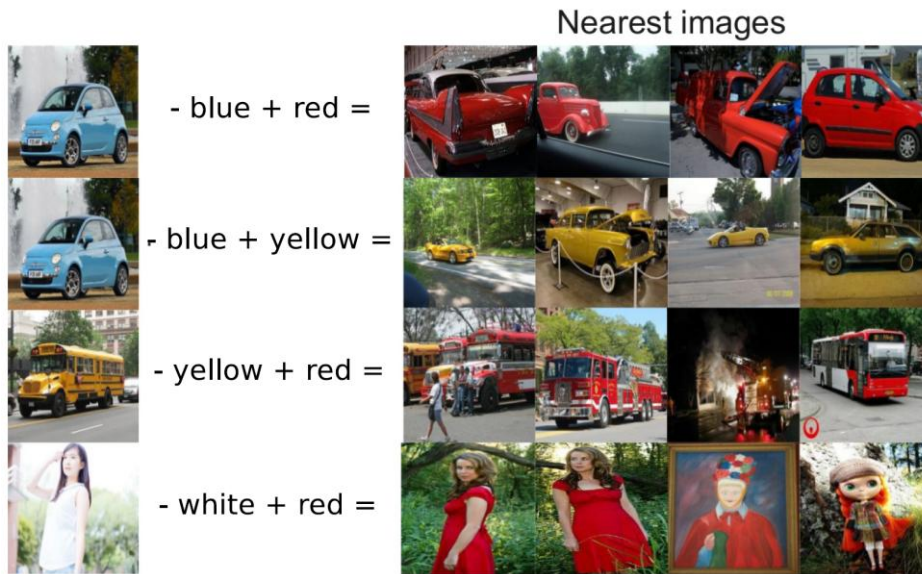
Example – Visual-Semantic Embeddings



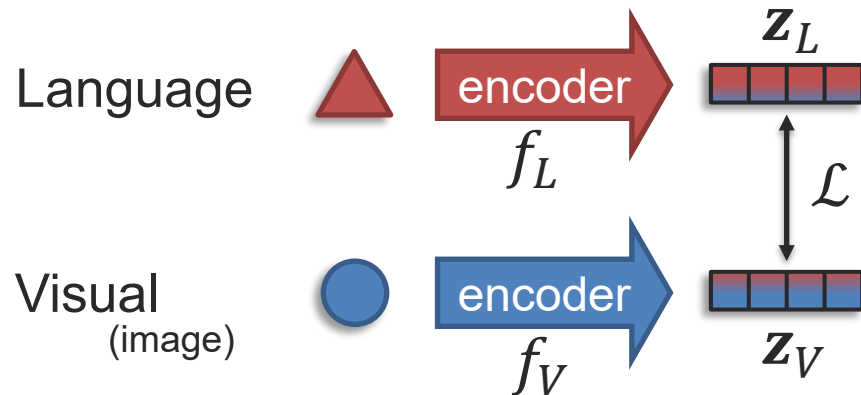
Two contrastive loss terms:

$$\max\{0, \alpha + \text{sim}(\mathbf{z}_L, \mathbf{z}_V^+) - \text{sim}(\mathbf{z}_L, \mathbf{z}_V^-)\}$$

$$+ \max\{0, \alpha + \text{sim}(\mathbf{z}_V, \mathbf{z}_L^+) - \text{sim}(\mathbf{z}_V, \mathbf{z}_L^-)\}$$



Example – CLIP (Contrastive Language–Image Pre-training)



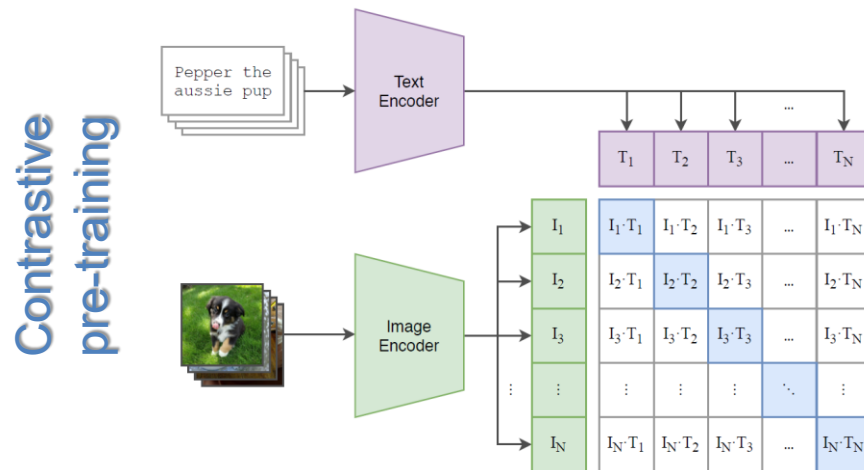
Popular contrastive loss: InfoNCE

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

Annotations:

- $\text{sim}(\mathbf{z}_A^l, \mathbf{z}_B^l)$ (green box) → positive pairs
- $\sum_{j=1}^N \text{sim}(\mathbf{z}_A^i, \mathbf{z}_B^j)$ (red box) → negative pairs and positive pairs
- Similarity function can be cosine similarity

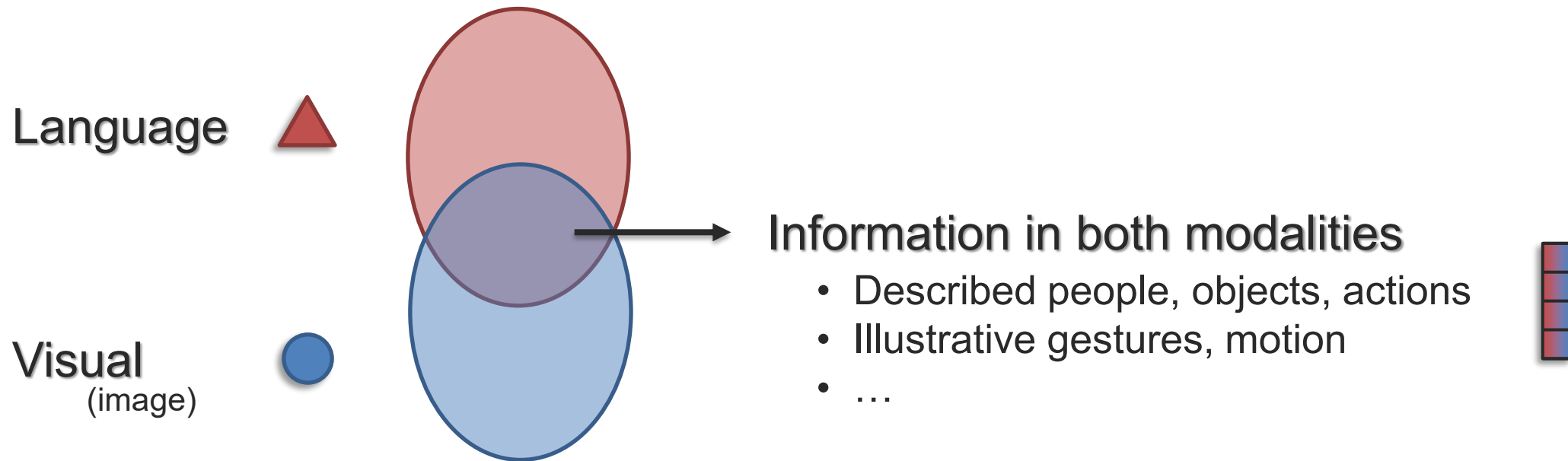
Positive and negative pairs:



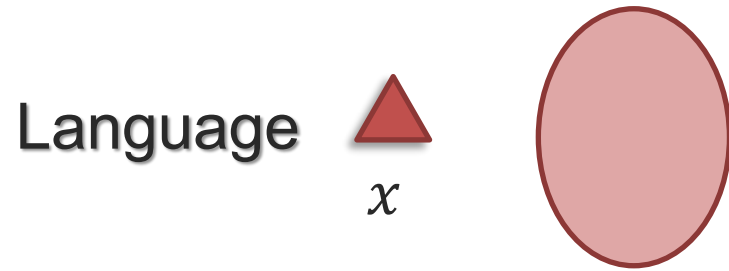
➔ CLIP encoders (f_L and f_V) are great for language-vision tasks

➔ $\mathbf{z}_L$ and $\mathbf{z}_V$ are coordinated but not identical representation spaces

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”

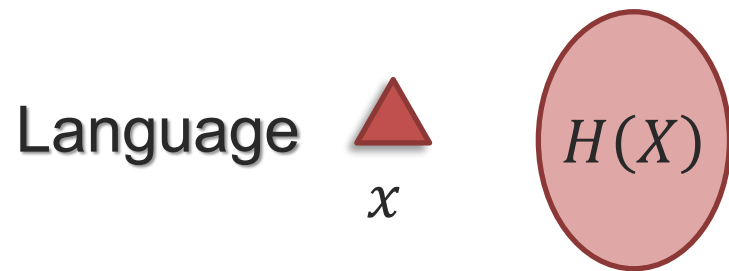
➡ More random... So, higher information

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

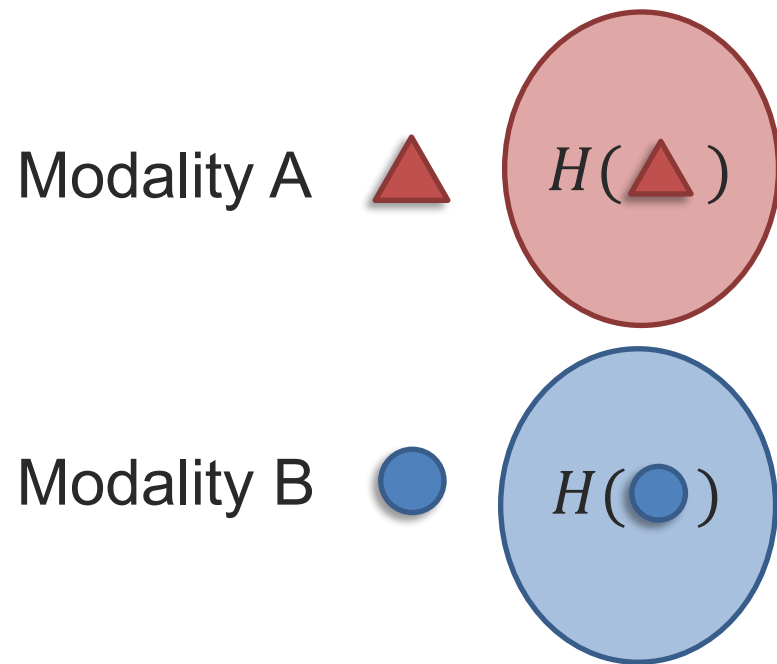
Entropy: weighted average of all possible outcomes from $\mathcal{X}$

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

➡ Entropy can also be defined for continuous random variables

Entropy with Two Modalities

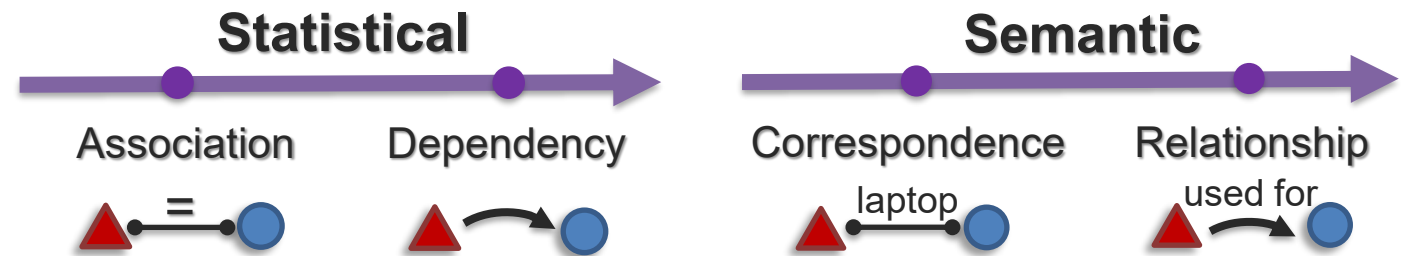
If no overlapping information



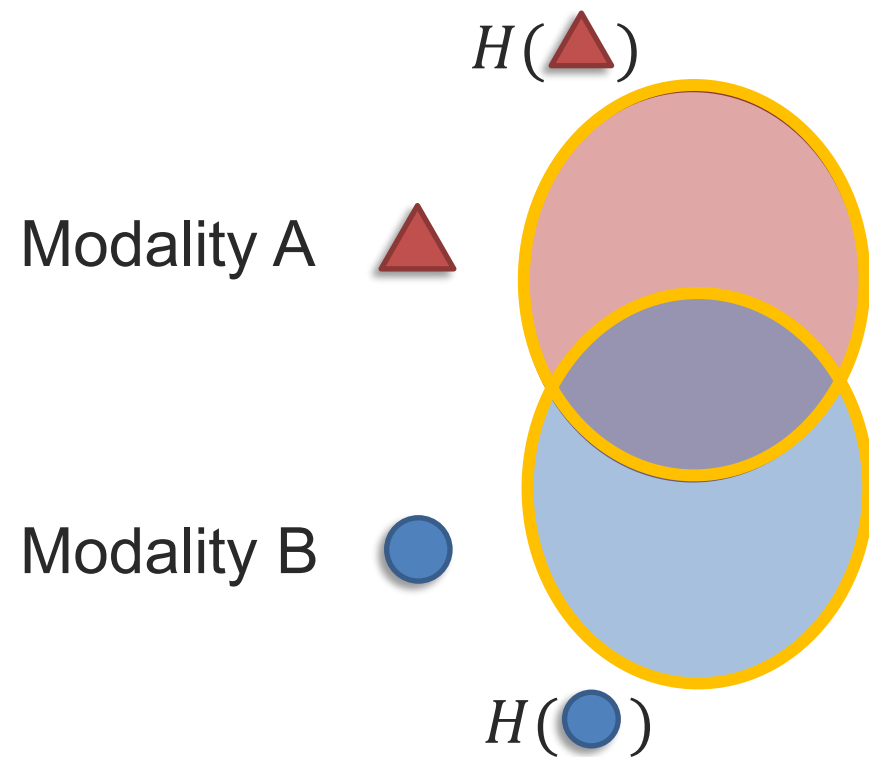
➡ But in most real-world scenarios, modalities are *inter-connected*



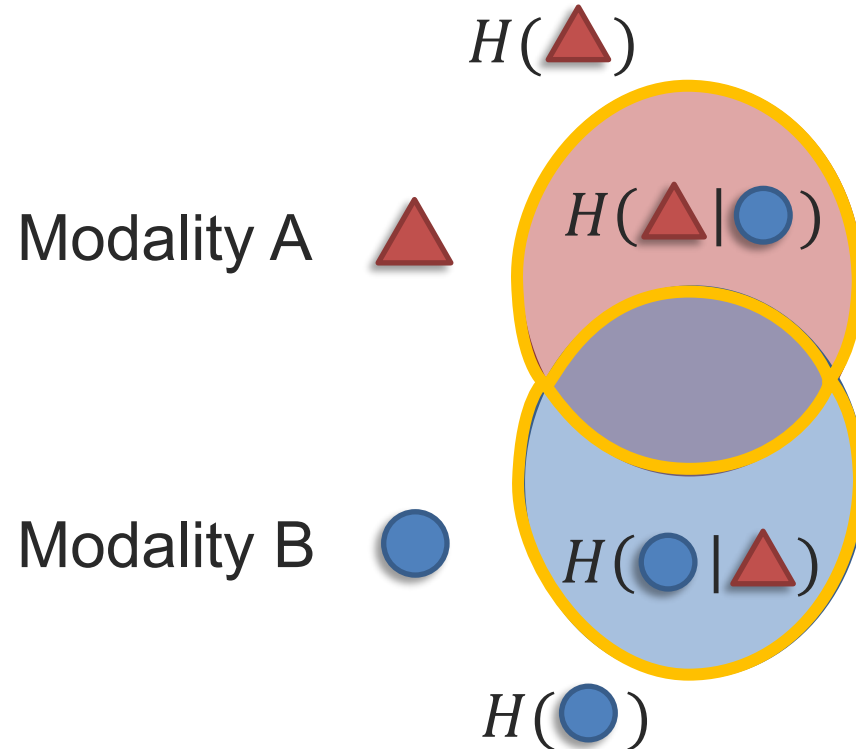
A **teacup** on the right of a **laptop** in a clean room.



Entropy with Two Modalities



Entropy with Two Modalities



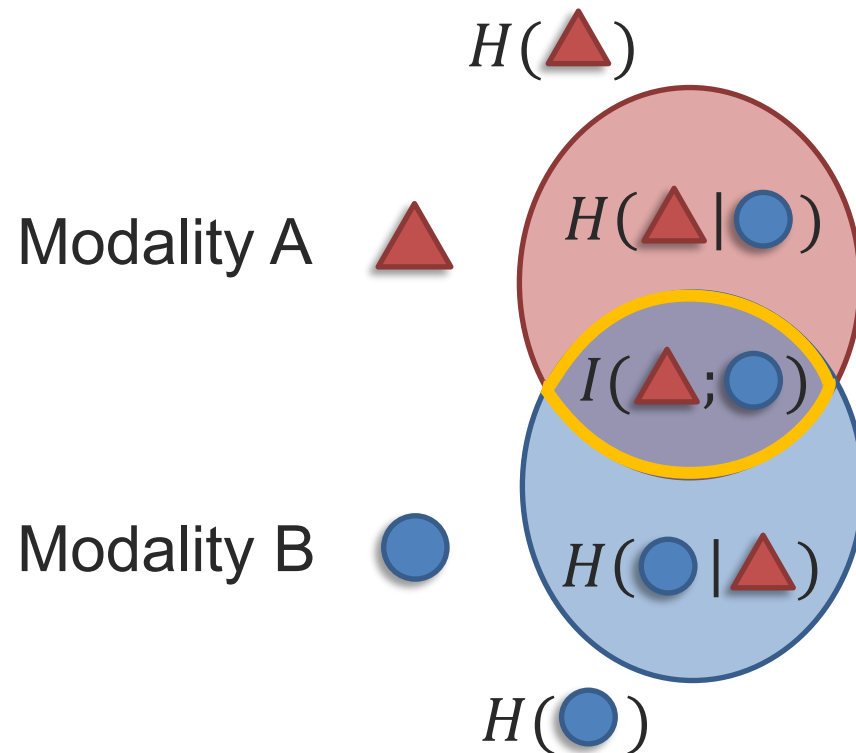
Conditional entropy $H(Y|X)$

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

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

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 $I(X; Y)$

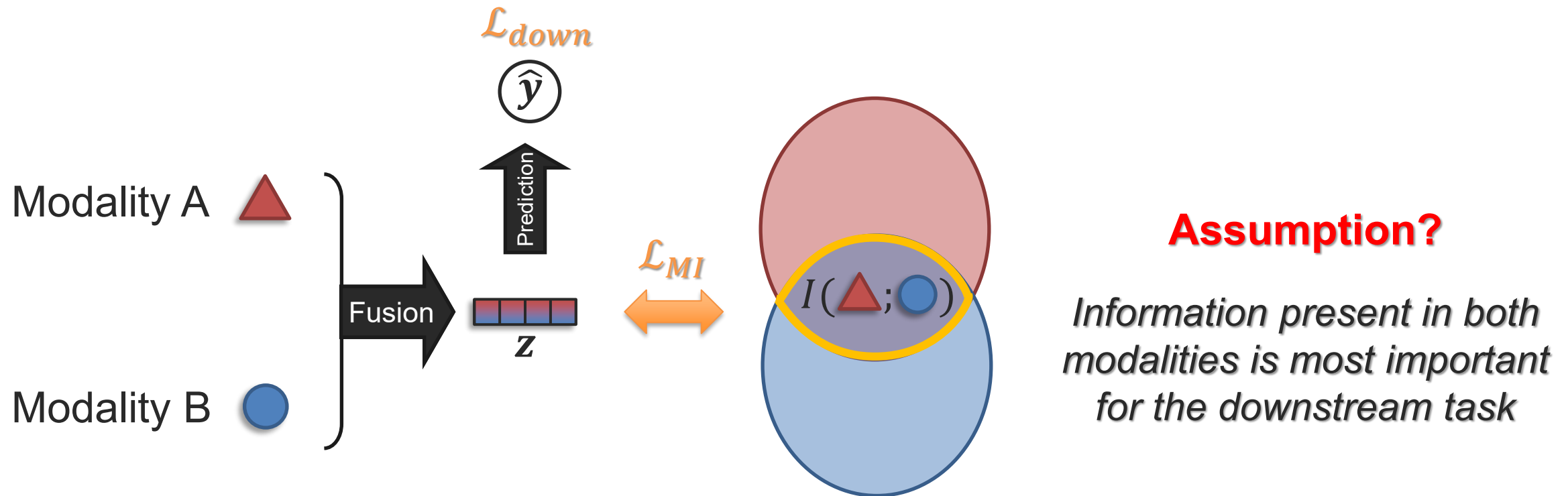
$$I(X; Y) = H(X) - H(X|Y)$$

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

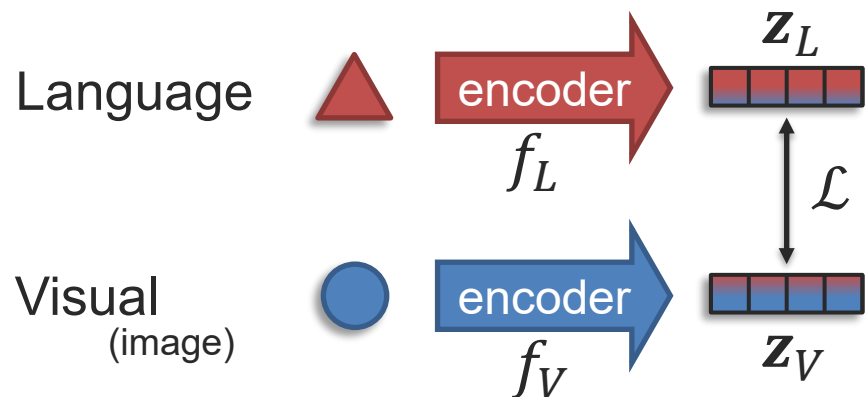
$$I(X; Y) = \mathbb{E}_{X,Y} \left[\log \frac{P_{XY}(x, y)}{P_X(x)P_Y(y)} \right]$$

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



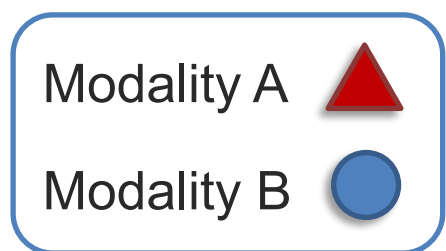
Contrastive Learning and Connected Modalities



Popular contrastive loss: InfoNCE

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

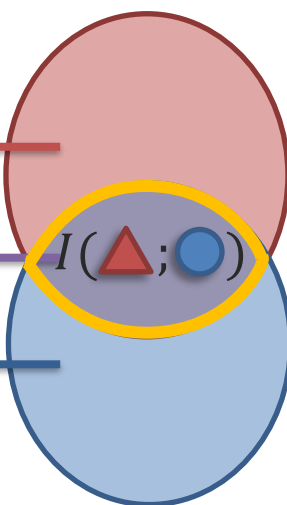
Connected modalities:



unique

shared

unique



Mutual information $I(X; Y)$

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

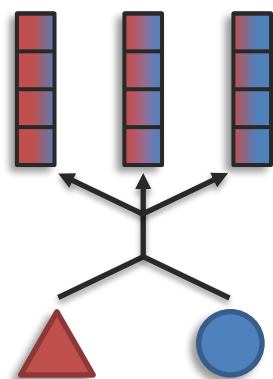


CLIP focuses on shared connections

Representation

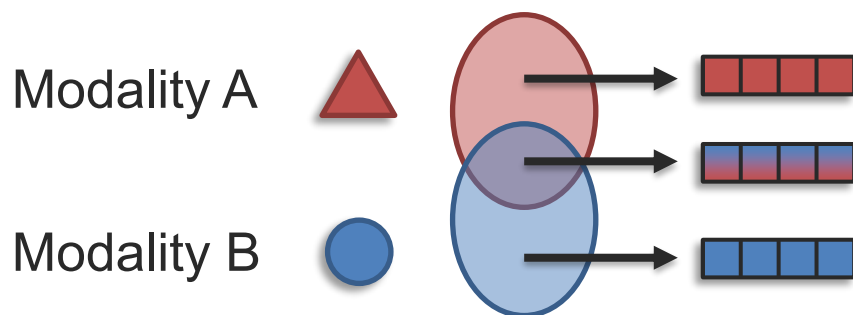
Fission

Sub-Challenge 1c: Representation Fission

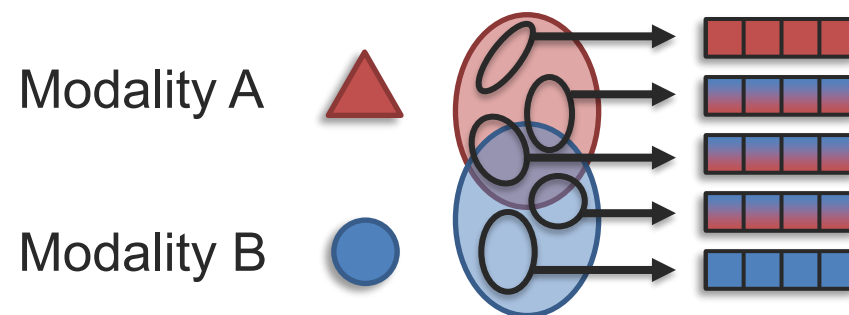


Definition: Learning a new set of representations that reflects multimodal internal structure such as data factorization or clustering

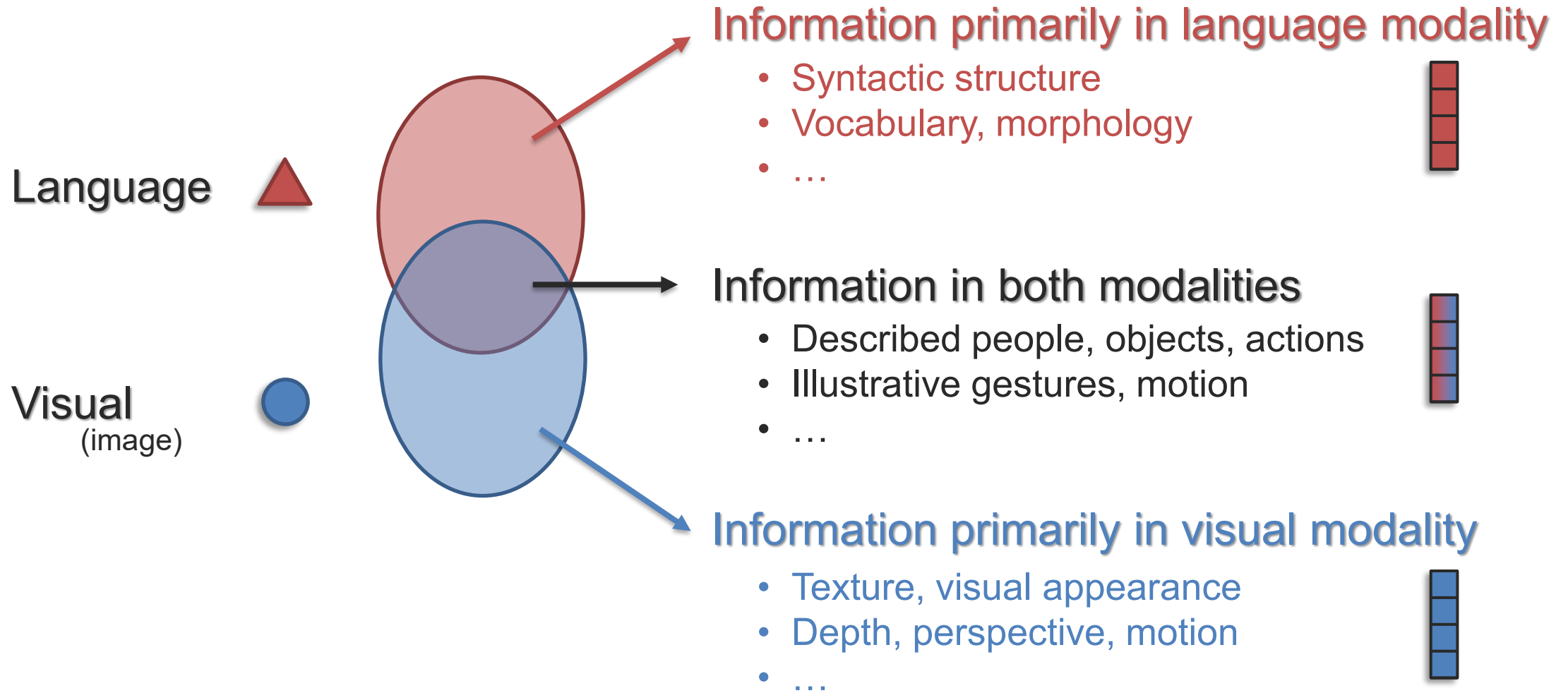
Modality-level fission:



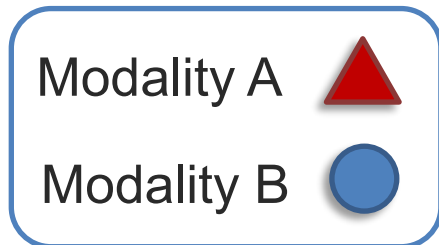
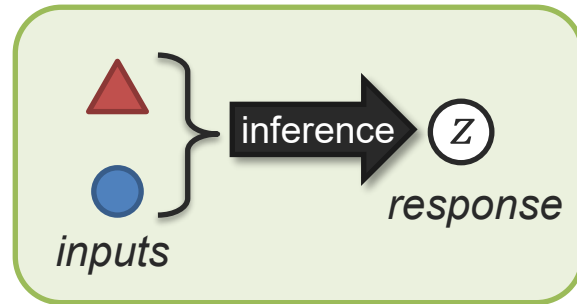
Fine-grained fission:



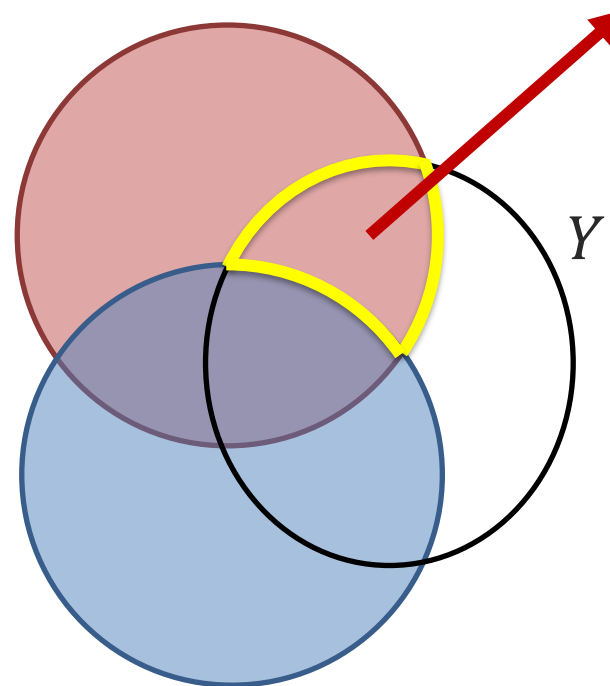
Modality-Level Fission



Modeling Multimodal Interactions

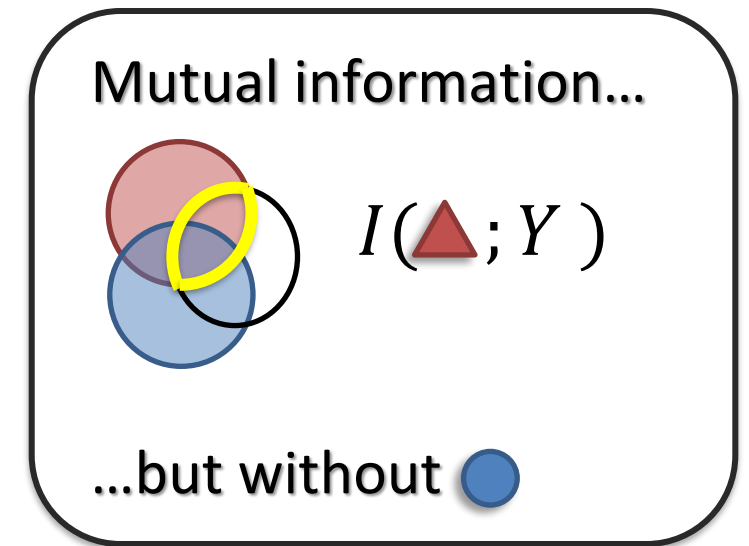


Information theory as a framework for modeling multimodal interactions between input modalities and responses

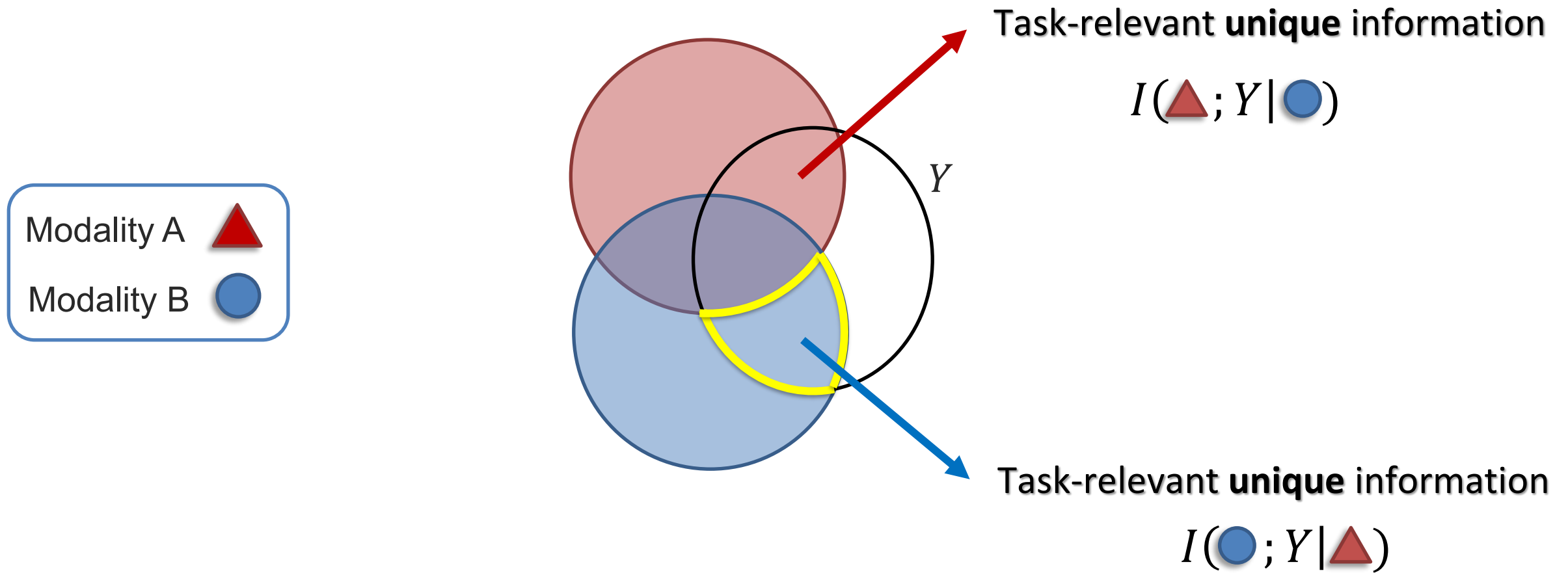


Task-relevant **unique** information

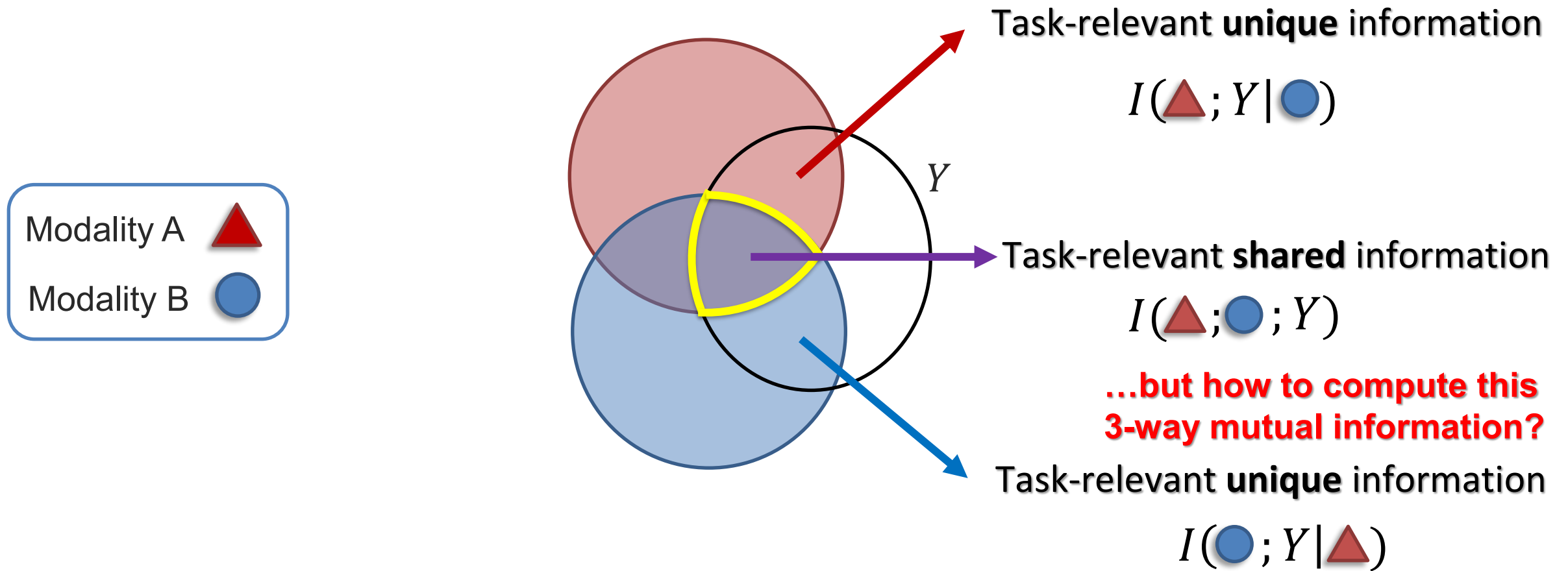
$$I(\triangle; Y | \bullet)$$



Modeling Multimodal Interactions



Modeling Multimodal Interactions



Modeling Multimodal Interactions

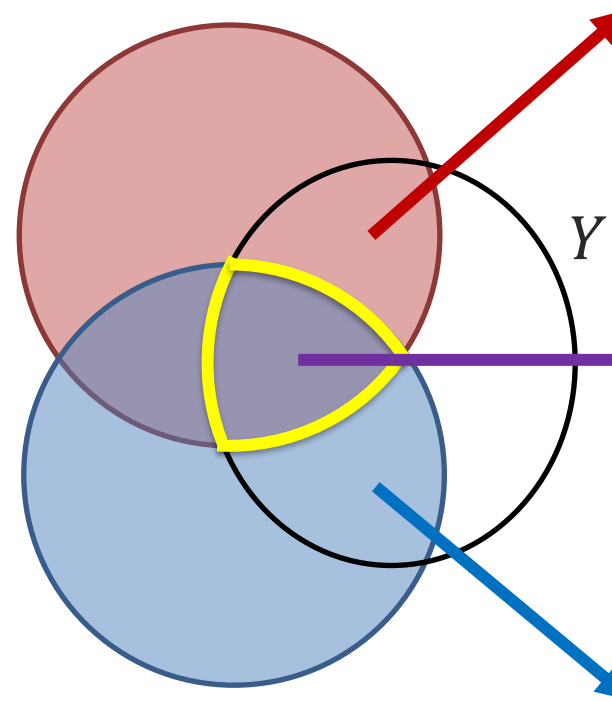
Partial Information Decomposition (PID)

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

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:

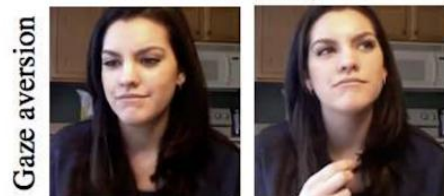
Liang et al., Quantifying & Modeling Feature Interactions: An Information Decomposition Framework. Neurips 2023

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.*

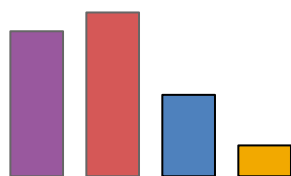
Vision:



Acoustic:

(frustrated voice)

Sentiment



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

Language/Agreement

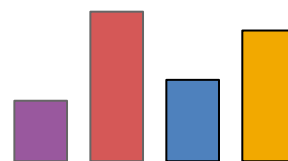
Sheldon :

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

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

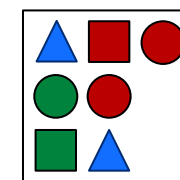


Sarcasm



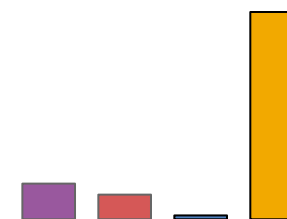
$R \ U_\ell U_{av} \ S$

Multimodal Transformer



Is there a
red shape
above a
circle?

VQA

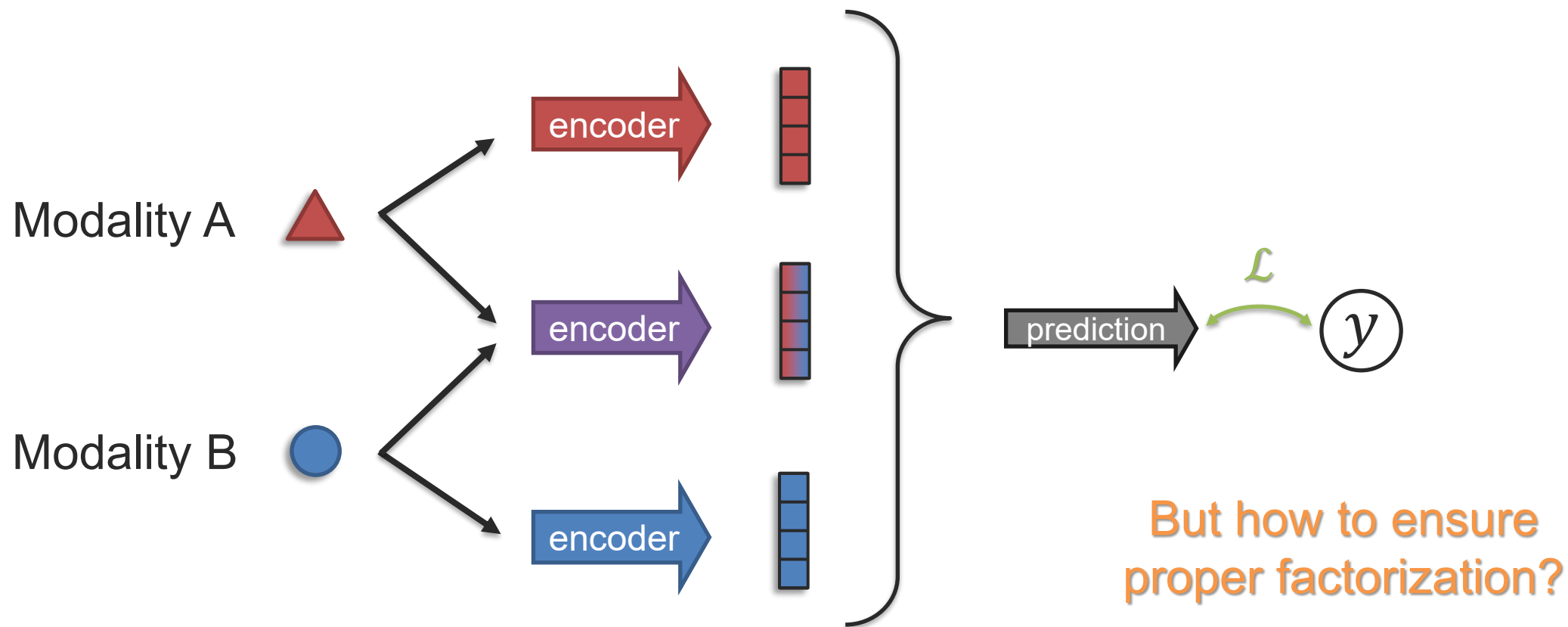


$R \ U_\ell U_i \ S$

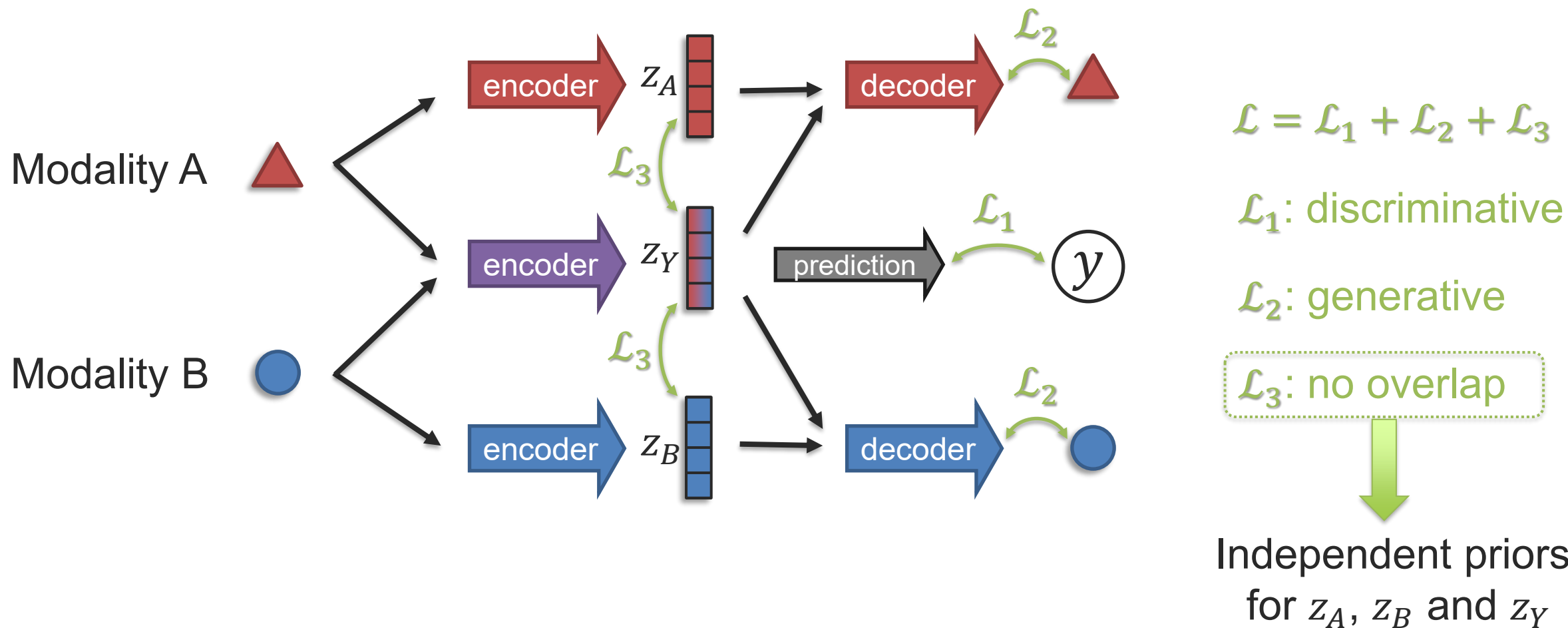
Multiplicative/Transformer

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

Factorized Multimodal Representations



A Generative-Discriminative Approach

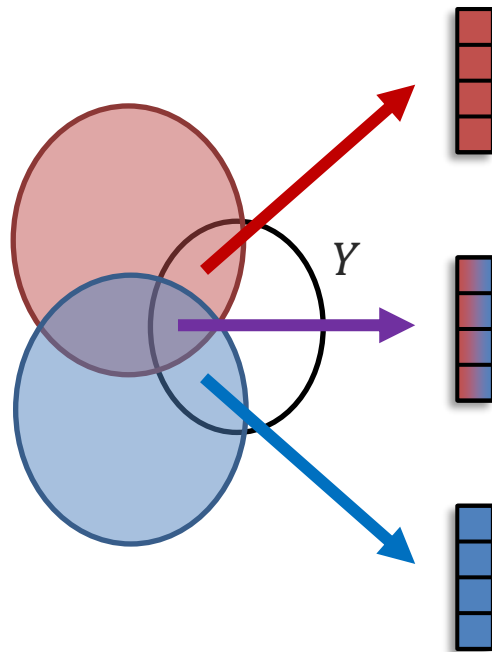


Learning Task-relevant Unique Information

Modeling task-relevant unique information



Can you please pass the cow?



- 2 Maximize task-relevant **unique** information

$$I(\mathbf{Z}; Y | \bullet)$$

- 1 Maximize task-relevant **shared** information

$$I(\mathbf{Z}; \bullet; Y) \quad \text{and} \quad I(\mathbf{Z}; \blacktriangle; Y)$$

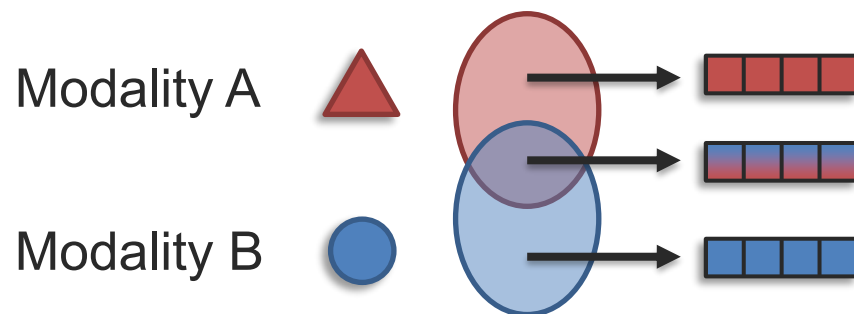
- 3 Maximize task-relevant **unique** information

$$I(\mathbf{Z}; Y | \blacktriangle)$$

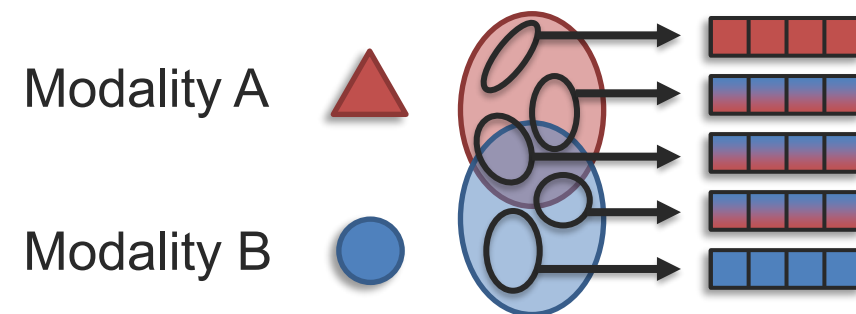
Fine-Grained Fission

How to automatically discover these internal clusters, factors?

Modality-level fission:

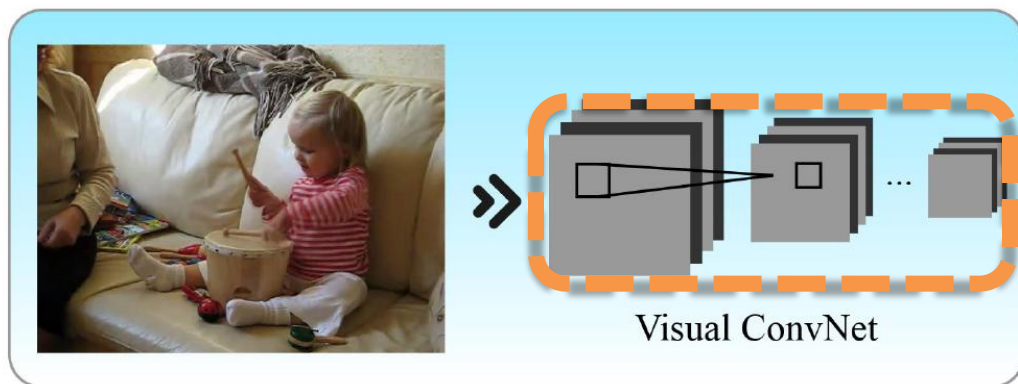


Fine-grained fission:

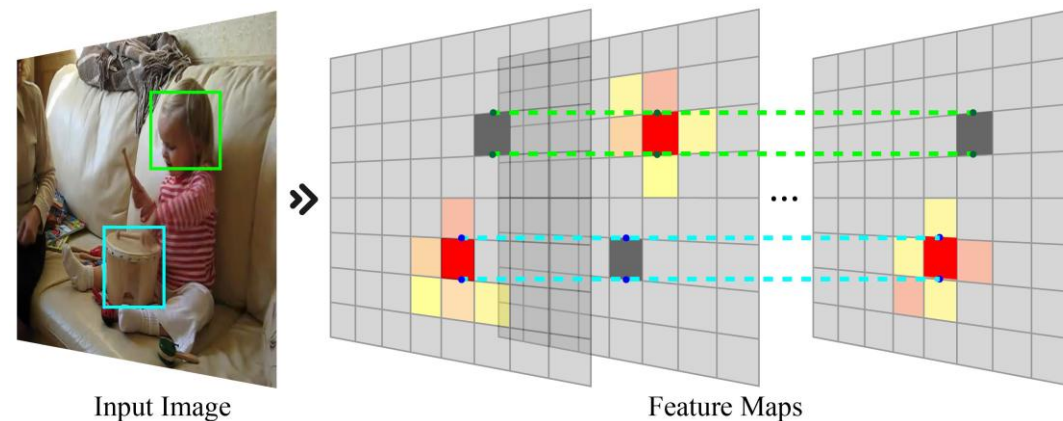
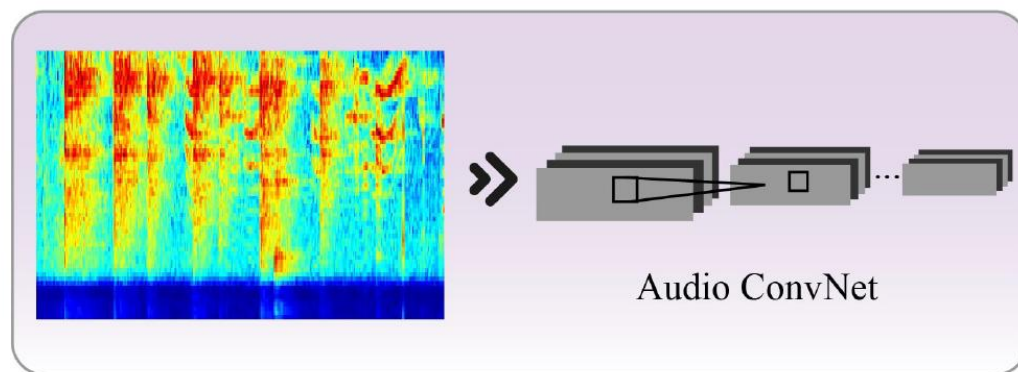


Fine-Grained Fission – A Clustering Approach

Unimodal Encoders

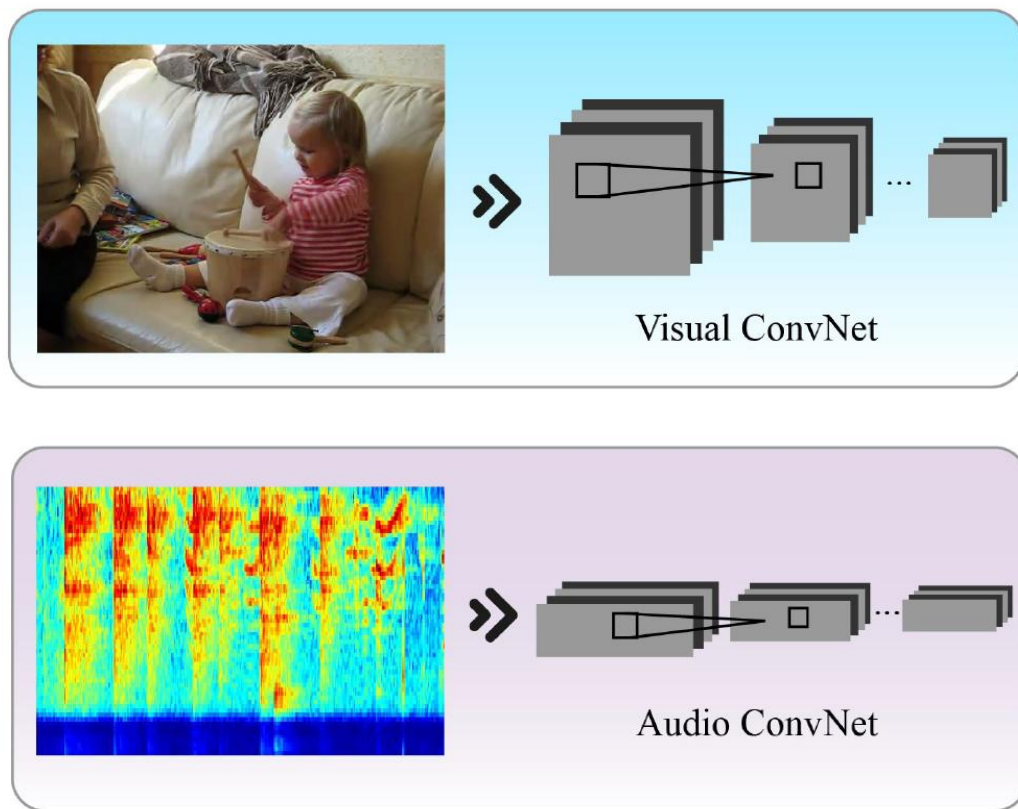


Localized activations for different objects



Fine-Grained Fission – A Clustering Approach

Unimodal Encoders



Multimodal Fission

