

DS-GA.3001

Embodied Learning and Vision

Mengye Ren

NYU

Spring 2025

embodied-learning-vision-course.github.io

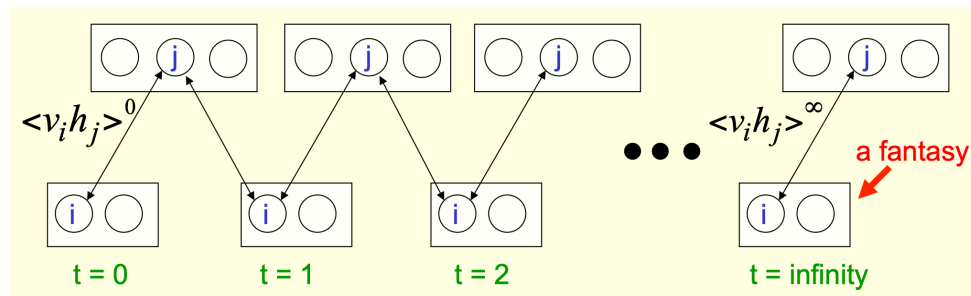


Lecture Slides for Note Taking

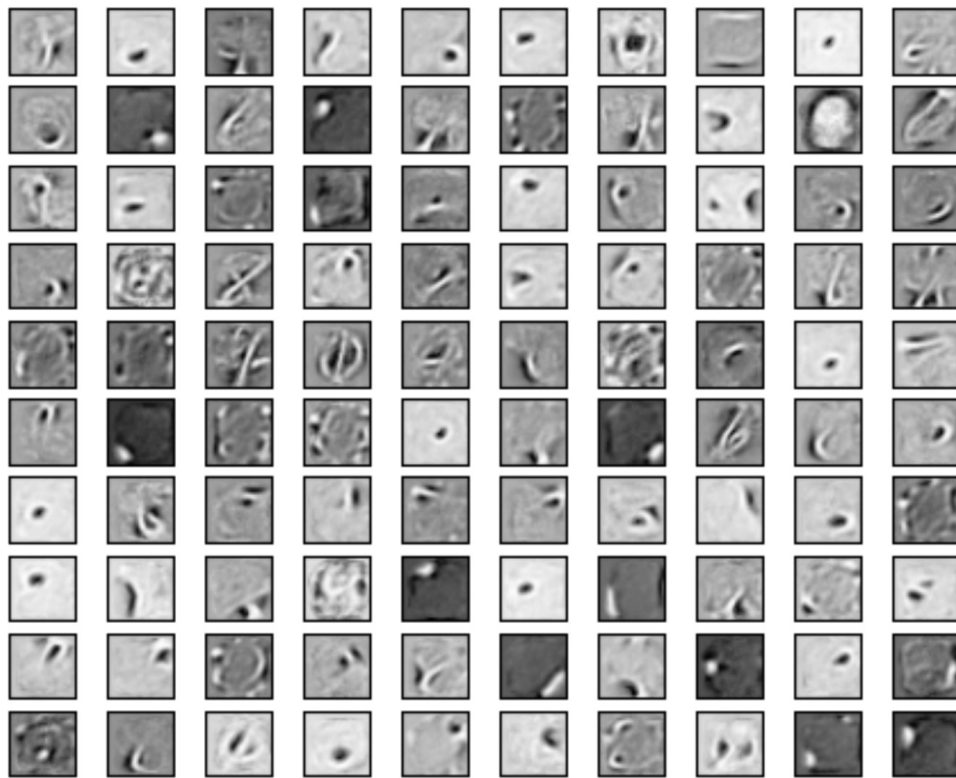


Energy-Based Learning

- Example: RBMs
- Energy: $E(v, h) = - \sum_{i,j} v_i h_j w_{ij}$.
 $p(h_j = 1 | v_i) = \sigma(\sum_{ij} v_i w_{ij)$.



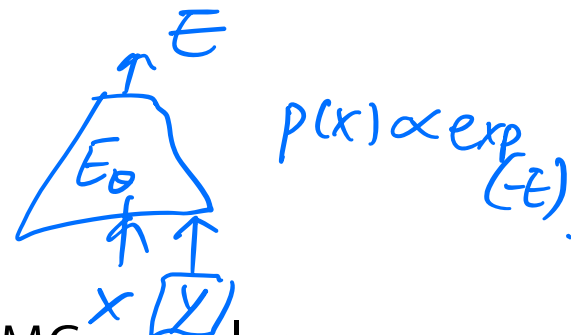
$$\frac{\partial \log p(v)}{\partial w_{ij}} = \langle v_i h_j \rangle^0 - \langle v_i h_j \rangle^\infty.$$



General EBM

- Inference requires running gradient descent and MCMC samples (Langevin samples).

General EBM



- Inference requires running gradient descent and MCMC samples (Langevin samples).

$$\tilde{\mathbf{x}}^k = \mathbf{x}^{k-1} - \frac{\lambda}{2} \nabla_{\mathbf{x}} E_{\theta}(\tilde{\mathbf{x}}^{k-1}) + \omega^k, \omega^k \sim \mathcal{N}(0, \lambda).$$

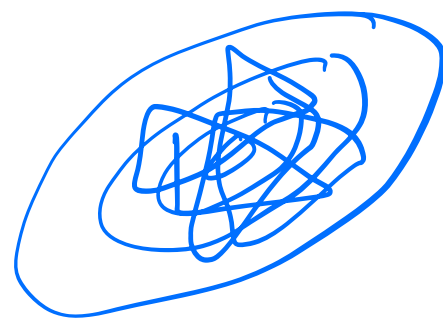
$$\nabla_{\theta} \mathcal{L} = \mathbb{E}_{\mathbf{x} \sim p_D} [\nabla_{\theta} E_{\theta}(\mathbf{x}^+)] - \mathbb{E}_{\mathbf{x} \sim q_{\theta}} [\nabla_{\theta} E_{\theta}(\mathbf{x}^-)].$$

data.

model.

negative.

→ gradient descent on input.



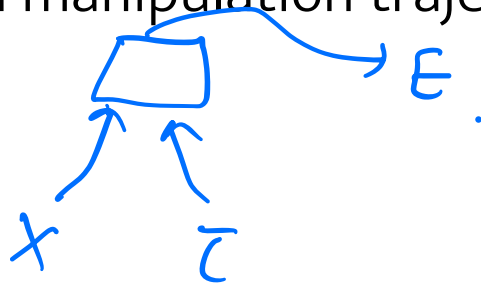
General EBM

- Inference requires running gradient descent and MCMC samples (Langevin samples).

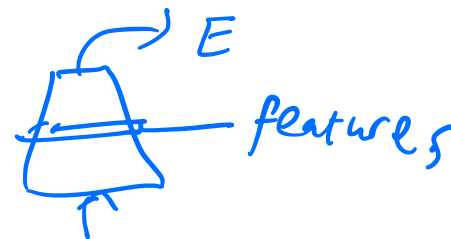
$$\tilde{\mathbf{x}}^k = \mathbf{x}^{\tilde{k}-1} - \frac{\lambda}{2} \nabla_{\mathbf{x}} E_{\theta}(\tilde{\mathbf{x}}^{k-1}) + \omega^k, \omega^k \sim \mathcal{N}(0, \lambda).$$

$$\nabla_{\theta} \mathcal{L} = \mathbb{E}_{\mathbf{x} \sim p_D} [\nabla_{\theta} E_{\theta}(\mathbf{x}^+)] - \mathbb{E}_{\mathbf{x} \sim q_{\theta}} [\nabla_{\theta} E_{\theta}(\mathbf{x}^-)].$$

- Can be applied on hand manipulation trajectory generation.



General EBM



- Inference requires running gradient descent and MCMC samples (Langevin samples).

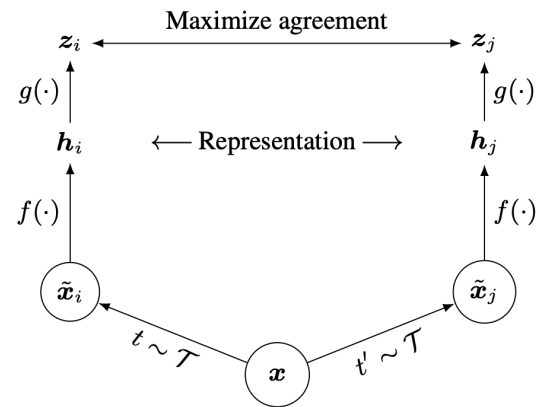
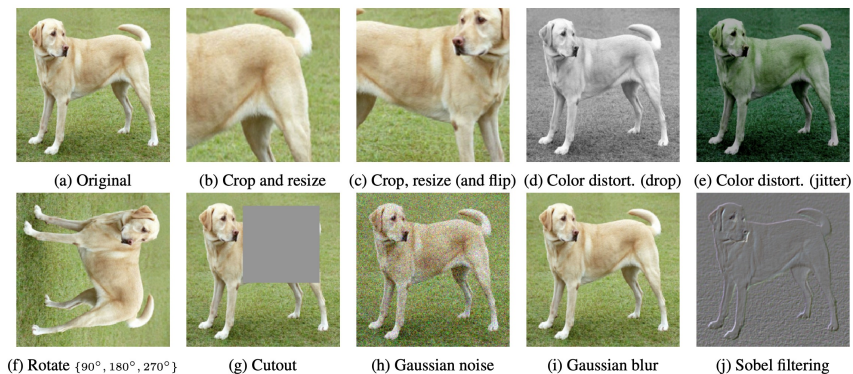
$$\tilde{\mathbf{x}}^k = \mathbf{x}^{\tilde{k}-1} - \frac{\lambda}{2} \nabla_{\mathbf{x}} E_{\theta}(\tilde{\mathbf{x}}^{k-1}) + \omega^k, \omega^k \sim \mathcal{N}(0, \lambda).$$

$$\nabla_{\theta} \mathcal{L} = \mathbb{E}_{\mathbf{x} \sim p_D} [\nabla_{\theta} E_{\theta}(\mathbf{x}^+)] - \mathbb{E}_{\mathbf{x} \sim q_{\theta}} [\nabla_{\theta} E_{\theta}(\mathbf{x}^-)].$$

- Can be applied on hand manipulation trajectory generation.
- Good results in generation but still not a generalized representation learning algorithm.

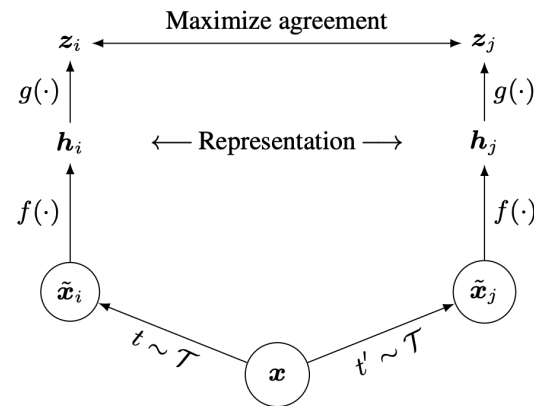
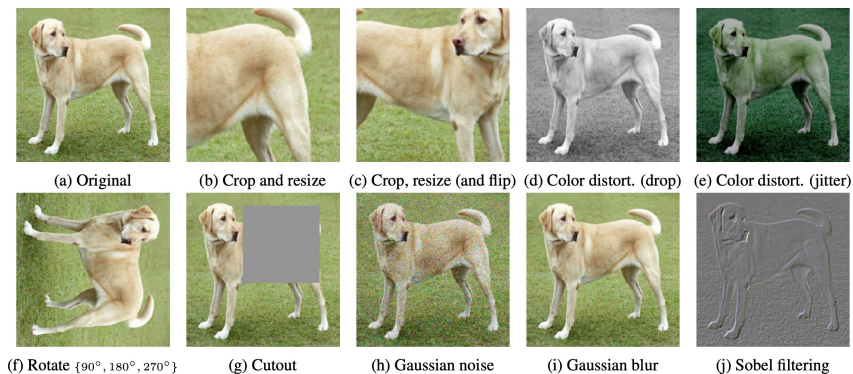
Self-Supervised Visual Learning

- Match the same image (with severe augmentation)



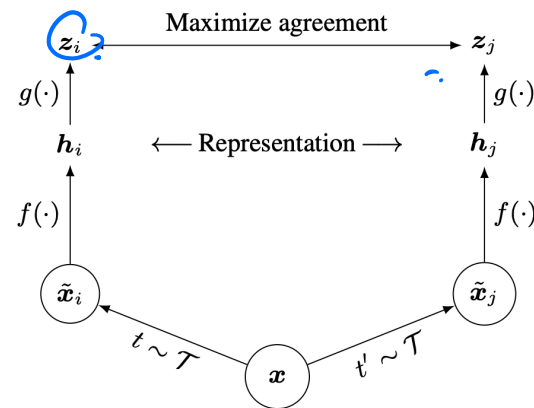
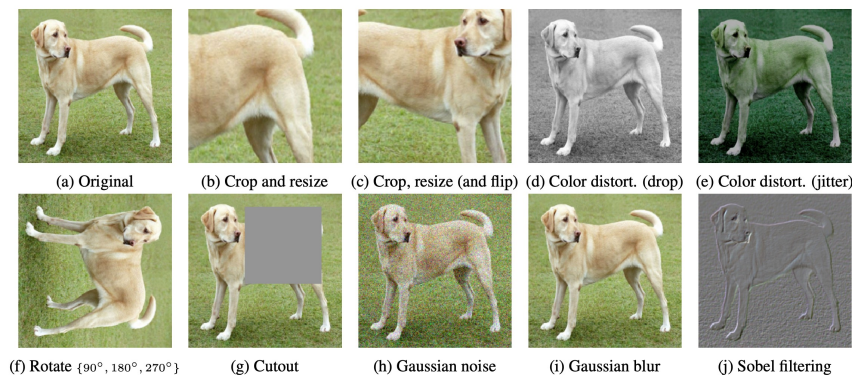
Self-Supervised Visual Learning

- Match the same image (with severe augmentation)
- Joint embedding approach: Apply loss on the embedding level.



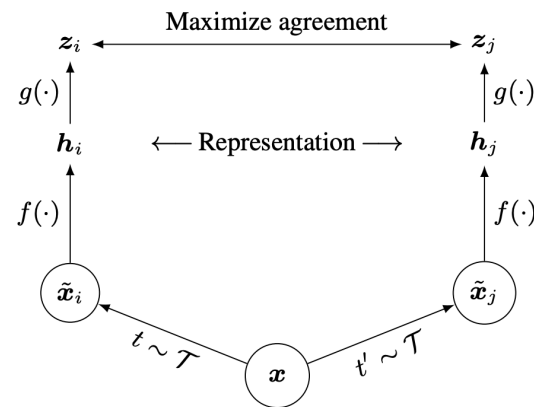
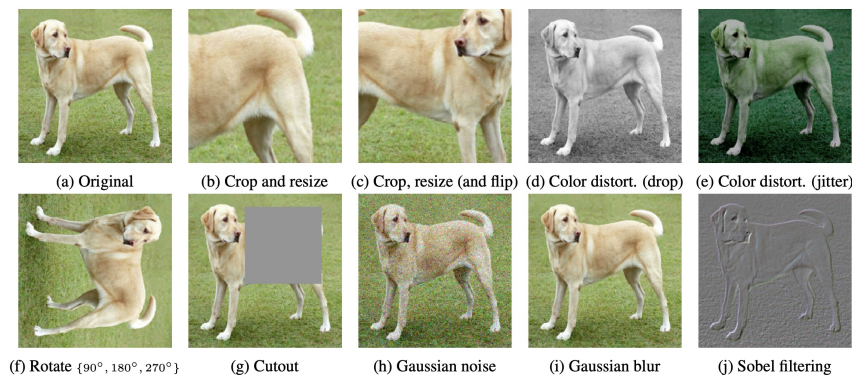
Self-Supervised Visual Learning

- Match the same image (with severe augmentation)
- Joint embedding approach: Apply loss on the embedding level.
- Use negative examples (contrastive) or not (non-contrastive).



Self-Supervised Visual Learning

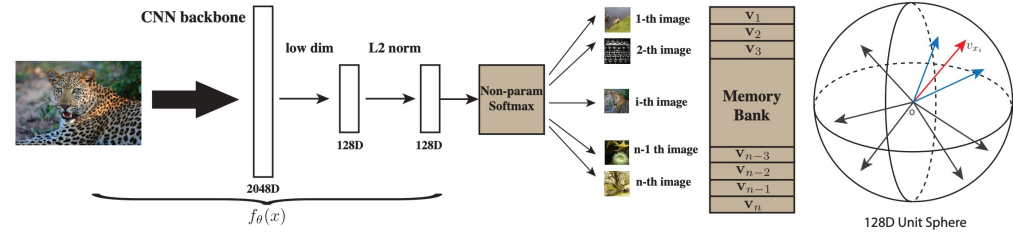
- Match the same image (with severe augmentation)
- Joint embedding approach: Apply loss on the embedding level.
- Use negative examples (contrastive) or not (non-contrastive).
- Energy is defined between a pair of images.



Several Embedding Loss Formulations

Wu et al., 2018

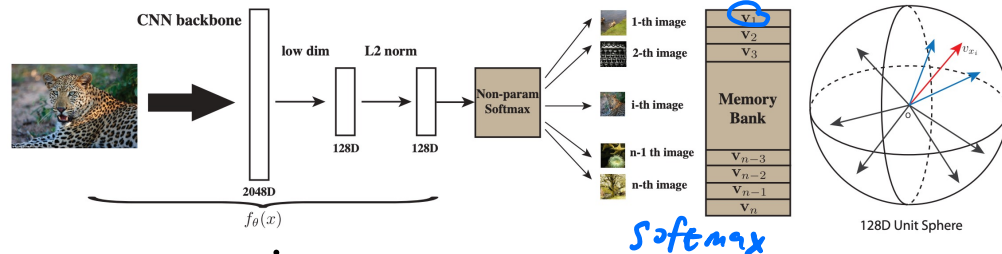
- Instance Classification:



Several Embedding Loss Formulations

Wu et al., 2018

- Instance Classification:



- Contrastive Learning: Cross entropy on pairs

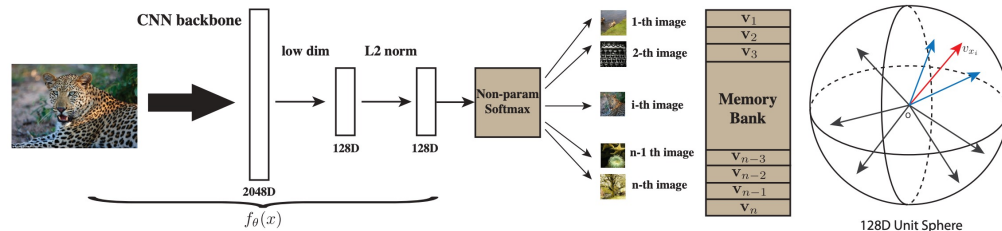
$$\ell_{i,j} = -\log \frac{\exp(\text{sim}(\mathbf{z}_i, \mathbf{z}_j)/\tau)}{\sum_{k=1}^{2N} \mathbb{1}_{[k \neq i]} \exp(\text{sim}(\mathbf{z}_i, \mathbf{z}_k)/\tau)}.$$

Chen et al. 2020

Several Embedding Loss Formulations

Wu et al., 2018

- Instance Classification:



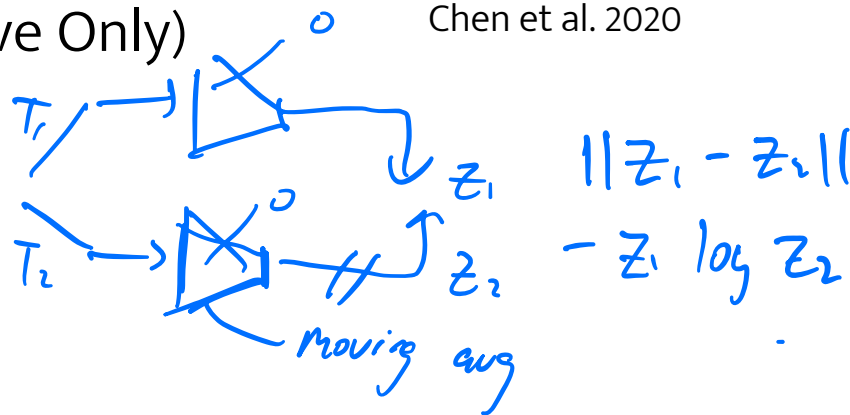
- Contrastive Learning: Cross entropy on pairs

$$\ell_{i,j} = -\log \frac{\exp(\text{sim}(\mathbf{z}_i, \mathbf{z}_j)/\tau)}{\sum_{k=1}^{2N} \mathbb{1}_{[k \neq i]} \exp(\text{sim}(\mathbf{z}_i, \mathbf{z}_k)/\tau)}.$$

Chen et al. 2020

- Non-contrastive Learning (Positive Only)

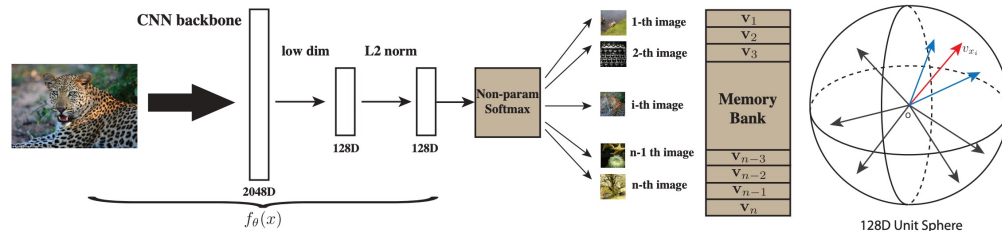
- Moving Average [Grill et al., 2020]
- Stop Gradient [Chen & He, 2020]



Several Embedding Loss Formulations

Wu et al., 2018

- Instance Classification:



- Contrastive Learning: Cross entropy on pairs

$$\ell_{i,j} = -\log \frac{\exp(\text{sim}(\mathbf{z}_i, \mathbf{z}_j)/\tau)}{\sum_{k=1}^{2N} \mathbb{1}_{[k \neq i]} \exp(\text{sim}(\mathbf{z}_i, \mathbf{z}_k)/\tau)}.$$

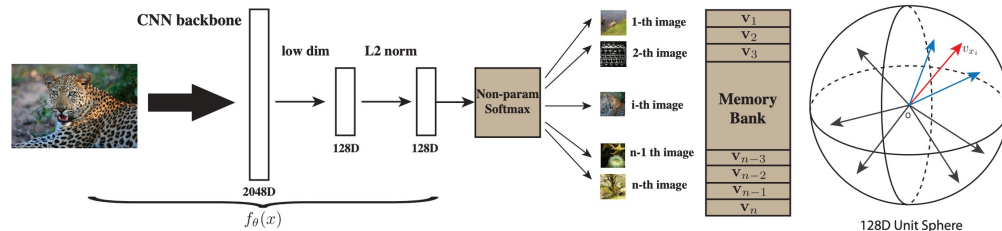
Chen et al. 2020

- Non-contrastive Learning (Positive Only)
 - Moving Average [Grill et al., 2020]
 - Stop Gradient [Chen & He, 2020]
- Use of projectors and predictors

Several Embedding Loss Formulations

Wu et al., 2018

- Instance Classification:



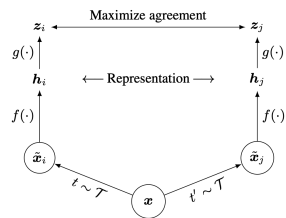
- Contrastive Learning: Cross entropy on pairs

$$\ell_{i,j} = -\log \frac{\exp(\text{sim}(\mathbf{z}_i, \mathbf{z}_j)/\tau)}{\sum_{k=1}^{2N} \mathbb{1}_{[k \neq i]} \exp(\text{sim}(\mathbf{z}_i, \mathbf{z}_k)/\tau)}.$$

Chen et al. 2020

- Non-contrastive Learning (Positive Only)
 - Moving Average [Grill et al., 2020]
 - Stop Gradient [Chen & He, 2020]

- Use of projectors and predictors

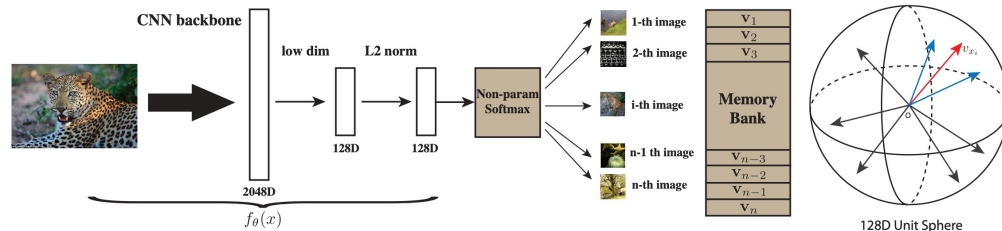


Chen et al. 2020

Several Embedding Loss Formulations

Wu et al., 2018

- Instance Classification:



- Contrastive Learning: Cross entropy on pairs

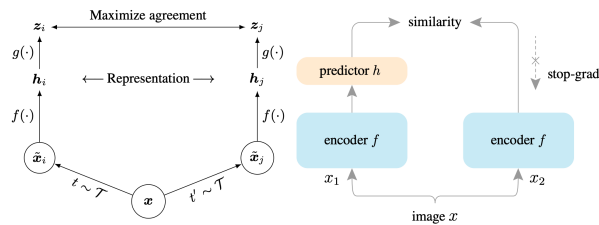
$$\ell_{i,j} = -\log \frac{\exp(\text{sim}(\mathbf{z}_i, \mathbf{z}_j)/\tau)}{\sum_{k=1}^{2N} \mathbb{1}_{[k \neq i]} \exp(\text{sim}(\mathbf{z}_i, \mathbf{z}_k)/\tau)}.$$

Chen et al. 2020

- Non-contrastive Learning (Positive Only)

- Moving Average [Grill et al., 2020]
- Stop Gradient [Chen & He, 2020]

- Use of projectors and predictors

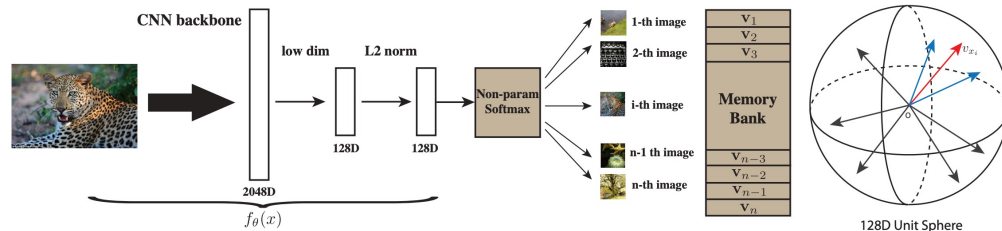


Chen et al. 2020 Chen & He, 2021

Several Embedding Loss Formulations

Wu et al., 2018

- Instance Classification:



- Contrastive Learning: Cross entropy on pairs

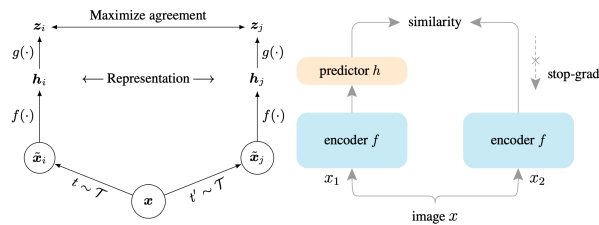
$$\ell_{i,j} = -\log \frac{\exp(\text{sim}(z_i, z_j)/\tau)}{\sum_{k=1}^{2N} \mathbb{1}_{[k \neq i]} \exp(\text{sim}(z_i, z_k)/\tau)}.$$

Chen et al. 2020

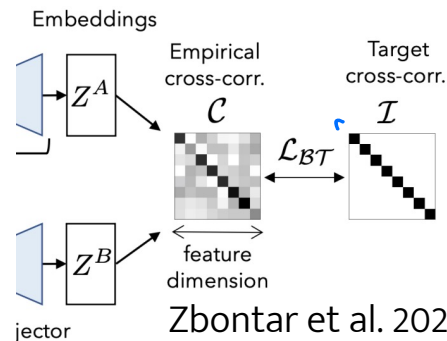
- Non-contrastive Learning (Positive Only)

- Moving Average [Grill et al., 2020]
- Stop Gradient [Chen & He, 2020]

- Use of projectors and predictors
- Use of co-variance regularization



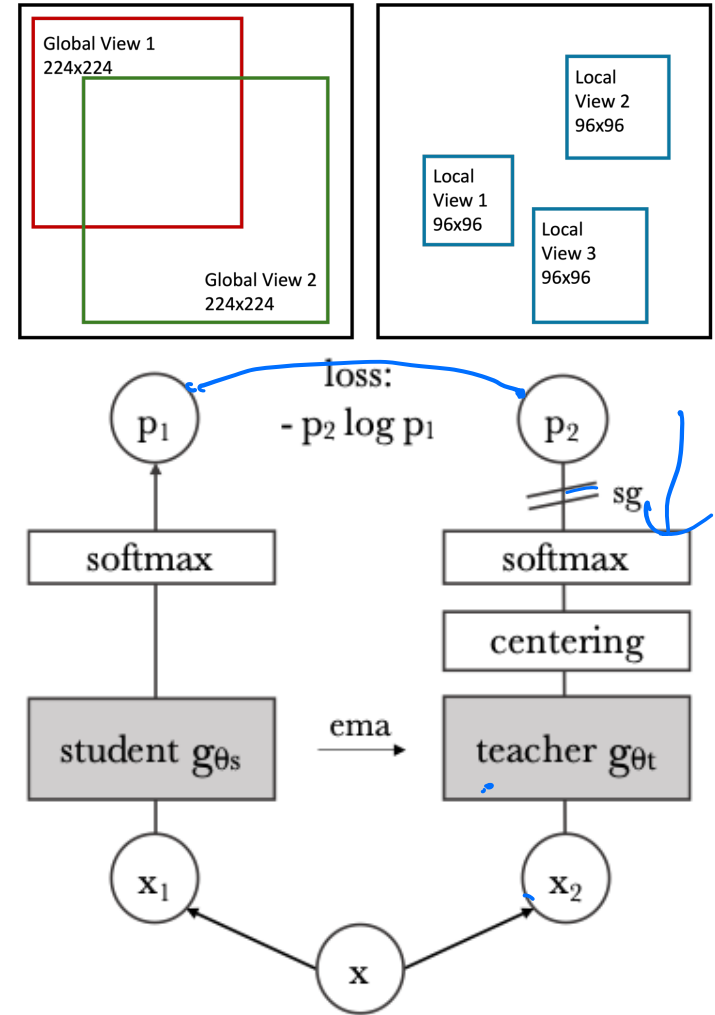
Chen et al. 2020 Chen & He, 2021



Zbontar et al. 2021
Bardes et al. 2021

DINO

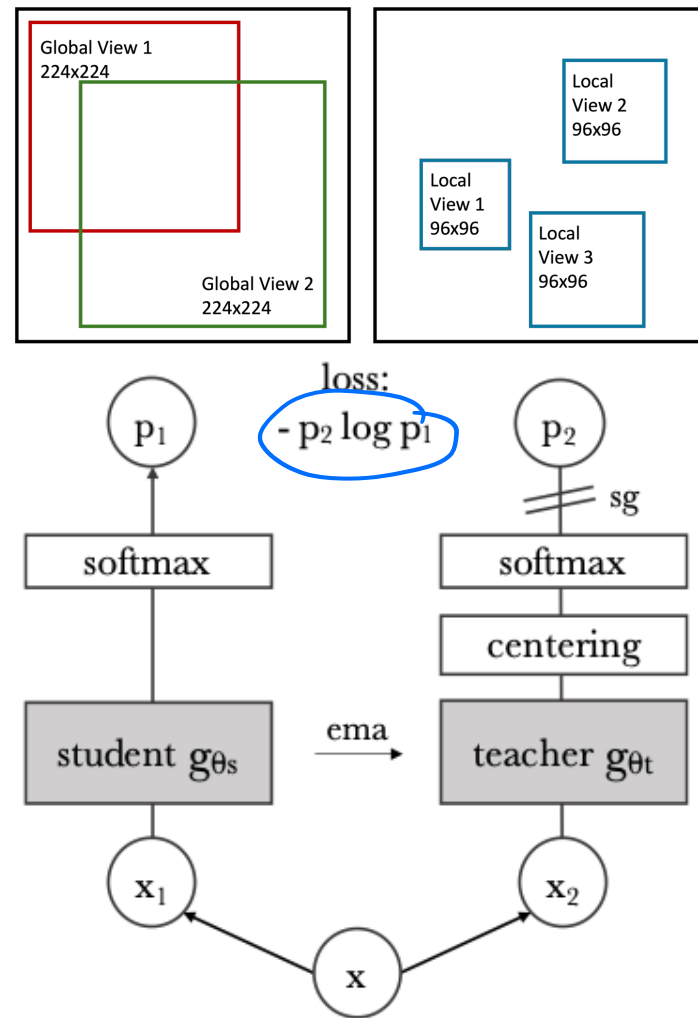
- Knowledge distillation between a student and a teacher network.



DINO

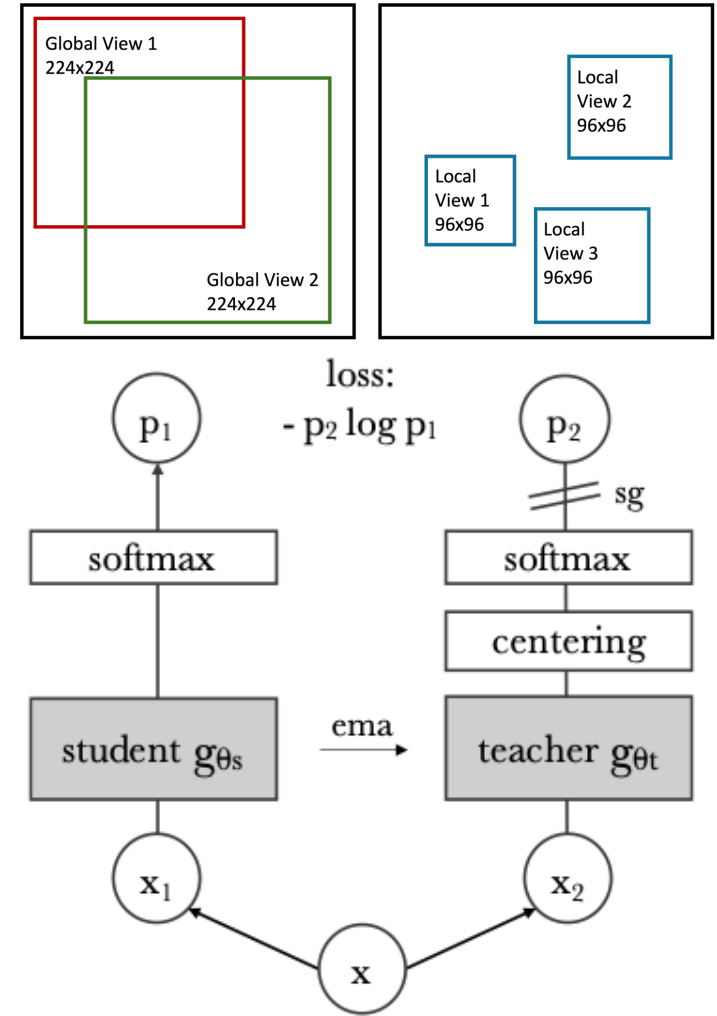
- Knowledge distillation between a student and a teacher network.

- Student: $p_s(x) = \frac{\exp(g_{\theta_s}(x)_i / \tau_s)}{\sum_k \exp(g_{\theta_s}(x)_k / \tau_s)}$



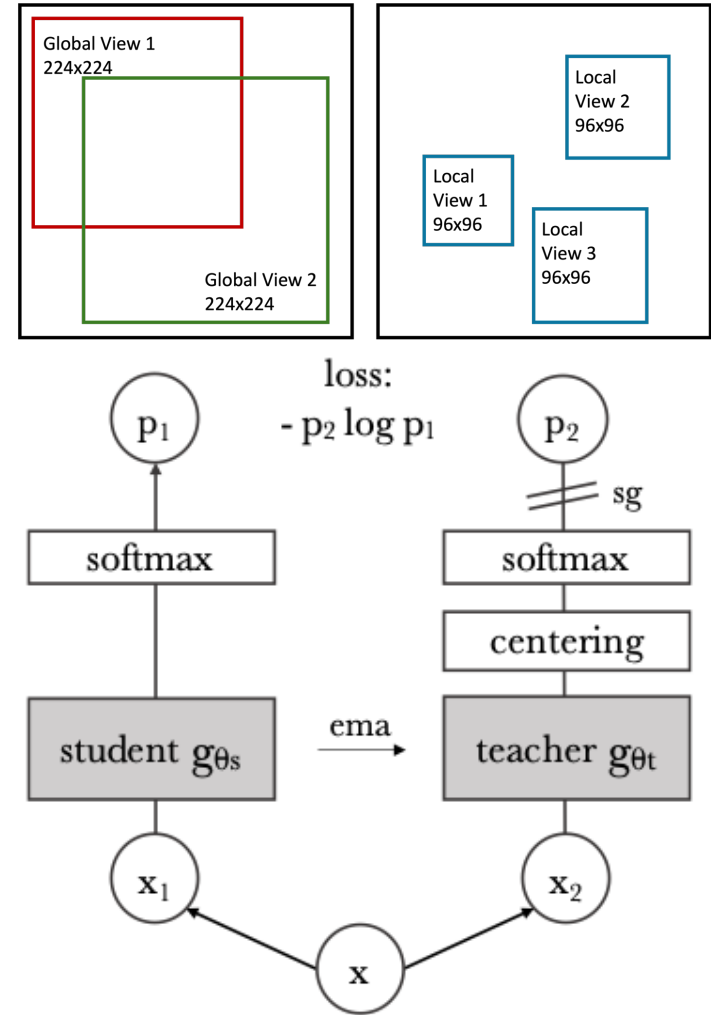
DINO

- Knowledge distillation between a student and a teacher network.
- Student: $p_s(x) = \frac{\exp(g_{\theta_s}(x)_i / \tau_s)}{\sum_k \exp(g_{\theta_s}(x)_k / \tau_s)}$.
- Minimize CE: $\min_{\theta} H(\underbrace{p_t(x)}, \underbrace{p_s(x)})$.



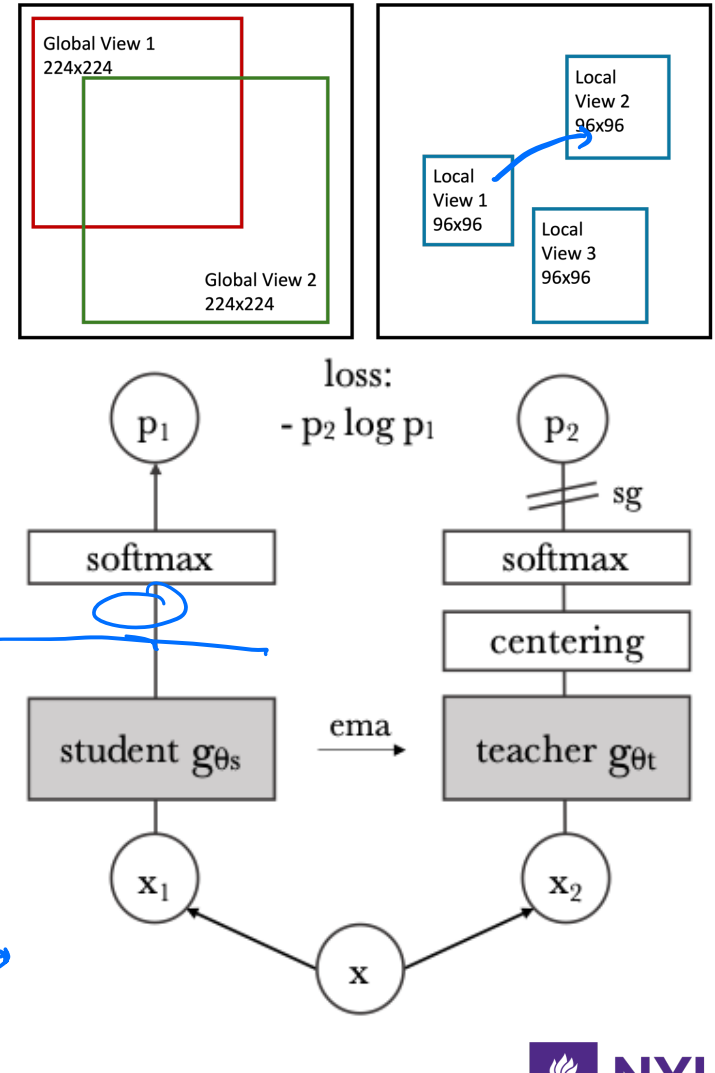
DINO

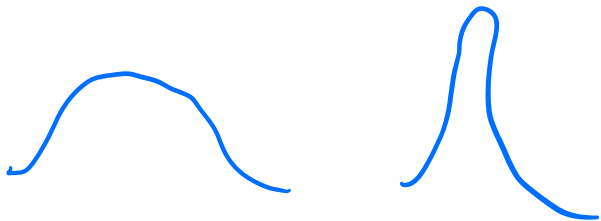
- Knowledge distillation between a student and a teacher network.
- Student: $p_s(x) = \frac{\exp(g_{\theta_s}(x)_i / \tau_s)}{\sum_k \exp(g_{\theta_s}(x)_k / \tau_s)}$.
- Minimize CE: $\min_{\theta} H(p_t(x), p_s(x))$.
- Stop gradient on the teacher (no true label).



DINO

- Knowledge distillation between a student and a teacher network.
- Student: $p_s(x) = \frac{\exp(g_{\theta_s}(x)_i / \tau_s)}{\sum_k \exp(g_{\theta_s}(x)_k / \tau_s)}$.
- Minimize CE: $\min_{\theta} H(p_t(x), p_s(x))$.
- Stop gradient on the teacher (no true label)
- Teacher network has EMA weights copied from student (prevent collapse).



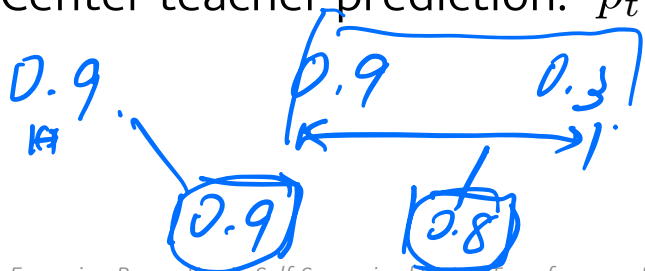


Preventing Collapse

- Cross entropy objective can make both sides collapse to uniform distribution.
 - Apply sharpening, apply a temperature term on both teacher and student.
 - $\text{softmax}(g/\tau)$ The higher the temperature, the more uniform.

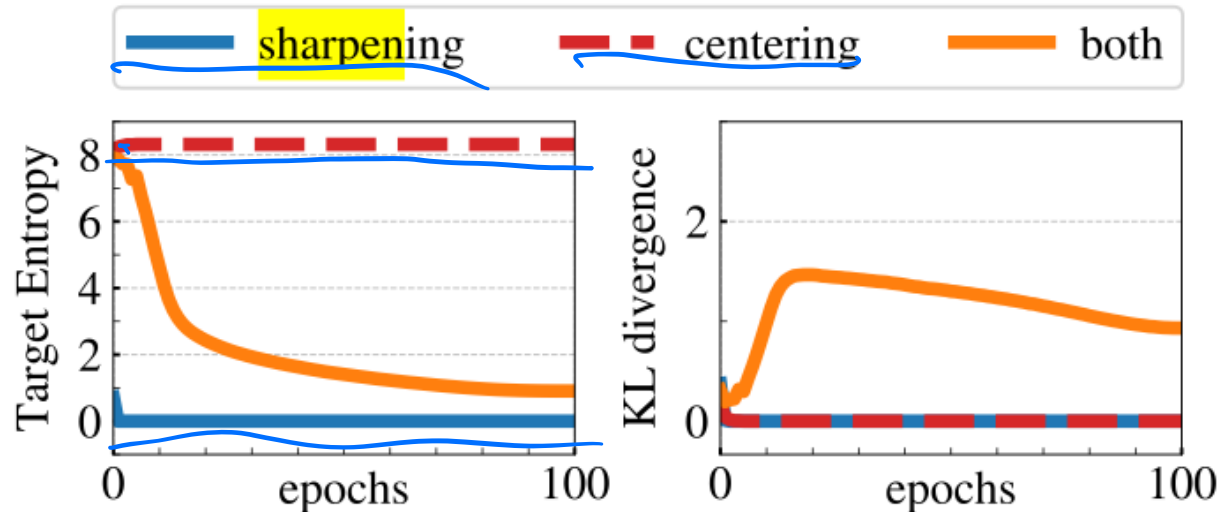
Preventing Collapse

- Cross entropy objective can make both sides collapse to uniform distribution.
 - Apply sharpening, apply a temperature term on both teacher and student.
 - $\text{softmax}(g/\tau)$ The higher the temperature, the more uniform.
- It can also collapse into always activating a single unit.
 - Mean statistics: $c_t = mc_{t-1} + (1 - m) \frac{1}{B} \sum_{i=1}^B g_{\theta_t}(x_i)$
 - Center teacher prediction: $p_t(x) = \frac{\exp((g_{\theta_t}(x)_i - c_t)/\tau_t)}{\sum_k \exp((g_{\theta_t}(x)_k - c_t)/\tau_t)}$



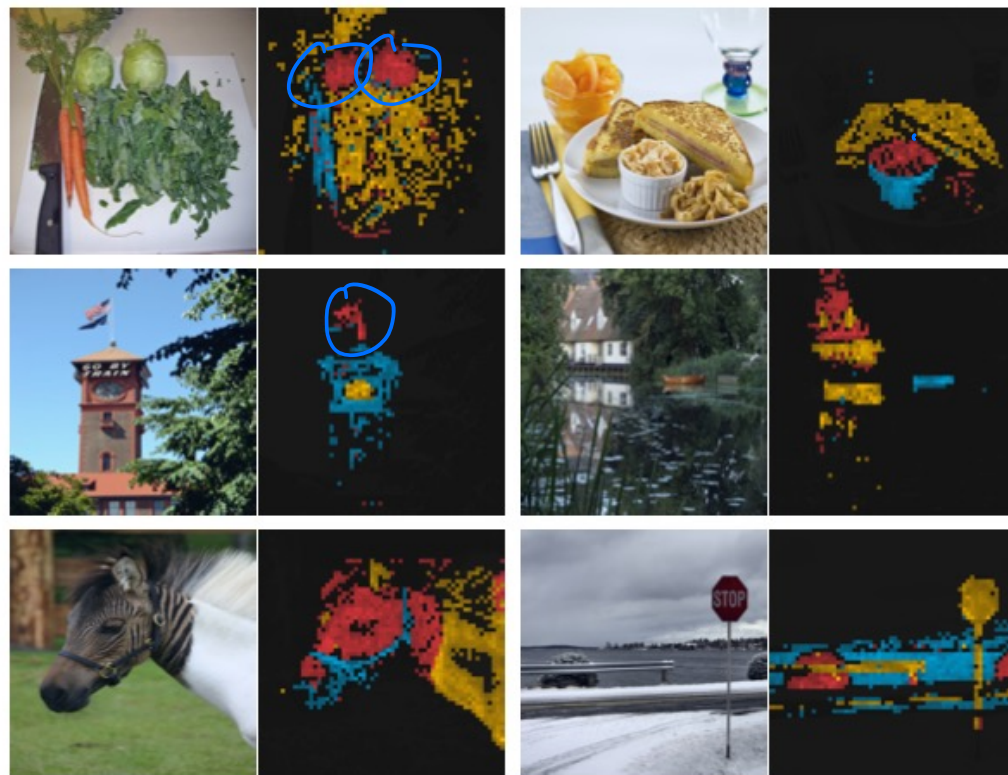
Centering and Sharpening

- Only centering: Always uniform distribution, high entropy, easy to guess.
- Only sharpening: Collapsed into one unit, easy to guess, low loss, but no real learning.



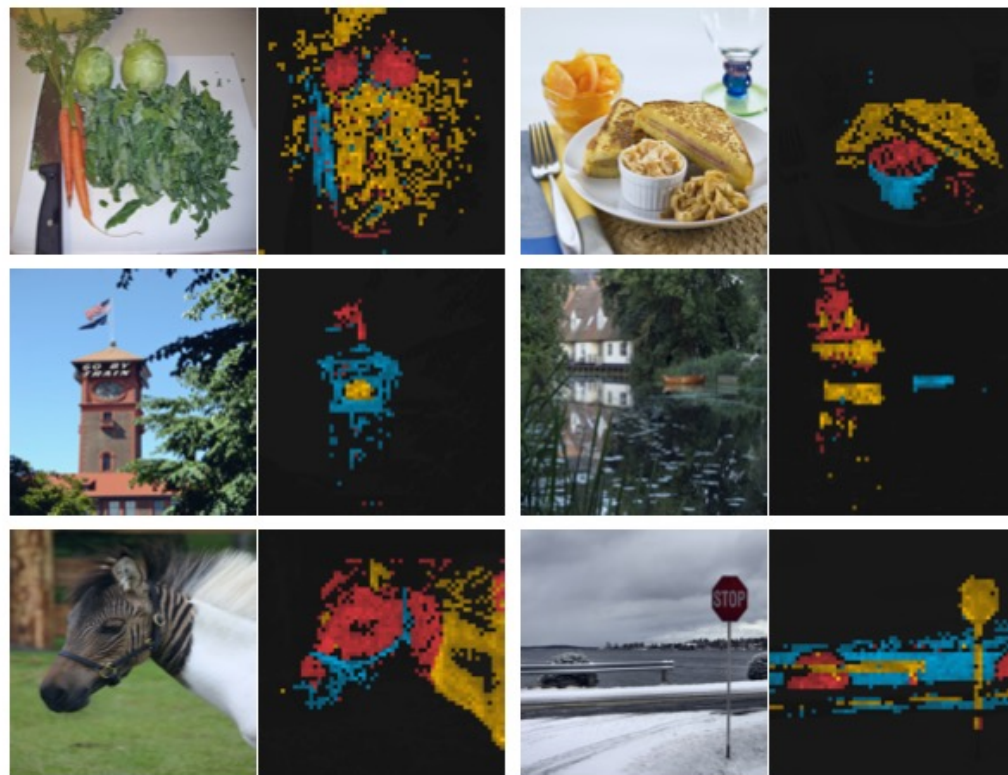
Visualizing Attention

- The [CLS] token is an extra token added to summarize the whole image into a vector.



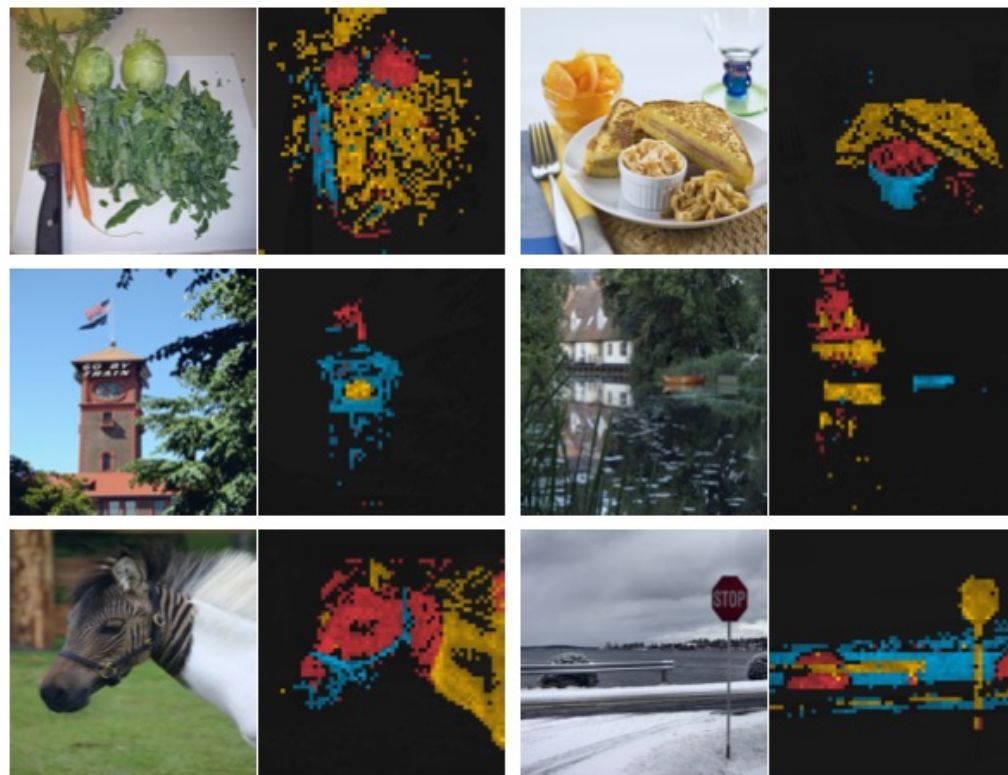
Visualizing Attention

- The [CLS] token is an extra token added to summarize the whole image into a vector.
- Visualize the attention map of different attention heads using different colors.



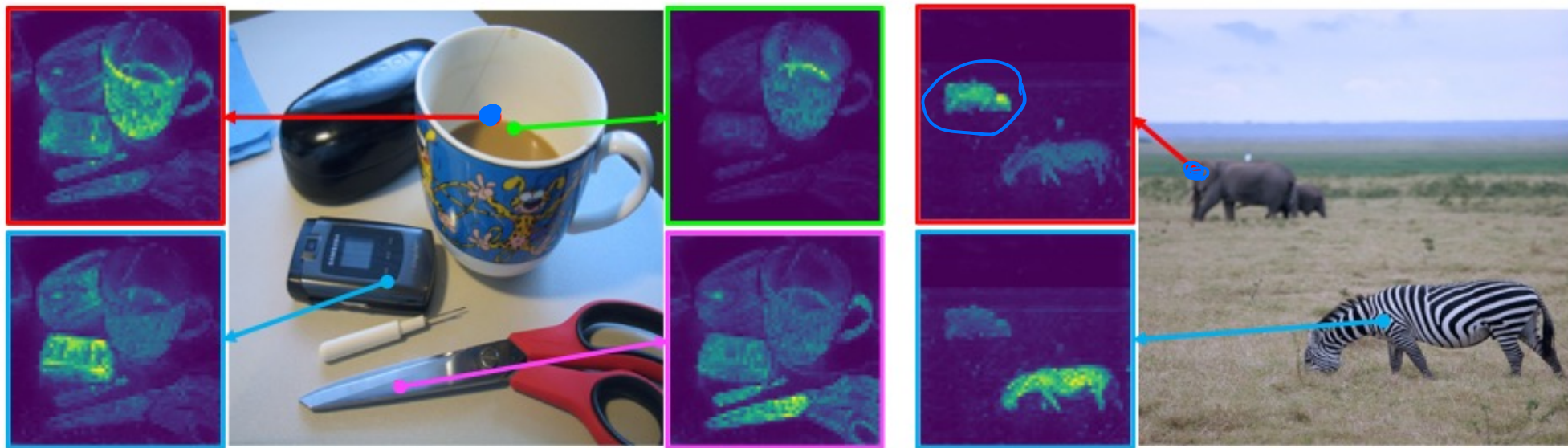
Visualizing Attention

- The [CLS] token is an extra token added to summarize the whole image into a vector.
- Visualize the attention map of different attention heads using different colors.
- Showing understanding of different objects and parts.



Visualizing Attention

- We can also visualize the attention by querying from a location.
- Weak separation of objects.

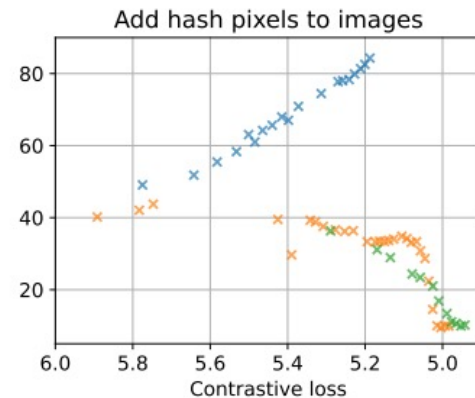
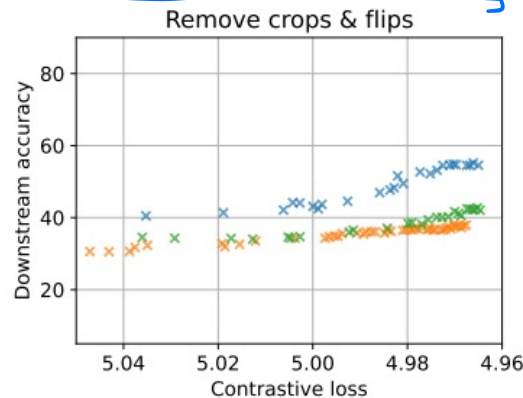
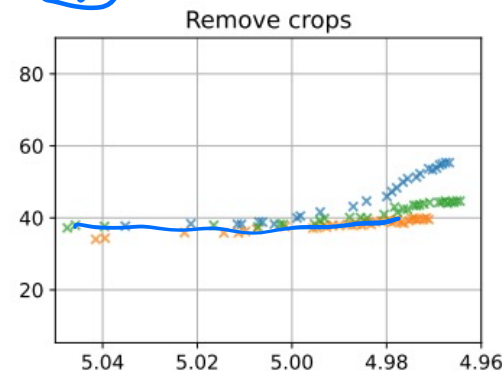
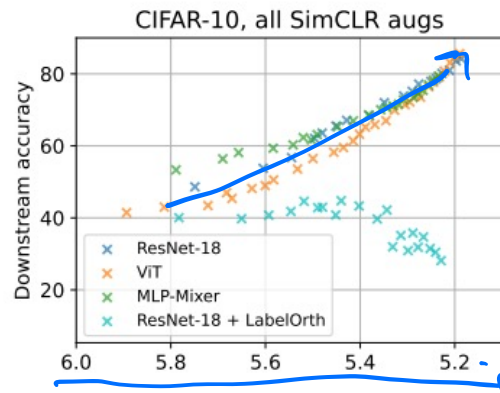


$z_i = z_i \rightarrow$ classification.
 T_1 | T_2

Why Does SSL Work?

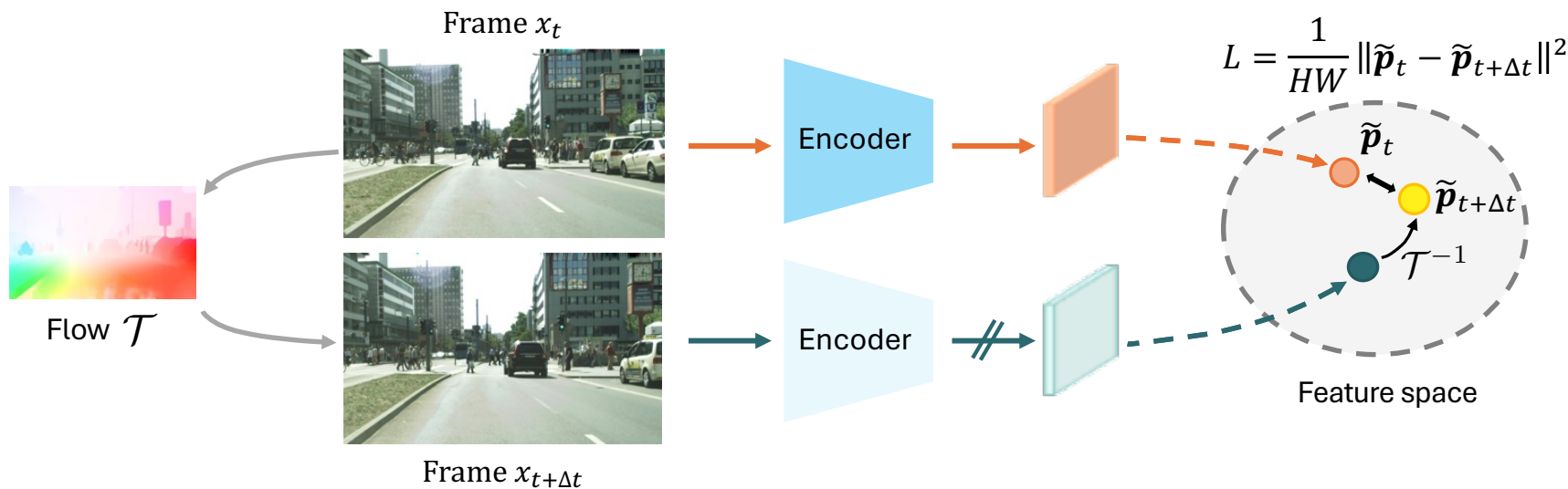


- The unsupervised loss is a surrogate. If an image belongs to a similarity class, it also belongs to the same semantic class.
- The choice of similarity class matters.



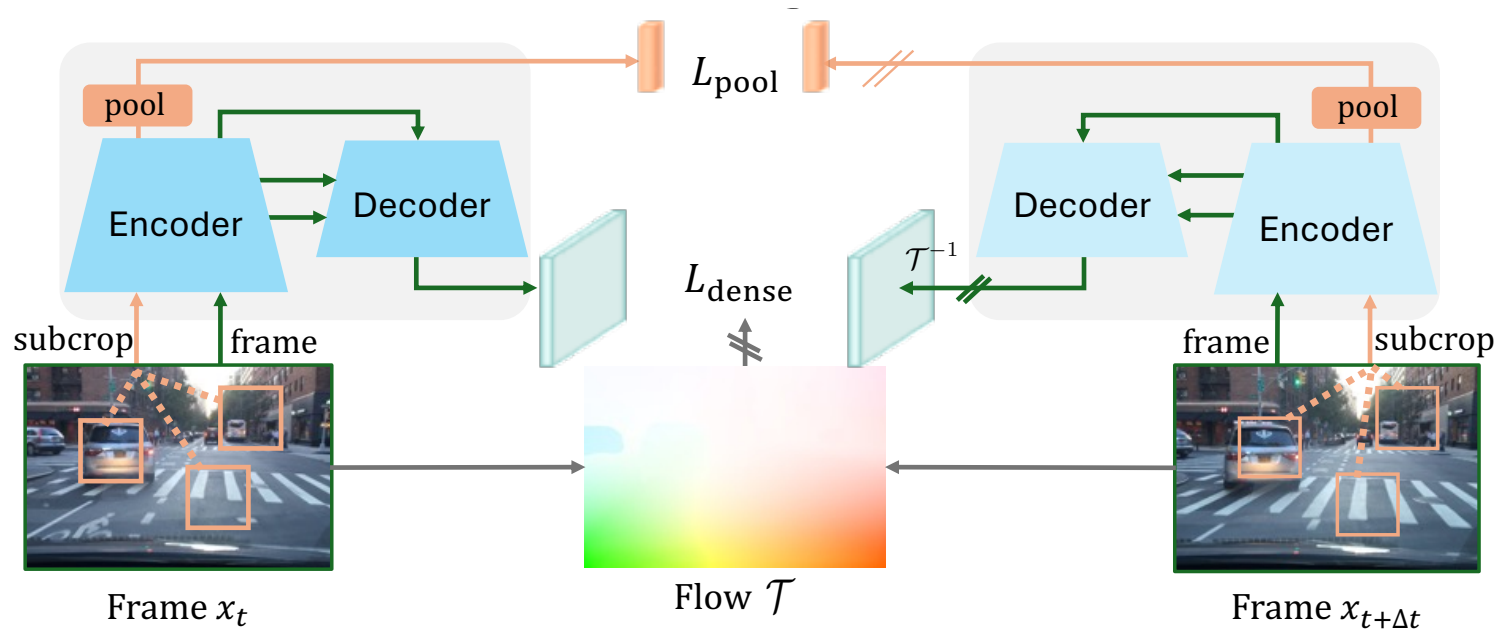
SSL with Motion

- Can we use adjacent frames as self-supervision?
- Objects move densely throughout the image.



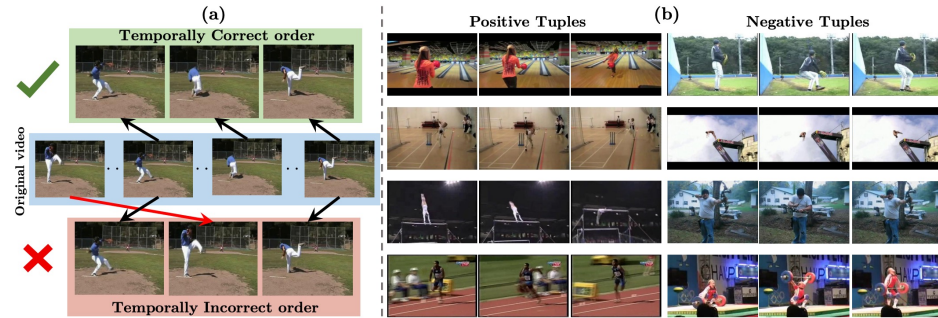
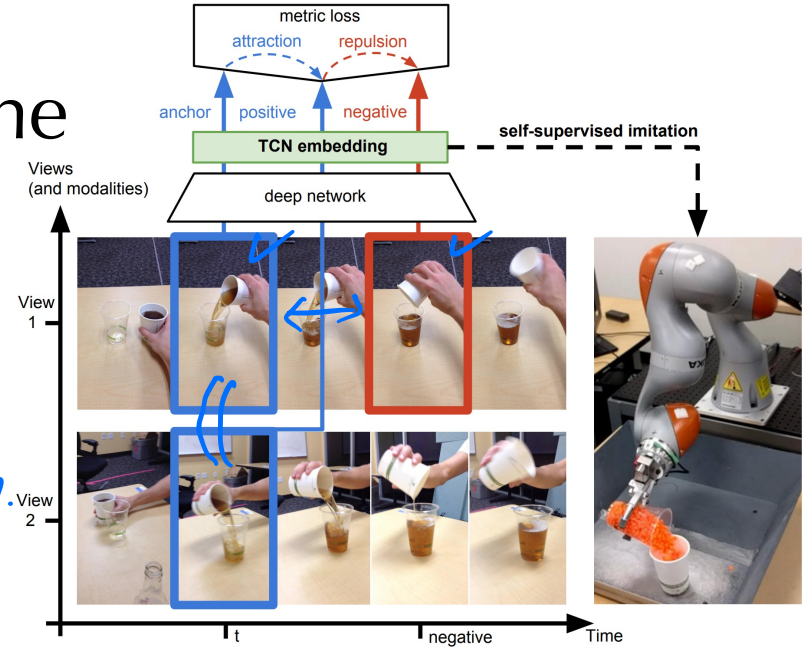
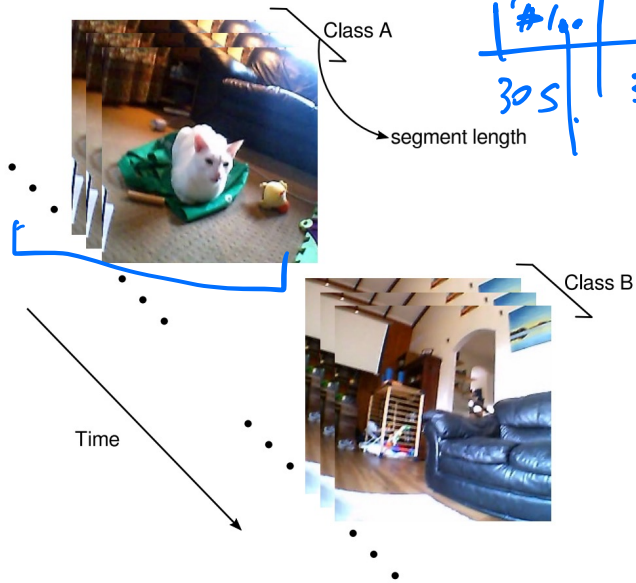
SSL with Motion

- Perform SSL in multiple scales (small objects vs. big regions).



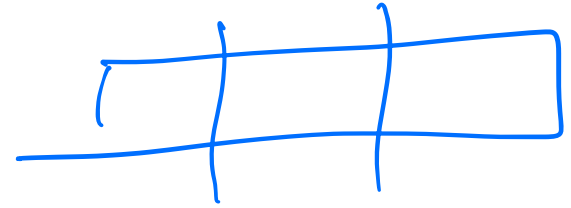
SSL with Time

- Use time as an additional source of supervision.

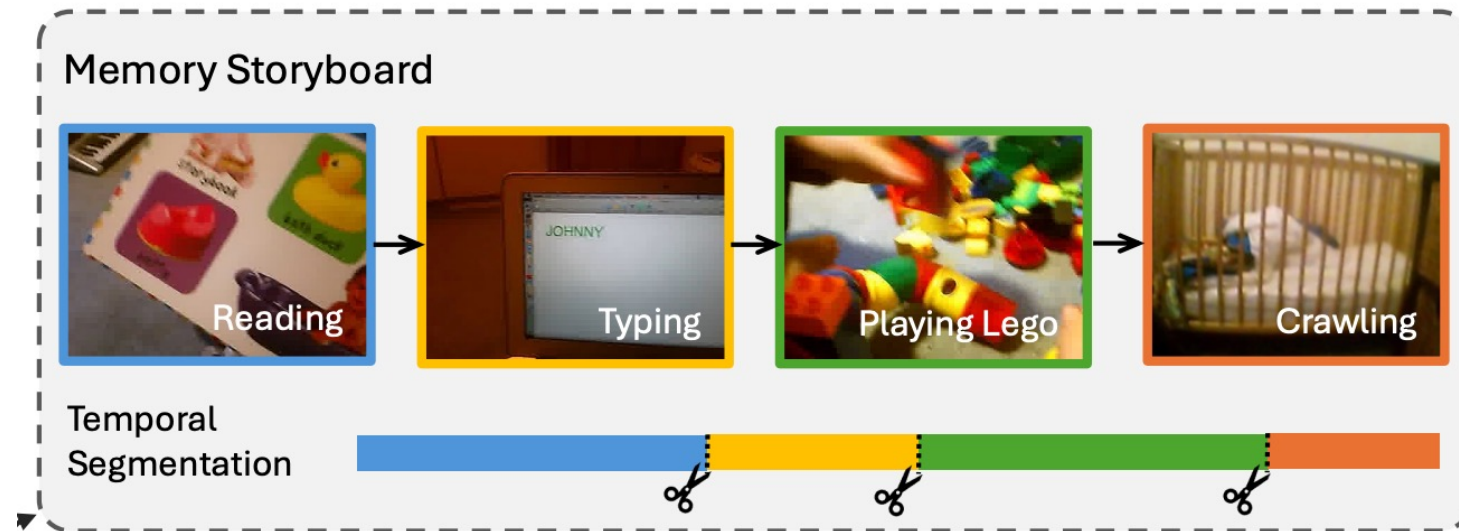
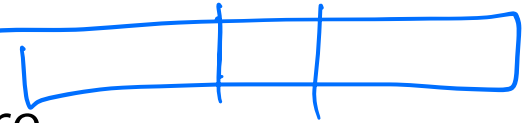


Misra et al. *Shuffle and Learn: Unsupervised Learning using Temporal Order Verification*. ECCV 2016.
 Sermanet et al. *Time-Contrastive Networks: Self-Supervised Learning from Video*. ICRA 2018.
 Orhan et al. *Self-Supervised Learning through the Eyes of a Child*. NeurIPS 2020.

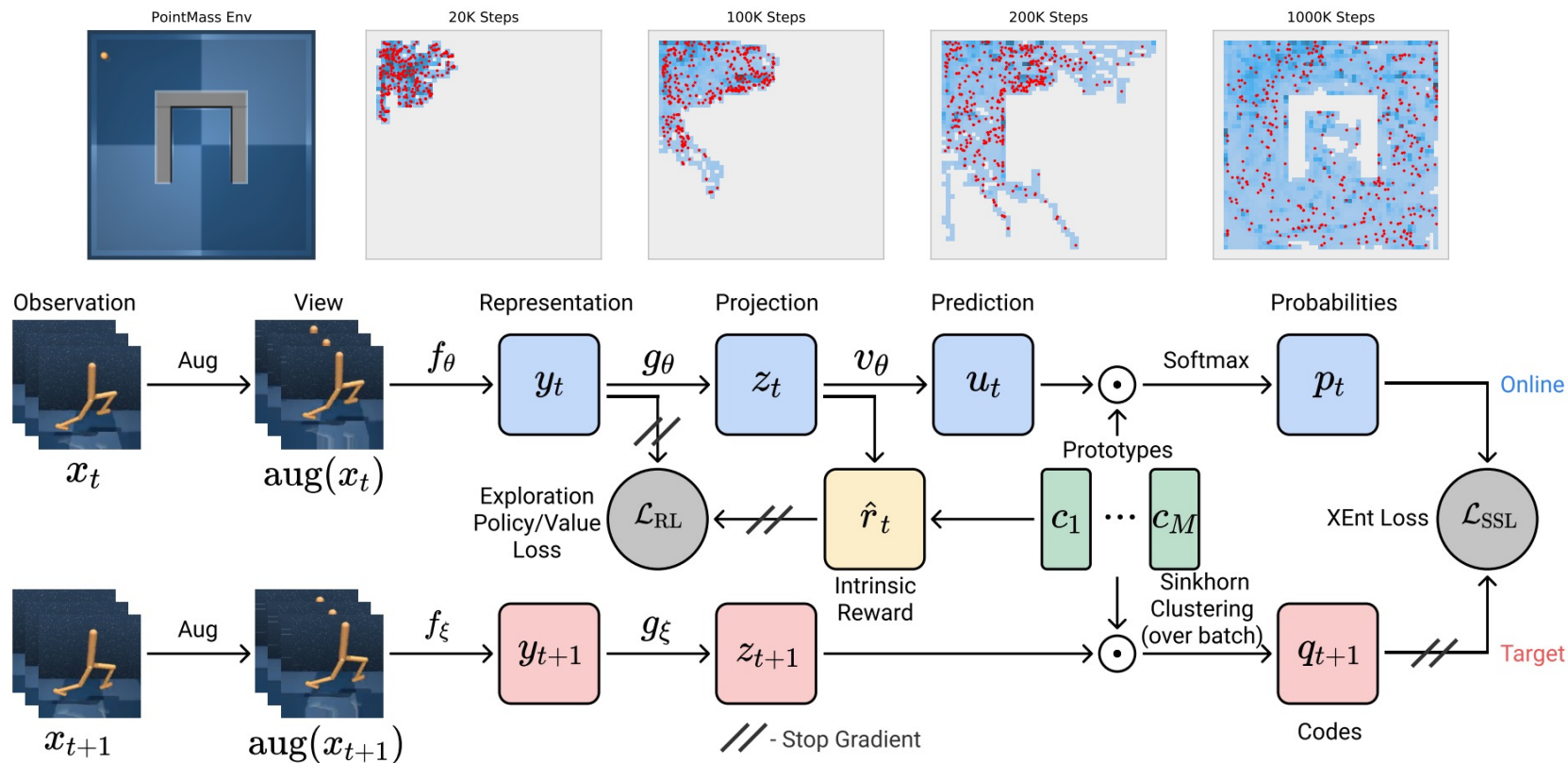
SSL with Time



- We can segment videos into meaningful events.
- Leverage the spatiotemporal continuity structure.



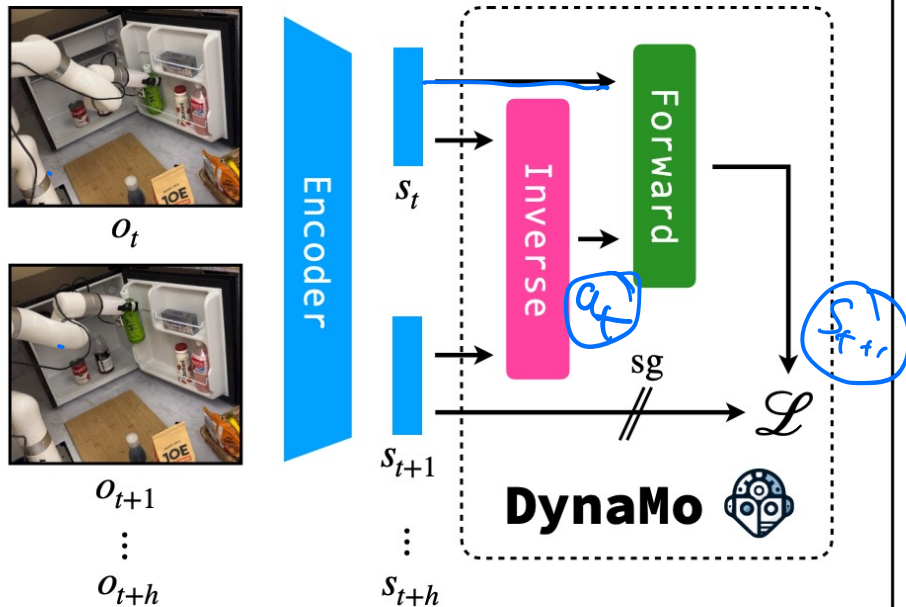
SSL for Visual Control



SSL for Visual Control

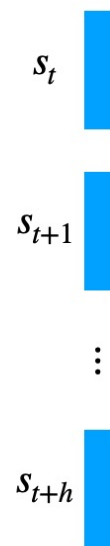
Not Augment just video frames.

Observations



(a) Representation learning

Embeddings



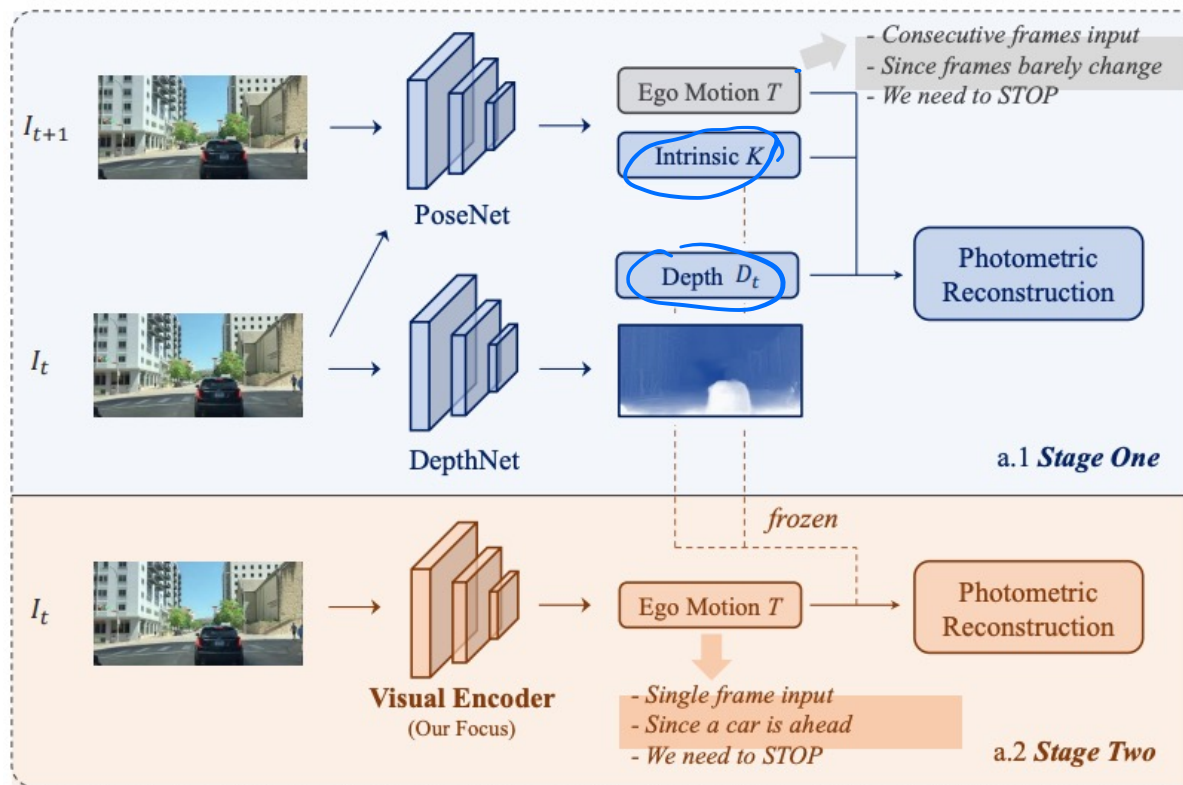
Environments



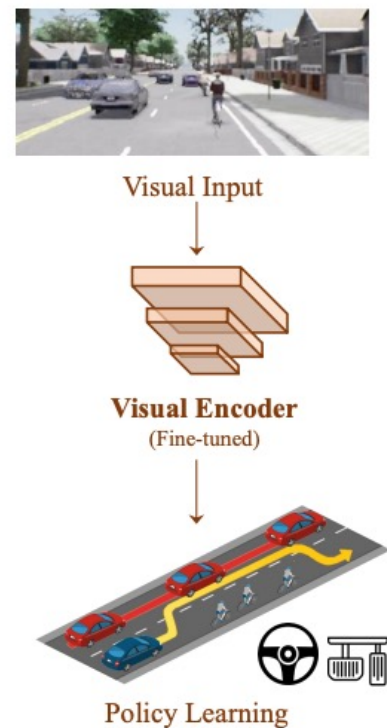
(b) Policy on pretrained representations

SSL for Visual Control

(a) Self-supervised Visuomotor Policy Pre-training

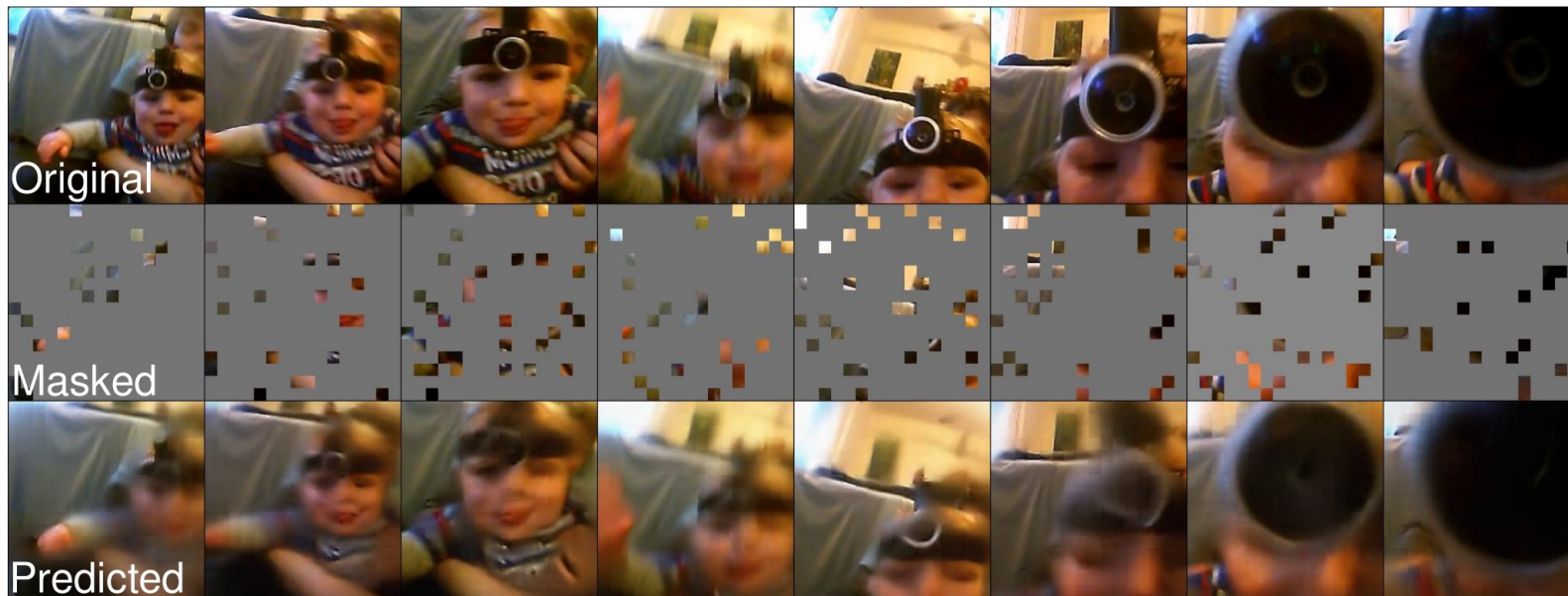
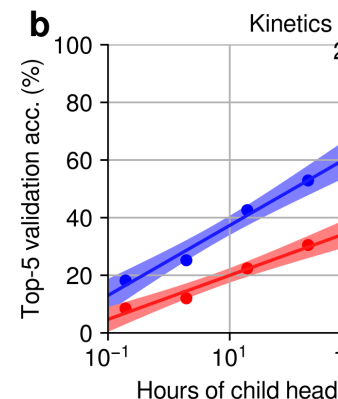
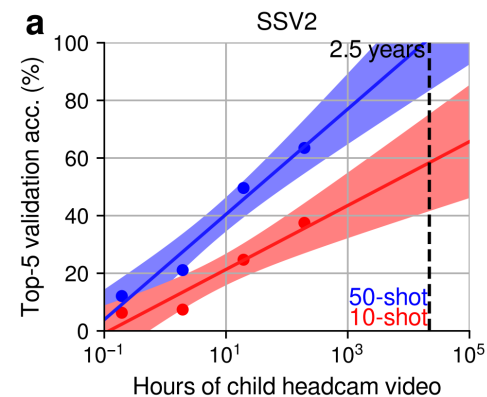


(b) Downstream Tasks



BabyCam

- Run visual learning algorithms on baby headcam videos.



Summary

- Representation learning leverage the information in unlabeled data.

Summary

- Representation learning leverage the information in unlabeled data.
- A foundation for sensorimotor learning.

Summary

- Representation learning leverage the information in unlabeled data.
- A foundation for sensorimotor learning.
- Inductive biases matter.

Summary

- Representation learning leverage the information in unlabeled data.
- A foundation for sensorimotor learning.
- Inductive biases matter.
- Possible learning objectives for egocentric videos.

Summary

- Representation learning leverage the information in unlabeled data.
- A foundation for sensorimotor learning.
- Inductive biases matter.
- Possible learning objectives for egocentric videos.
- Incorporate 3D vision and actions for downstream planning.

Emergent Attention, Object Discovery

- SSL representations show awareness of object classes and instance identities.

Emergent Attention, Object Discovery

- SSL representations show awareness of object classes and instance identities.
- Why does attention show awareness of objects?

Emergent Attention, Object Discovery

- SSL representations show awareness of object classes and instance identities.
- Why does attention show awareness of objects?
- The network is encouraged to associate different parts of the objects together in order to identify whether two inputs belong to the same image or not.

Emergent Attention, Object Discovery

- SSL representations show awareness of object classes and instance identities.
- Why does attention show awareness of objects?
- The network is encouraged to associate different parts of the objects together in order to identify whether two inputs belong to the same image or not.
- Attending to semantically similar parts facilitates the process.

Emergent Attention, Object Discovery

- SSL representations show awareness of object classes and instance identities.
- Why does attention show awareness of objects?
- The network is encouraged to associate different parts of the objects together in order to identify whether two inputs belong to the same image or not.
- Attending to semantically similar parts facilitates the process.
- The network is a hierarchical information processing pipeline – Lower layers integrate more granular and smaller neighborhood.

Weak-to-Strong Supervision

- General idea: Use self-supervised learning to learn good features, which allow us to generate low-quality masks.

Weak-to-Strong Supervision

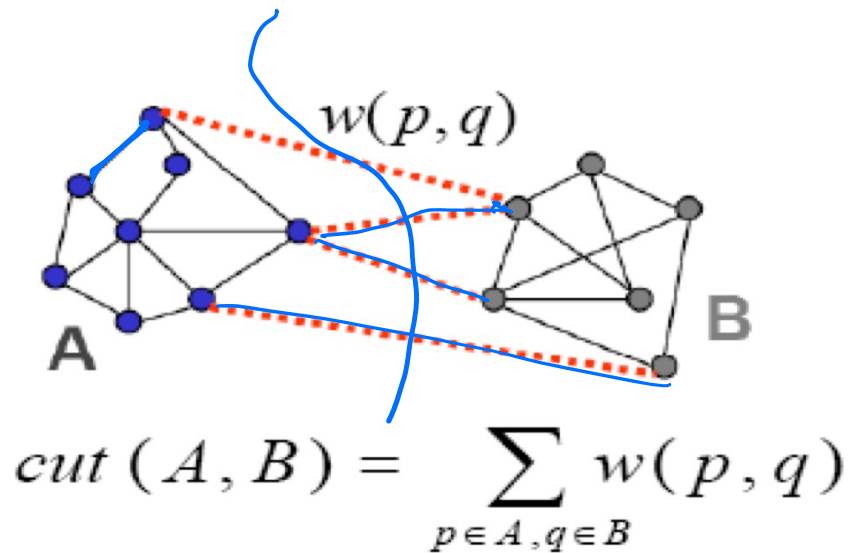
- General idea: Use self-supervised learning to learn good features, which allow us to generate low-quality masks.
- Then use these masks as pseudo labels and supervise the network to predict these low-quality masks.

Weak-to-Strong Supervision

- General idea: Use self-supervised learning to learn good features, which allow us to generate low-quality masks.
- Then use these masks as pseudo labels and supervise the network to predict these low-quality masks.
- Question: how do we come up with masks? What loss is used to supervise the network?

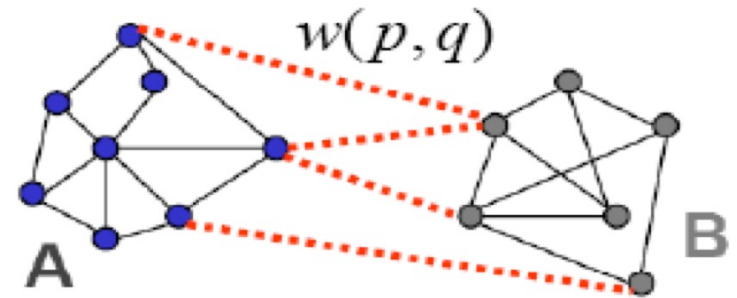
Graph Cut

- Segmentation is essentially a clustering problem.



Graph Cut

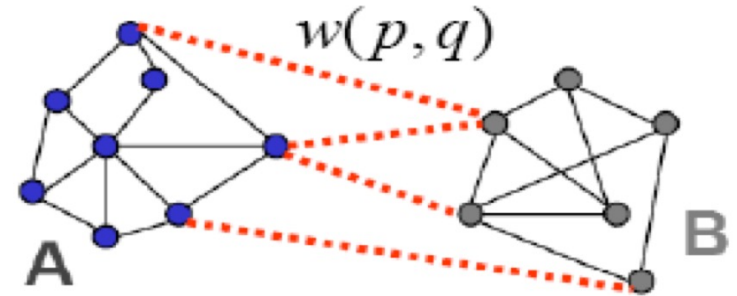
- Segmentation is essentially a clustering problem.
- We can transform the clustering problem with the graph cut problem.



$$cut(A, B) = \sum_{p \in A, q \in B} w(p, q)$$

Graph Cut

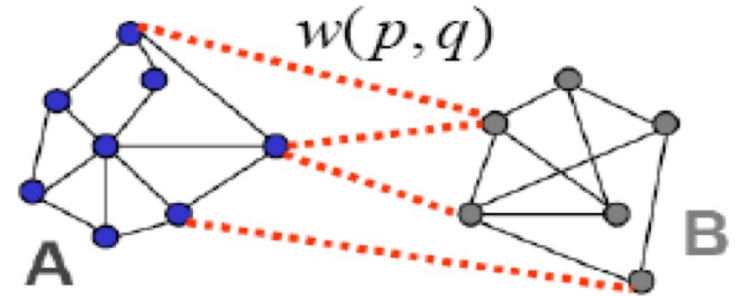
- Segmentation is essentially a clustering problem.
- We can transform the clustering problem with the graph cut problem.
- Pixel = node.



$$cut(A, B) = \sum_{p \in A, q \in B} w(p, q)$$

Graph Cut

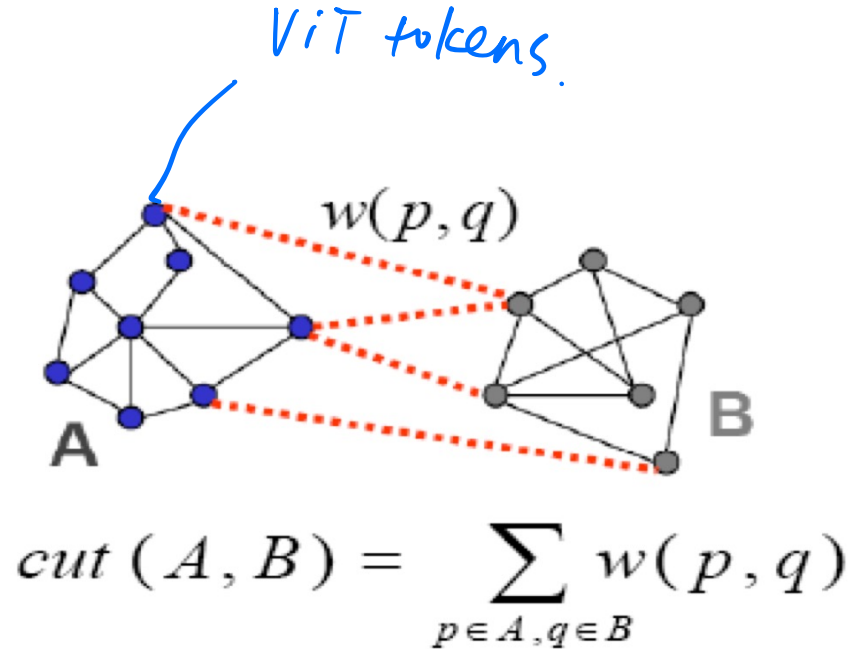
- Segmentation is essentially a clustering problem.
- We can transform the clustering problem with the graph cut problem.
- Pixel = node.
- Affinity between the two pixels = edge value (flow).



$$cut(A, B) = \sum_{p \in A, q \in B} w(p, q)$$

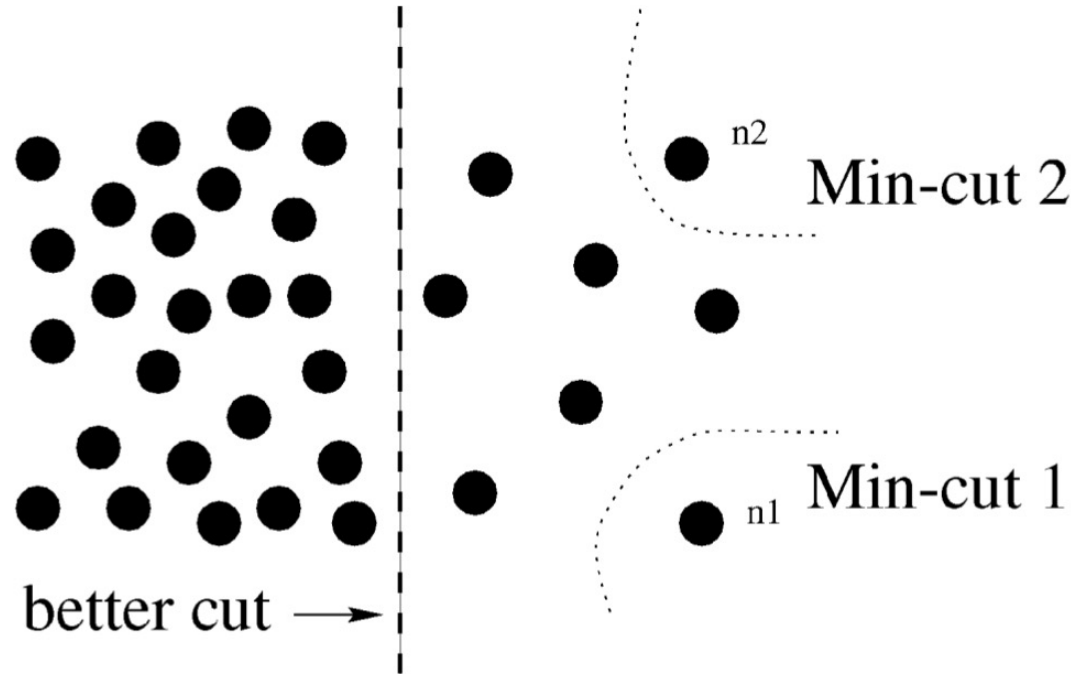
Graph Cut

- Segmentation is essentially a clustering problem.
- We can transform the clustering problem with the graph cut problem.
- Pixel = node.
- Affinity between the two pixels = edge value (flow).
- Objective: Cut the graph into disconnected components with a minimum sum of edge values.



Normalized Graph Cut (NCut)

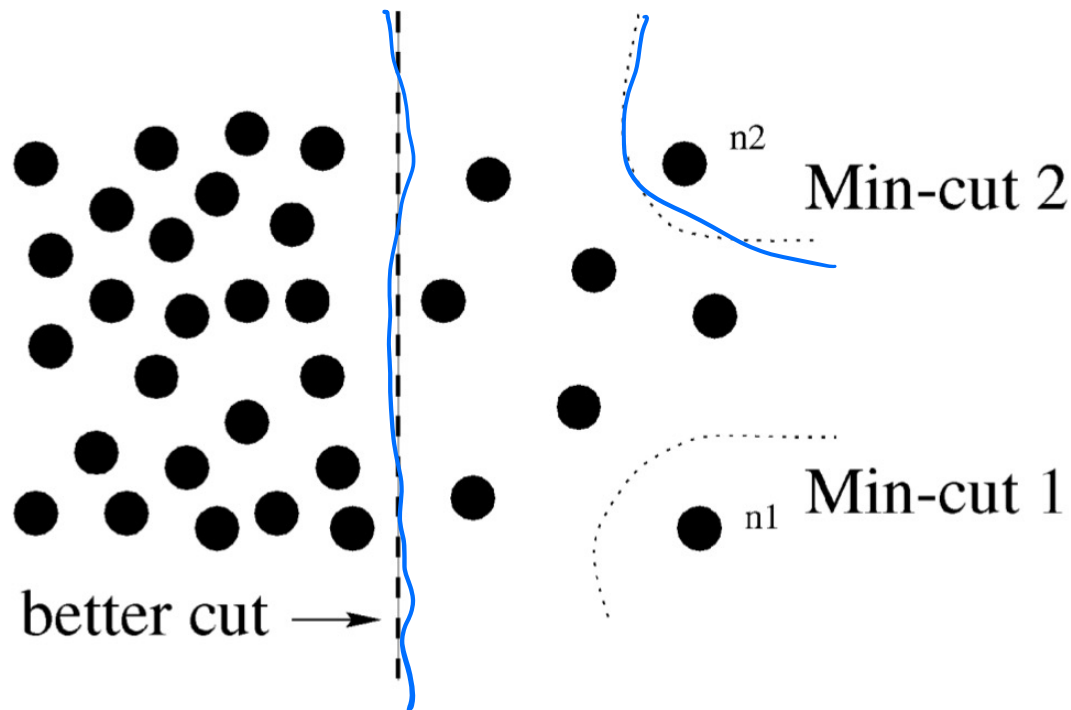
- How to prevent cutting small isolated nodes?



Normalized Graph Cut (NCut)

- How to prevent cutting small isolated nodes?
- Normalize by the total edge connections of a group to all the nodes.

$$Ncut(A, B) = \frac{cut(A, B)}{assoc(A, V)} + \frac{cut(A, B)}{assoc(B, V)}$$



NCut Details (Optional)

- A form of spectral clustering.
- Degree matrix D $N \times N$ with d_i on the diagonal.
- Weight matrix W $N \times N$ symmetric w_{ij} .
- Selection vector $x_i = 1$ if $i \in A$ otherwise -1 .
- Solve the minimization: $\min_y \frac{y^\top (D - W)y}{y^\top D y}$ $y = (1 + x) - \frac{\sum_{i|x_i > 0} d_i}{\sum_{i|x_i < 0} d_i} (1 - x)$.
- Generalized eigenvalue system: $(D - W)y = \lambda D y$.
- Let $z = D^{1/2}y$ $D^{-\frac{1}{2}}(D - W)D^{-\frac{1}{2}}z = \lambda z$.

NCut

- Sort the eigenvectors from the smallest to the largest.



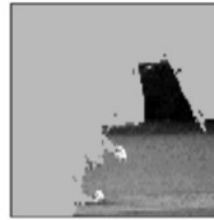
(a)



(b)



(c)



(d)



(e)



(f)



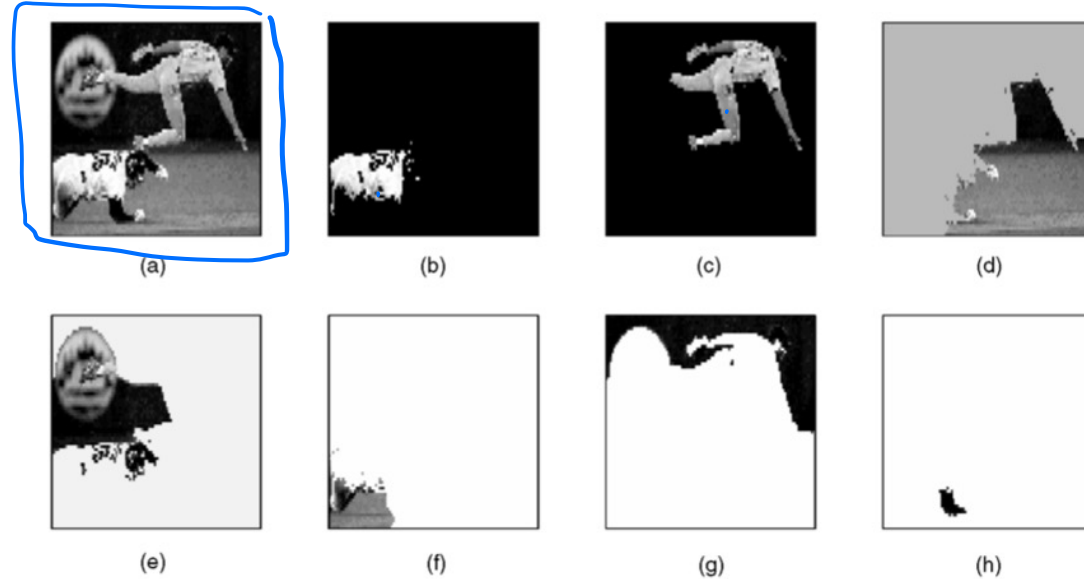
(g)



(h)

NCut

- Sort the eigenvectors from the smallest to the largest.
- This was a classic image segmentation technique operating directly on image intensity.



NCut

- Sort the eigenvectors from the smallest to the largest.
- This was a classic image segmentation technique operating directly on image intensity.
- Now, instead of segmenting pixels, we can directly segment semantically meaningful representations from self-supervision.



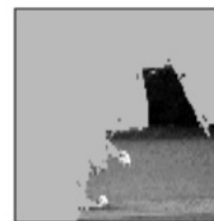
(a)



(b)



(c)



(d)



(e)



(f)



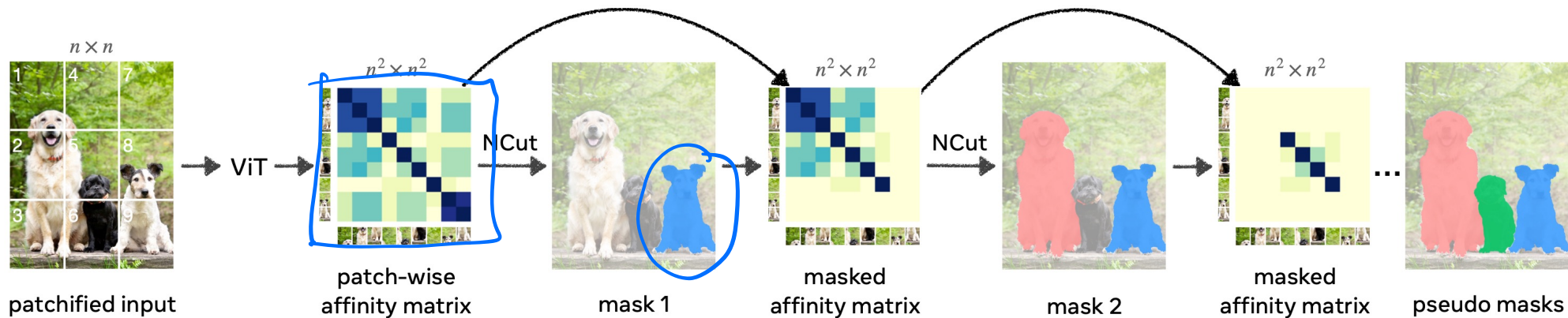
(g)



(h)

MaskCut

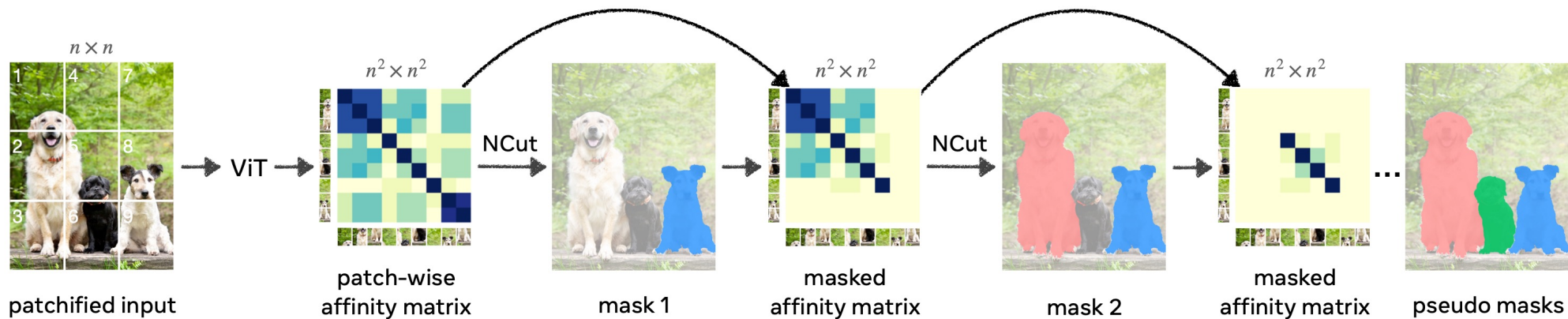
- Use a pretrained DINO ViT network.



MaskCut

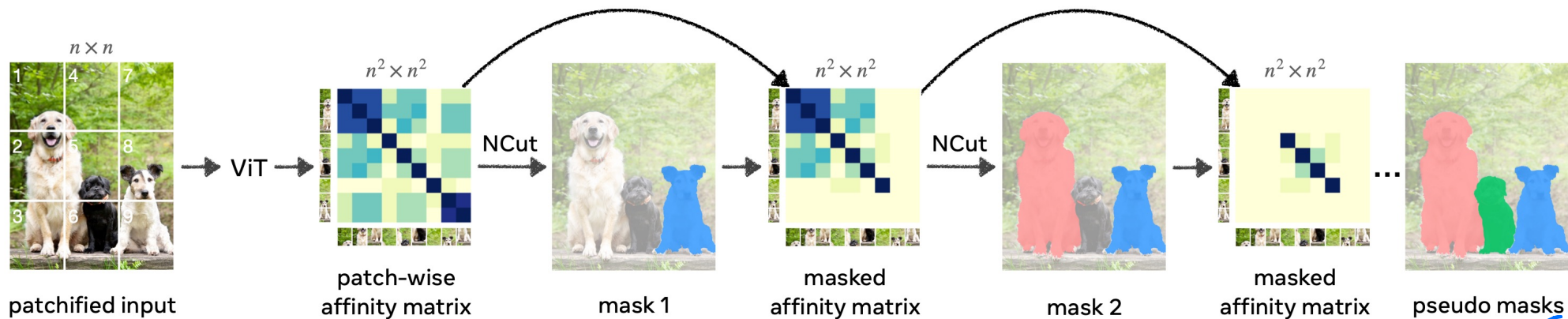
- Use a pretrained DINO ViT network.

- Use the “key” features from the last attention layer: $W_{ij} = \frac{K_i K_j}{\|K_i\|_2 \|K_j\|_2}$



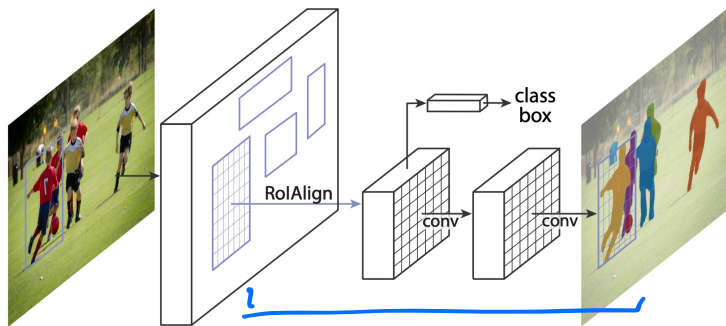
MaskCut

- Use a pretrained DINO ViT network.
- Use the “key” features from the last attention layer: $W_{ij} = \frac{K_i K_j}{\|K_i\|_2 \|K_j\|_2}$
- Iterative NCut on the pairwise matrix by masking out the regions from previous stages.



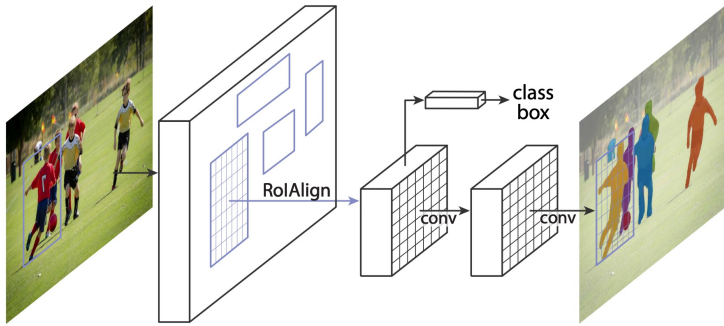
Iterative Self-Training

- Now add a MaskRCNN structure on top of the pretrained network.



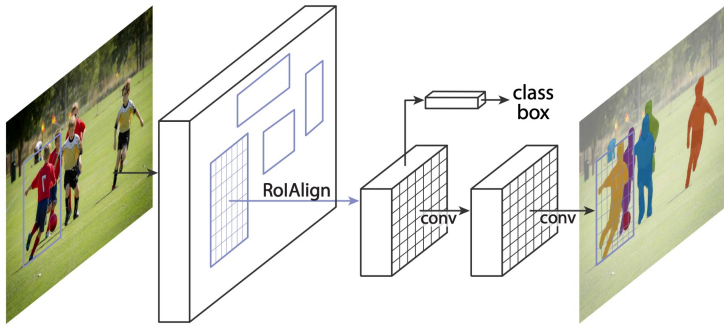
Iterative Self-Training

- Now add a MaskRCNN structure on top of the pretrained network.
- Select the predictions with the highest confidence score and use them as labels.

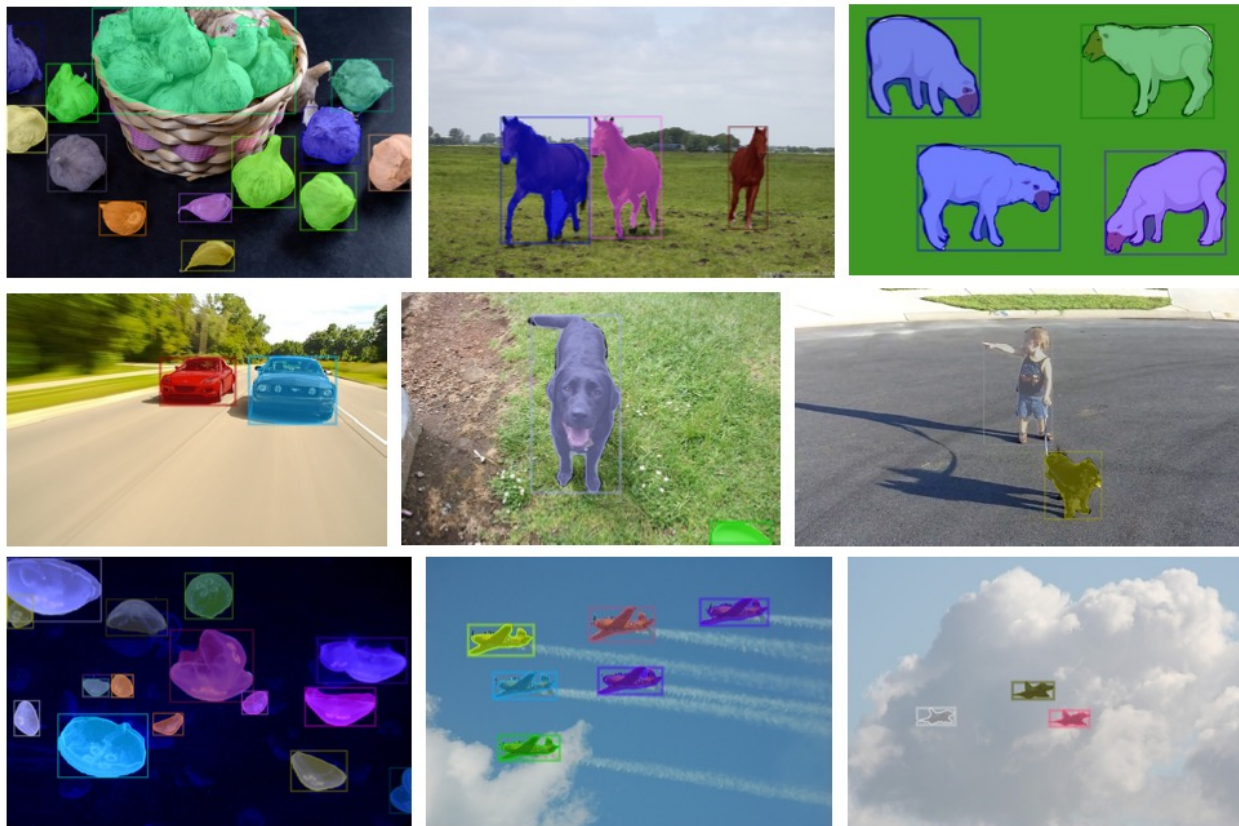


Iterative Self-Training

- Now add a MaskRCNN structure on top of the pretrained network.
- Select the predictions with the highest confidence score and use them as labels.
- Neural networks can learn from the noisy labels and output smoother predictions.

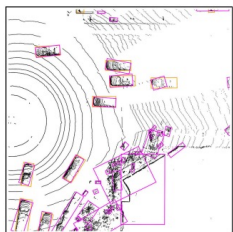


More Visualization

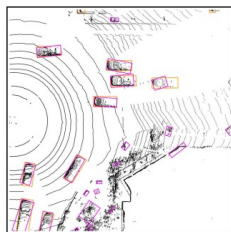


Pseudo Labels in 3D

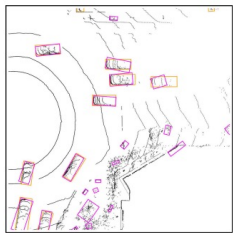
1 Point clustering
pseudo-labels



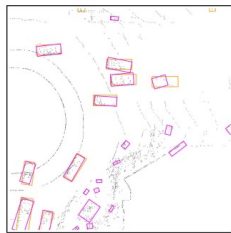
2 Filter out temporally inconsistent tracks



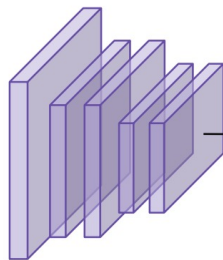
3 Randomly drop lidar beams



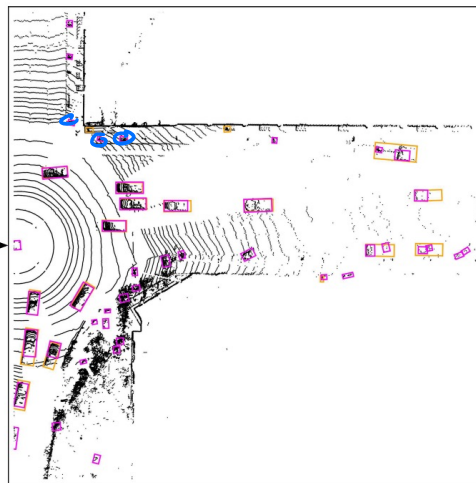
4 Randomly drop spherical rows/cols



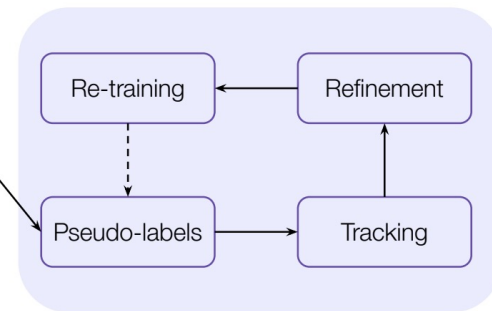
5 Train CNN in short range



6 Zero-shot generalization to long-range



7 Self-training in long-range

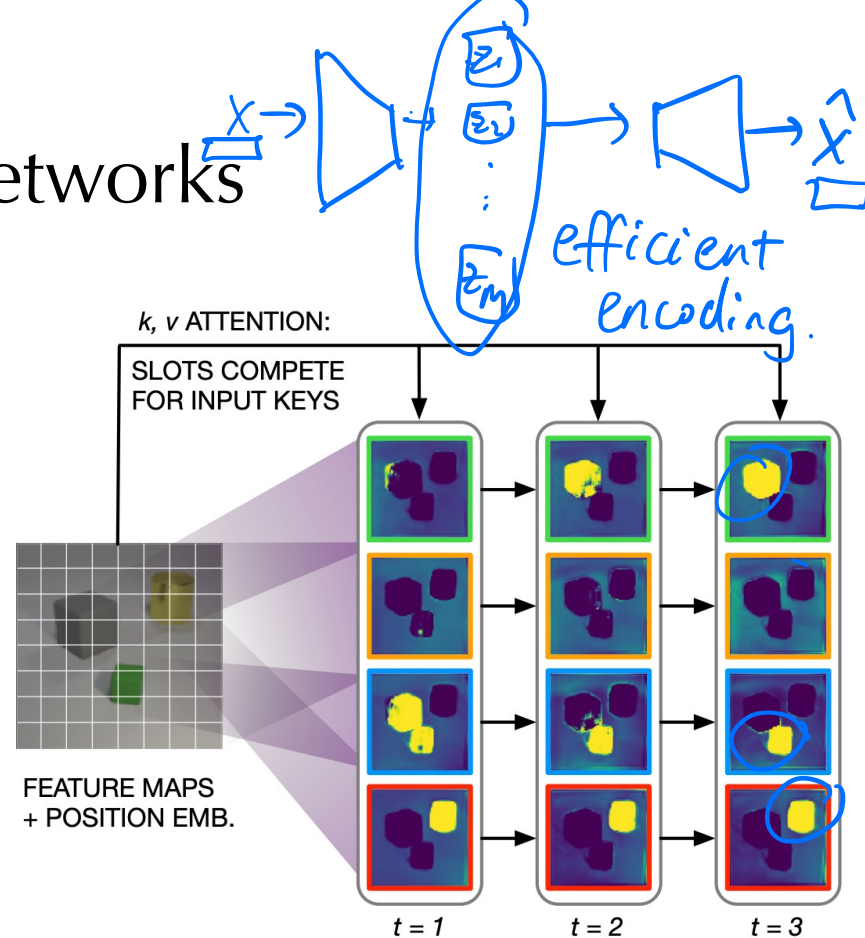


Iterative Refinement of Pseudo Labels



Slot Attention Networks

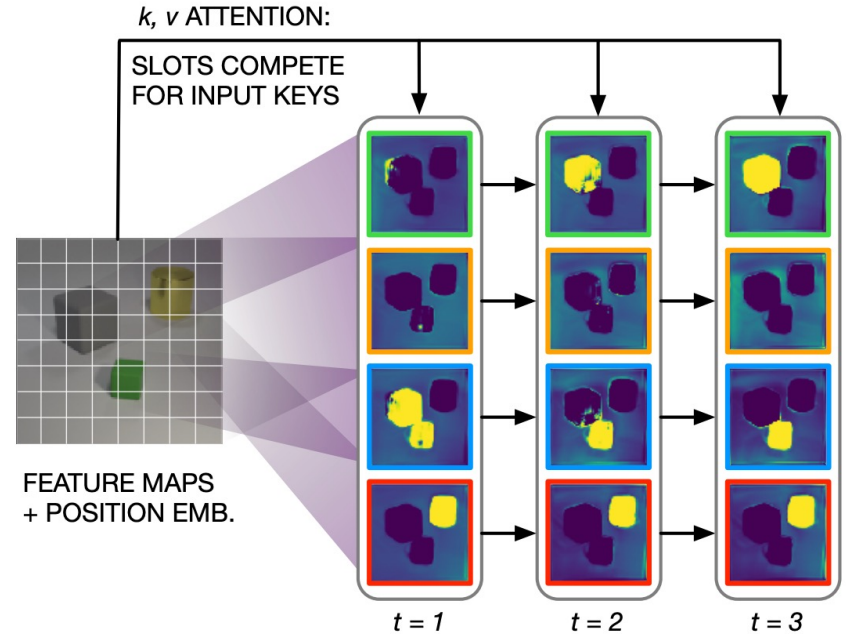
- Can we learn clustering as an end-to-end operation?



(a) Slot Attention module.

Slot Attention Networks

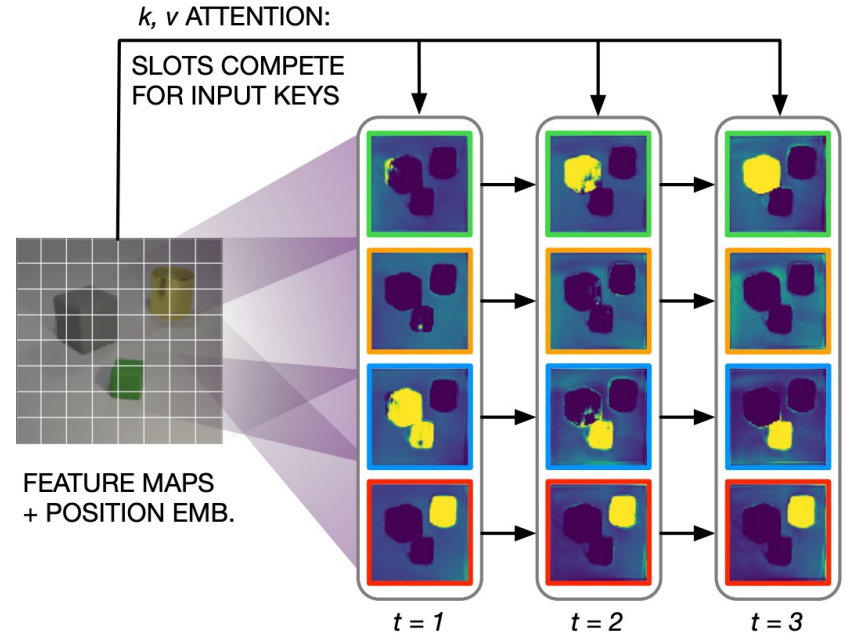
- Can we learn clustering as an end-to-end operation?
- Slot attention is inspired by the success of the attention mechanism.



(a) Slot Attention module.

Slot Attention Networks

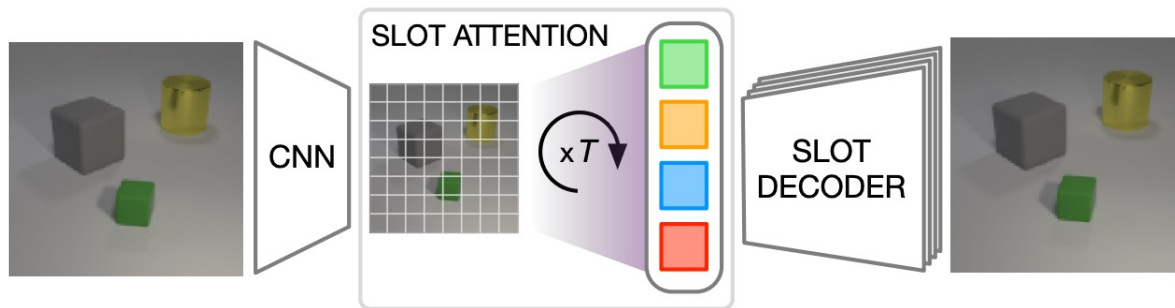
- Can we learn clustering as an end-to-end operation?
- Slot attention is inspired by the success of the attention mechanism.
- Each “slot” attends to a region of the image and stores an object centric representation.



(a) Slot Attention module.

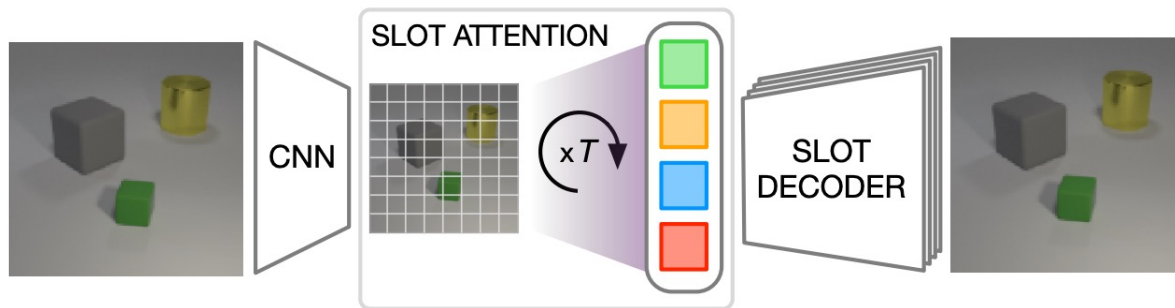
Slot Attention Networks

- Goal: Reconstruct the image with a concise slot-based representation.



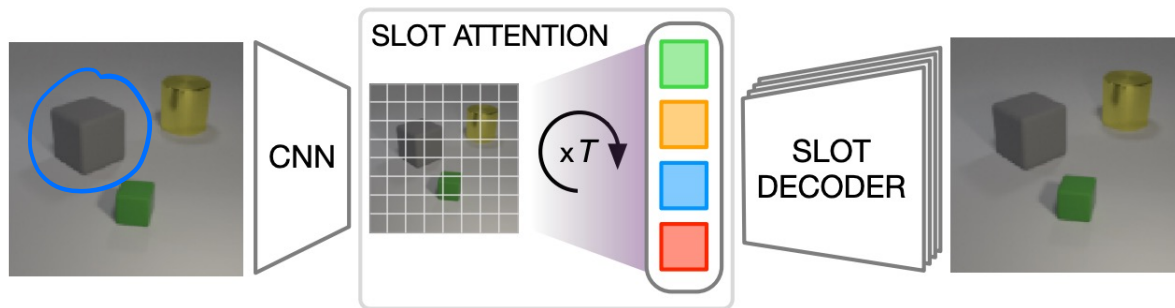
Slot Attention Networks

- Goal: Reconstruct the image with a concise slot-based representation.
- Input: $x \in \mathbb{R}^{N \times D}$ (after encoder), Slots: $m \in \mathbb{R}^{M \times D}$. Normalize:
 $\tilde{m}_{t-1} = \text{LN}(m_{t-1})$.



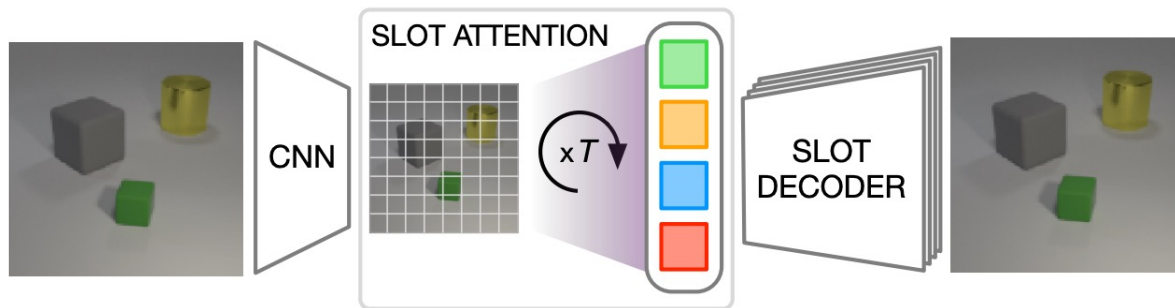
Slot Attention Networks

- Goal: Reconstruct the image with a concise slot-based representation.
- Input: $x \in \mathbb{R}^{N \times D}$ (after encoder), Slots: $m \in \mathbb{R}^{M \times D}$. Normalize:
 $\tilde{m}_{t-1} = LN(m_{t-1})$.
- Attention over slots: $a_{t,i,j} = \frac{\frac{1}{\sqrt{D}} k(x_i) \cdot q(\tilde{m}_j)^\top}{\sum_j \frac{1}{\sqrt{D}} k(x_i) \cdot q(\tilde{m}_j)^\top}$.



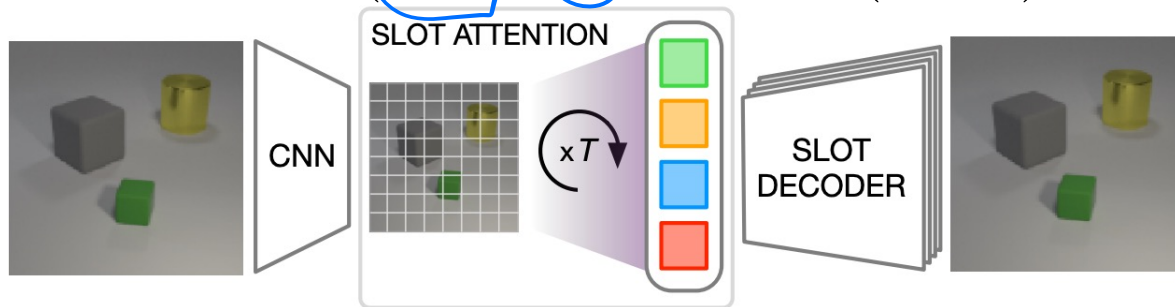
Slot Attention Networks

- Goal: Reconstruct the image with a concise slot-based representation.
- Input: $x \in \mathbb{R}^{N \times D}$ (after encoder), Slots: $m \in \mathbb{R}^{M \times D}$. Normalize:
 $\tilde{m}_{t-1} = LN(m_{t-1})$.
- Attention over slots: $a_{t,i,j} = \frac{\frac{1}{\sqrt{D}} k(x_i) \cdot q(\tilde{m}_j)^\top}{\sum_j \frac{1}{\sqrt{D}} k(x_i) \cdot q(\tilde{m}_j)^\top}$.
- Updates: $u_{tj} = \sum_i a_{tij} v(x_i)$.

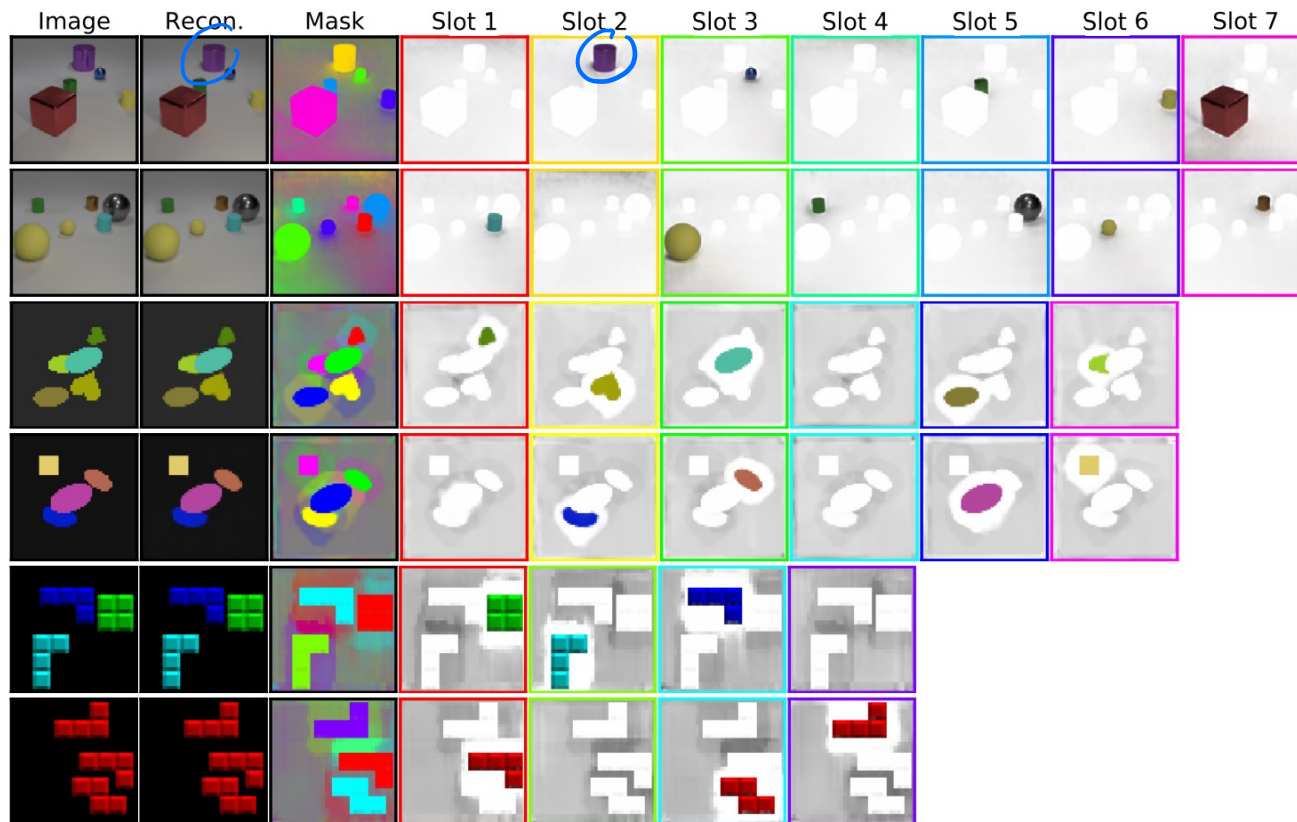


Slot Attention Networks

- Goal: Reconstruct the image with a concise slot-based representation.
- Input: $x \in \mathbb{R}^{N \times D}$ (after encoder), Slots: $m \in \mathbb{R}^{M \times D}$. Normalize: $\tilde{m}_{t-1} = LN(m_{t-1})$.
- Attention over slots: $a_{t,i,j} = \frac{\frac{1}{\sqrt{D}} k(x_i) \cdot q(\tilde{m}_j)^\top}{\sum_j \frac{1}{\sqrt{D}} k(x_i) \cdot q(\tilde{m}_j)^\top}$.
- Updates: $u_{tj} = \sum_i a_{tij} v(x_i)$.
- Write into slots: $m_t = GRU(\text{prev mem } m_{t-1}, \text{updates } u_t) + MLP(\tilde{m}_{t-1})$.

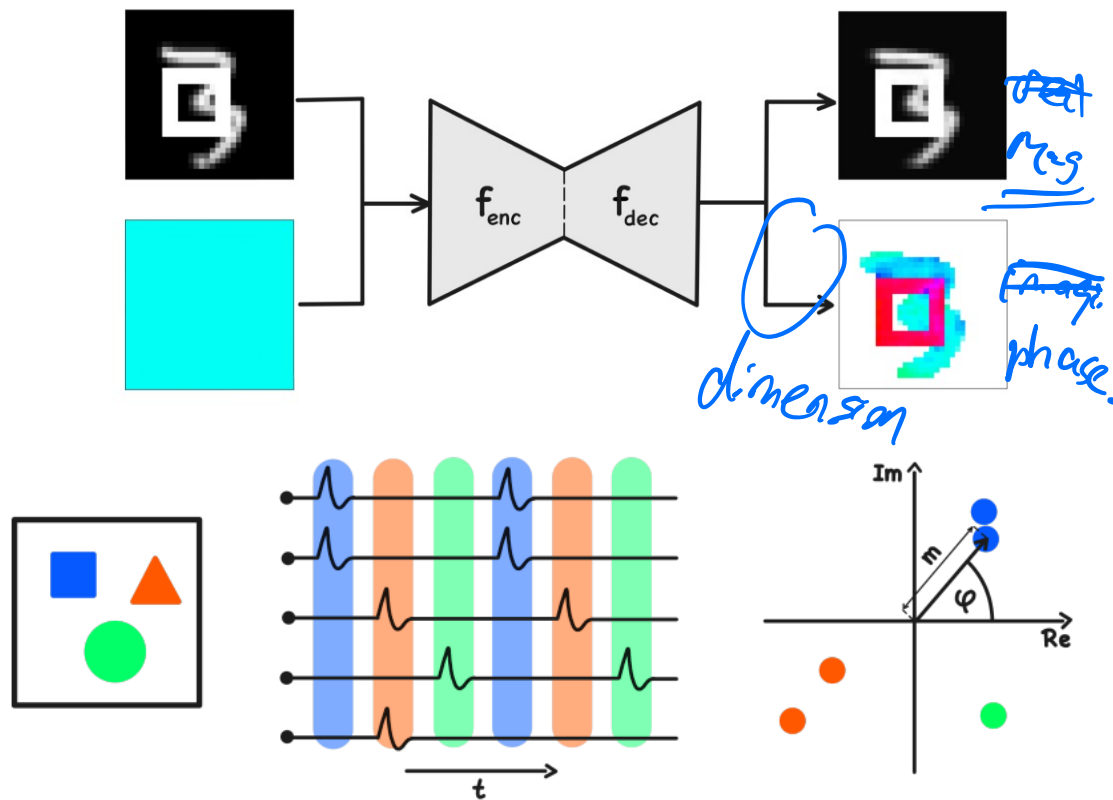


Slot Attention Networks



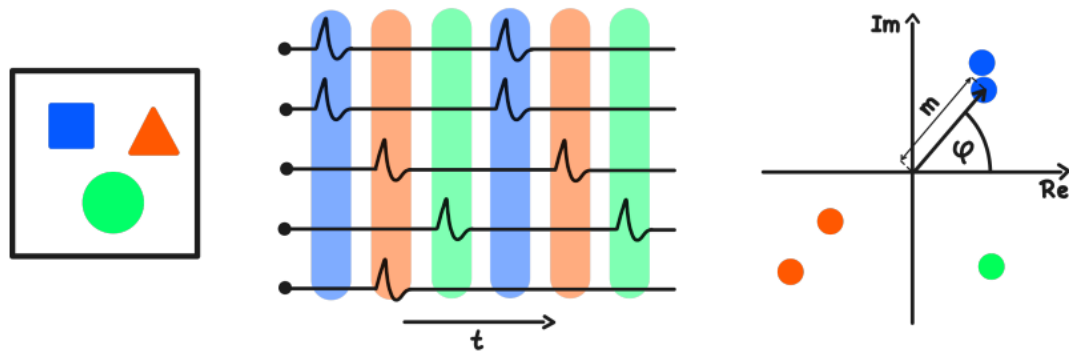
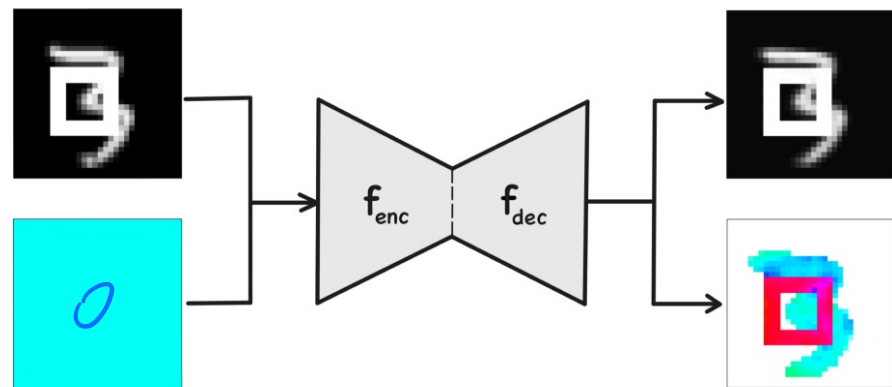
Complex-Valued Autoencoders (CAEs)

- The complex number can represent magnitude and phase: $z = m \cdot e^{i\varphi} \in \mathbb{C}$.



Complex-Valued Autoencoders (CAEs)

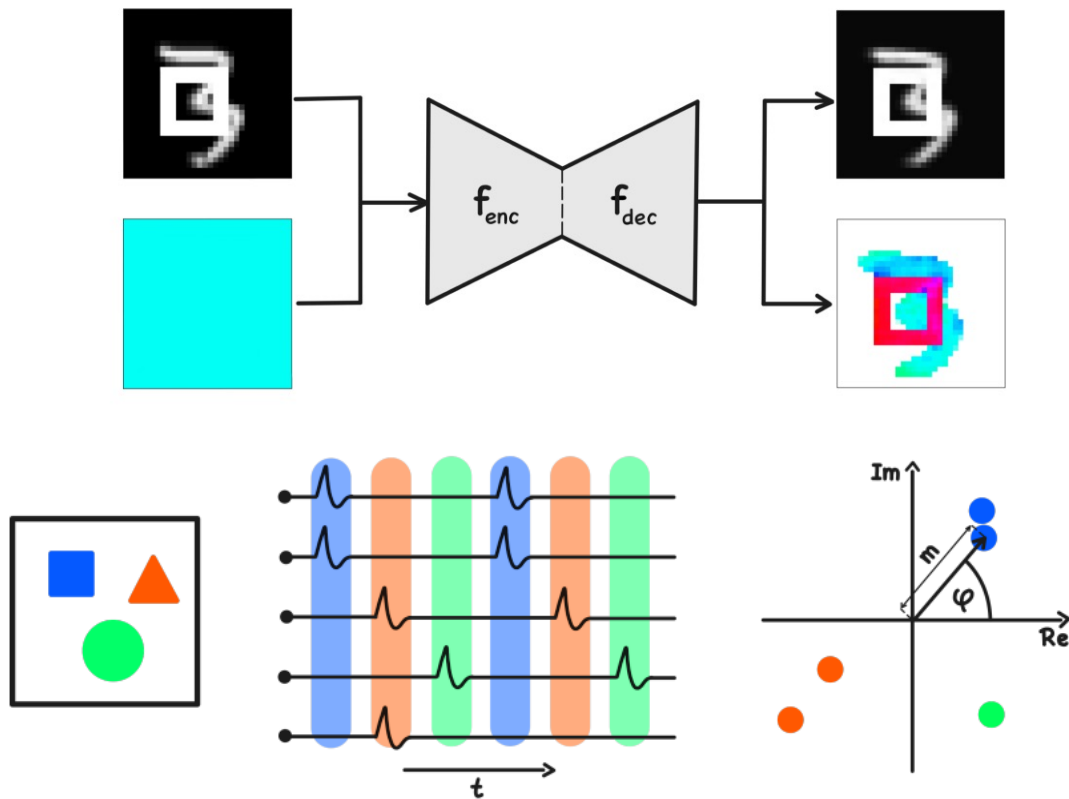
- The complex number can represent magnitude and phase: $z = m \cdot e^{i\varphi} \in \mathbb{C}$.
- Each pixel starts with an initial phase 0.



Complex-Valued Autoencoders (CAEs)

- The complex number can represent magnitude and phase: $z = m \cdot e^{i\varphi} \in \mathbb{C}$.
- Each pixel starts with an initial phase 0.

$$\hat{\mathbf{z}} = \underbrace{f_{\text{dec}}}_{\text{decoder}}(\underbrace{f_{\text{enc}}}_{\text{encoder}}(\underbrace{\mathbf{x}}_{\text{input}})) \in \mathbb{C}^{h \times w}.$$

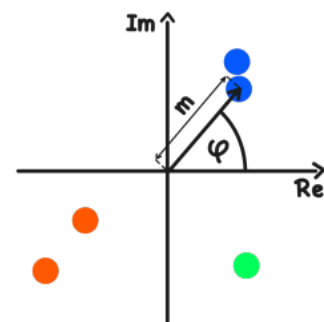
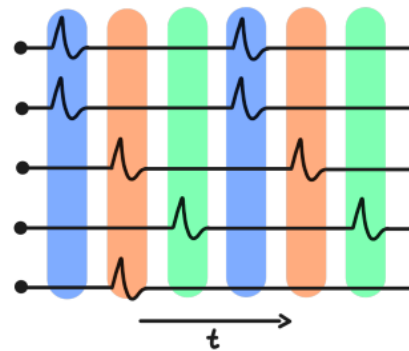
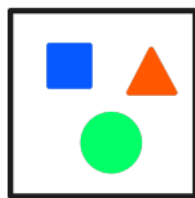
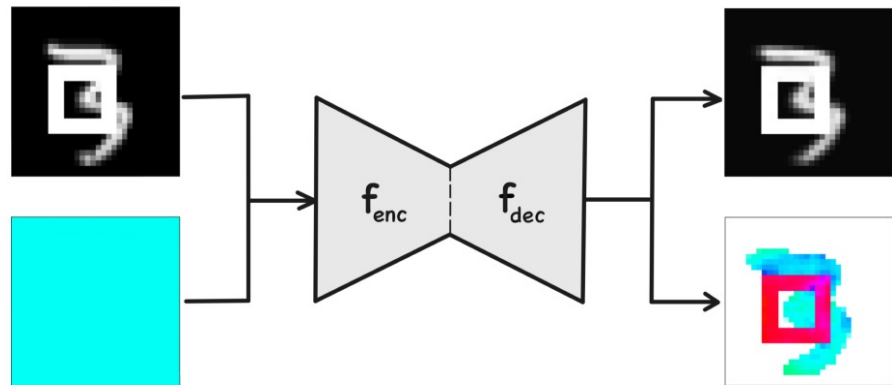


Complex-Valued Autoencoders (CAEs)

- The complex number can represent magnitude and phase: $z = m \cdot e^{i\varphi} \in \mathbb{C}$.
- Each pixel starts with an initial phase 0.

$$\hat{\mathbf{z}} = f_{\text{dec}}(f_{\text{enc}}(\mathbf{x})) \in \mathbb{C}^{h \times w}.$$

$$\hat{\mathbf{x}} = f_{\text{out}}(\hat{\mathbf{z}}) \in \mathbb{R}^{h \times w}.$$



CAE: More Details

- Apply weights separately to real and imaginary:

$$\psi = f_{\mathbf{w}}(\mathbf{z}) = \underbrace{f_{\mathbf{w}}(\text{Re}(\mathbf{z}))}_{\text{real part}} + \underbrace{f_{\mathbf{w}}(\text{Im}(\mathbf{z}))}_{\text{imaginary part}} \cdot i \in \mathbb{C}_{d_{\text{out}}}$$

CAE: More Details

- Apply weights separately to real and imaginary:

$$\psi = f_{\mathbf{w}}(\mathbf{z}) = f_{\mathbf{w}}(\text{Re}(\mathbf{z})) + f_{\mathbf{w}}(\text{Im}(\mathbf{z})) \cdot i \in \mathbb{C}_{d_{\text{out}}}$$

- Bias on magnitude and phase:

$$\mathbf{m}_{\psi} = |\psi| + \mathbf{b}_m \in \mathbb{R}^{d_{\text{out}}} \quad \varphi_{\psi} = \arg(\psi) + \mathbf{b}_{\varphi} \in \mathbb{R}^{d_{\text{out}}}$$

CAE: More Details

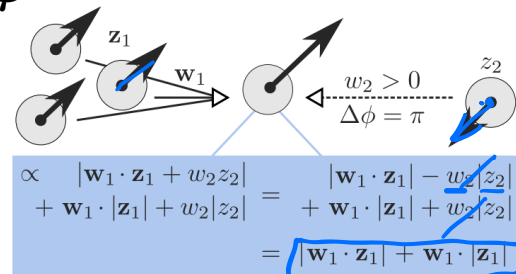
- Apply weights separately to real and imaginary:

$$\psi = f_{\mathbf{w}}(\mathbf{z}) = f_{\mathbf{w}}(\text{Re}(\mathbf{z})) + f_{\mathbf{w}}(\text{Im}(\mathbf{z})) \cdot i \in \mathbb{C}_{d_{\text{out}}}$$

- Bias on magnitude and phase:

$$\mathbf{m}_{\psi} = |\psi| + \mathbf{b}_m \in \mathbb{R}^{d_{\text{out}}} \quad \varphi_{\psi} = \arg(\psi) + \mathbf{b}_{\varphi} \in \mathbb{R}^{d_{\text{out}}}$$

- Gating: $\chi = \underline{f_{\mathbf{w}}(|\mathbf{z}|)} + \mathbf{b}_m \in \mathbb{R}^{d_{\text{out}}}$
 $\mathbf{m}_{\mathbf{z}} = 0.5 \cdot \mathbf{m}_{\psi} + 0.5 \cdot \chi \in \mathbb{R}^{d_{\text{out}}}$



CAE: More Details

- Apply weights separately to real and imaginary:

$$\psi = f_{\mathbf{w}}(\mathbf{z}) = f_{\mathbf{w}}(\text{Re}(\mathbf{z})) + f_{\mathbf{w}}(\text{Im}(\mathbf{z})) \cdot i \in \mathbb{C}_{d_{\text{out}}}$$

- Bias on magnitude and phase:

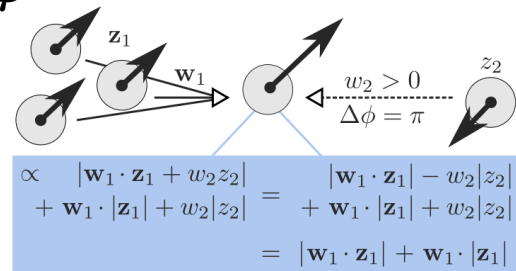
$$\mathbf{m}_{\psi} = |\psi| + \mathbf{b}_m \in \mathbb{R}^{d_{\text{out}}} \quad \varphi_{\psi} = \arg(\psi) + \mathbf{b}_{\varphi} \in \mathbb{R}^{d_{\text{out}}}$$

- Gating: $\chi = f_{\mathbf{w}}(|\mathbf{z}|) + \mathbf{b}_m \in \mathbb{R}^{d_{\text{out}}}$

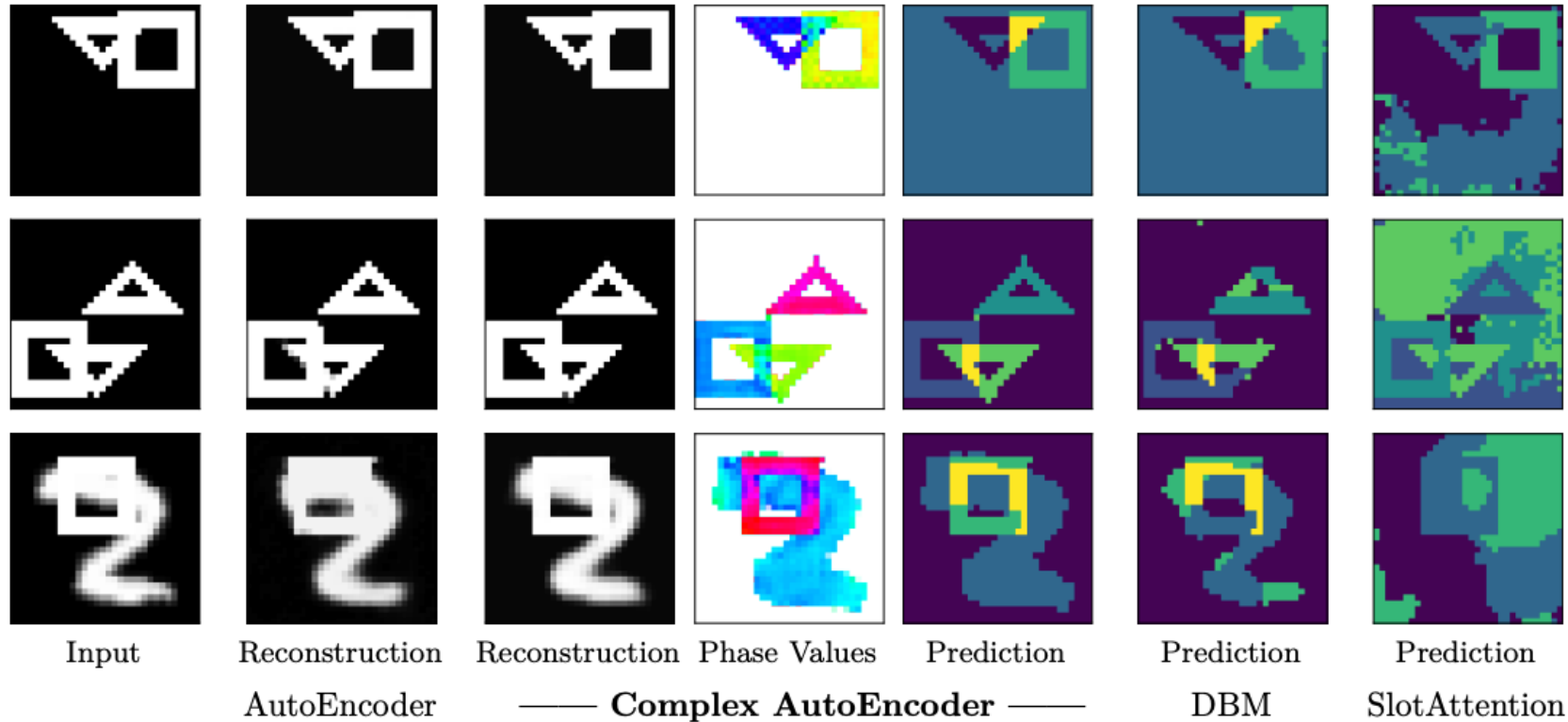
$$\mathbf{m}_{\mathbf{z}} = 0.5 \cdot \mathbf{m}_{\psi} + 0.5 \cdot \chi \in \mathbb{R}^{d_{\text{out}}}$$

- Activation: $\mathbf{z}' = \text{ReLU}(\text{BatchNorm}(\mathbf{m}_{\mathbf{z}})) \cdot e^{i\varphi_{\psi}} \in \mathbb{C}^{d_{\text{out}}}$

mag *phase*



Complex-Valued Autoencoders



Summary: Object Discovery

- Combine deep features with clustering algorithms.

Summary: Object Discovery

- Combine deep features with clustering algorithms.
- Pseudo-labels to train detector networks.

Summary: Object Discovery

- Combine deep features with clustering algorithms.
- Pseudo-labels to train detector networks.
- Creative end-to-end learning-based solutions exist, but there are still plenty room for improvement.
 - Possible to train from scratch!

Summary: Object Discovery

- Combine deep features with clustering algorithms.
- Pseudo-labels to train detector networks.
- Creative end-to-end learning-based solutions exist, but there are still plenty room for improvement.
 - Possible to train from scratch!
- What do we make use of the discovered objects? Is it better to keep the awareness in the latent space?

Module 4:

World Models and End-to-End Planning

World Models: Predicting Future

World Models: Predicting Future

- There is a debate whether predicting future is necessary.

World Models: Predicting Future

- There is a debate whether predicting future is necessary.
- Data efficiency
 - Predicting future can “simulate” roll out without querying for outcomes.
 - Leveraging the massive amount of data in past experiences (not just the final reward).

World Models: Predicting Future

- There is a debate whether predicting future is necessary.
- Data efficiency
 - Predicting future can “simulate” roll out without querying for outcomes.
 - Leveraging the massive amount of data in past experiences (not just the final reward).
- Long-horizon planning: Predicting high-level future steps is needed.

World Models: Predicting Future

- There is a debate whether predicting future is necessary.
- Data efficiency
 - Predicting future can “simulate” roll out without querying for outcomes.
 - Leveraging the massive amount of data in past experiences (not just the final reward).
- Long-horizon planning: Predicting high-level future steps is needed.
- Can representations learned from SSL help us build better prediction?

Model-Predictive Control (MPC)

- A classic example of world models.

Model-Predictive Control (MPC)

- A classic example of world models.
- Analytical forms, complete knowledge of the dynamical system.

Model-Predictive Control (MPC)

- A classic example of world models.
- Analytical forms, complete knowledge of the dynamical system.

General Form:

$$\begin{aligned}\frac{d\mathbf{x}}{dt} &= f(\mathbf{x}, \mathbf{u}, t) \\ \mathbf{y} &= h(\mathbf{x}, \mathbf{u}, t) \\ \mathbf{x}(t_0) &= \mathbf{x}_0\end{aligned}$$

Handwritten blue annotations:
- A blue arrow points from the word "state" to the $\mathbf{x}$ in the first equation.
- A blue arrow points from the word "control" to the $\mathbf{u}$ in the first equation.
- The $\mathbf{u}$ in the first equation is circled in blue.
- The $\mathbf{x}_0$ in the third equation is circled in blue.
- The entire set of equations is underlined in blue.

Model-Predictive Control (MPC)

- A classic example of world models.
- Analytical forms, complete knowledge of the dynamical system.

General Form:

$$\frac{d\mathbf{x}}{dt} = f(\mathbf{x}, \mathbf{u}, t)$$

$$\mathbf{y} = h(\mathbf{x}, \mathbf{u}, t)$$

$$\mathbf{x}(t_0) = \mathbf{x}_0$$

Linear Form:

$$\boxed{\frac{dx}{dt}} = \underline{A\mathbf{x} + B\mathbf{u}}$$

$$\mathbf{y} = C\mathbf{x} + D\mathbf{u}$$

$$\mathbf{x}(t_0) = \mathbf{x}_0$$

Model-Predictive Control (MPC)

- A classic example of world models.
- Analytical forms, complete knowledge of the dynamical system.

General Form:

$$\frac{d\mathbf{x}}{dt} = f(\mathbf{x}, \mathbf{u}, t)$$

$$\mathbf{y} = h(\mathbf{x}, \mathbf{u}, t)$$

$$\mathbf{x}(t_0) = \mathbf{x}_0$$

Linear Form:

$$\frac{dx}{dt} = A\mathbf{x} + B\mathbf{u}$$

$$\mathbf{y} = C\mathbf{x} + D\mathbf{u}$$

$$\mathbf{x}(t_0) = \mathbf{x}_0$$

Optimize Value/Cost Function:

$$\min_{\mathbf{u}} J(\mathbf{x}(0), \mathbf{u})$$

Model-Predictive Control (MPC)

- A classic example of world models.
- Analytical forms, complete knowledge of the dynamical system.

General Form:

$$\frac{d\mathbf{x}}{dt} = f(\mathbf{x}, \mathbf{u}, t)$$

$$\mathbf{y} = h(\mathbf{x}, \mathbf{u}, t)$$

$$\mathbf{x}(t_0) = \mathbf{x}_0$$

Linear Form:

$$\frac{dx}{dt} = A\mathbf{x} + B\mathbf{u}$$

$$\mathbf{y} = C\mathbf{x} + D\mathbf{u}$$

$$\mathbf{x}(t_0) = \mathbf{x}_0$$

Optimize Value/Cost Function:

$$\min_{\mathbf{u}} J(\mathbf{x}(0), \mathbf{u})$$

Quadratic Form (LQR):

$$J = \int_0^\infty \underbrace{\mathbf{x}^\top Q \mathbf{x}}_{\text{reach?}} + \underbrace{\mathbf{u}^\top R \mathbf{u}}_{\text{smooth?}} dt.$$

Model-Predictive Control (MPC)

- A classic example of world models.
- Analytical forms, complete knowledge of the dynamical system.

General Form:

$$\frac{d\mathbf{x}}{dt} = f(\mathbf{x}, \mathbf{u}, t)$$

$$\mathbf{y} = h(\mathbf{x}, \mathbf{u}, t)$$

$$\mathbf{x}(t_0) = \mathbf{x}_0$$

Linear Form:

$$\frac{dx}{dt} = A\mathbf{x} + B\mathbf{u}$$

$$\mathbf{y} = C\mathbf{x} + D\mathbf{u}$$

$$\mathbf{x}(t_0) = \mathbf{x}_0$$

Optimize Value/Cost Function:

$$\min_{\mathbf{u}} J(\mathbf{x}(0), \mathbf{u})$$

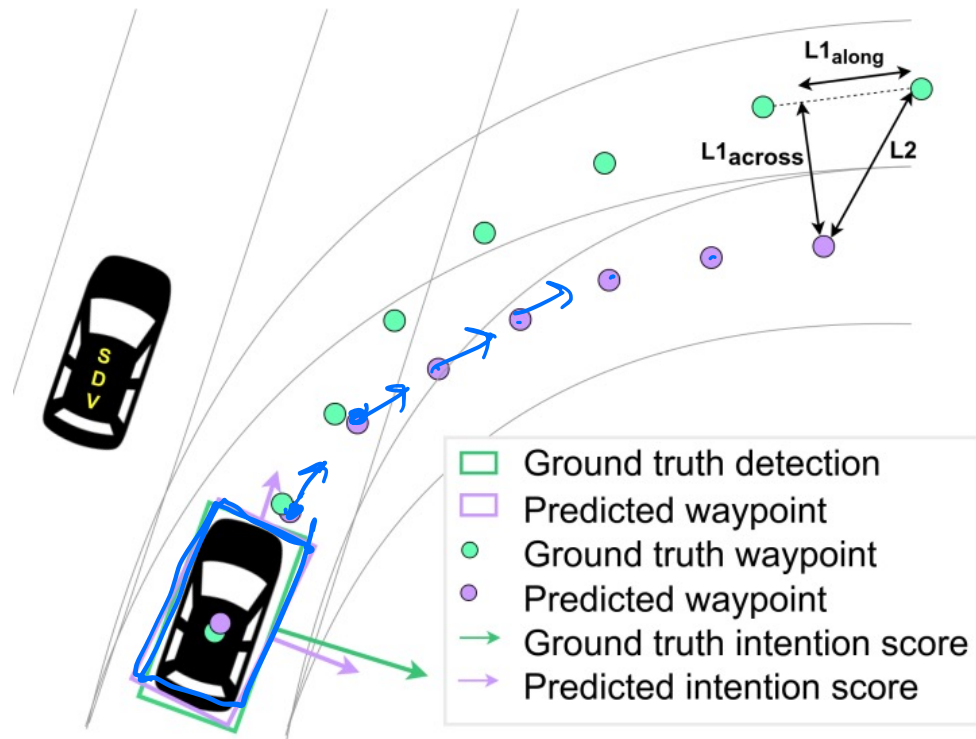
Quadratic Form (LQR):

$$J = \int_0^\infty \mathbf{x}^\top Q \mathbf{x} + \mathbf{u}^\top R \mathbf{u} dt.$$

- Example: Cars: x: position velocity angle angular velocity; u: jerk and angular accel.

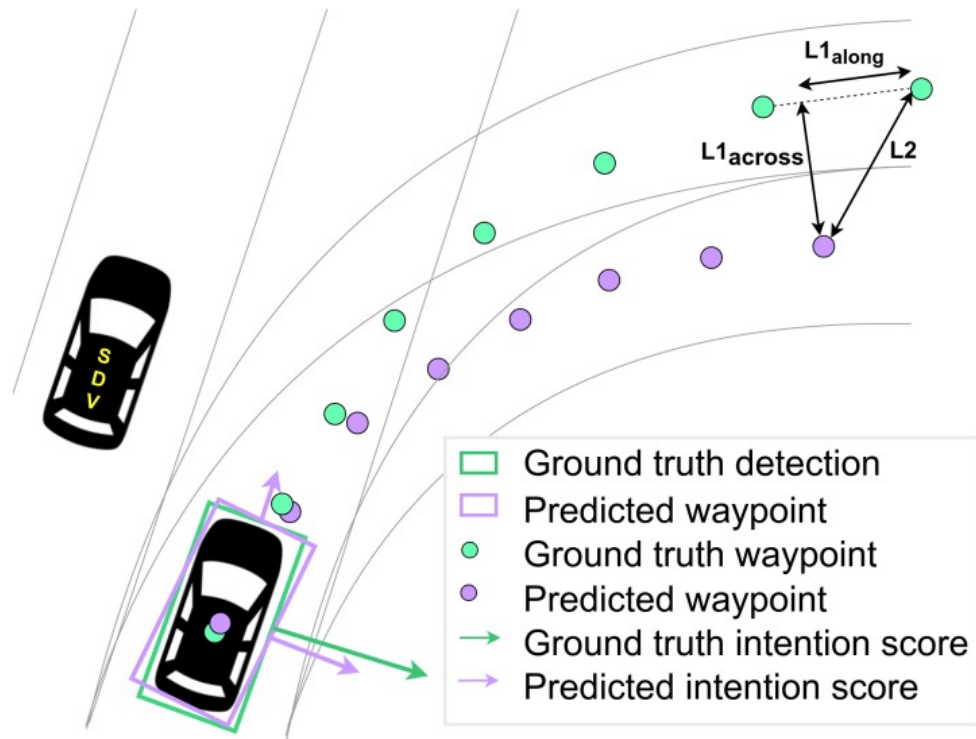
Trajectory Prediction as Object Detection

- We also need to predict external dynamic objects.



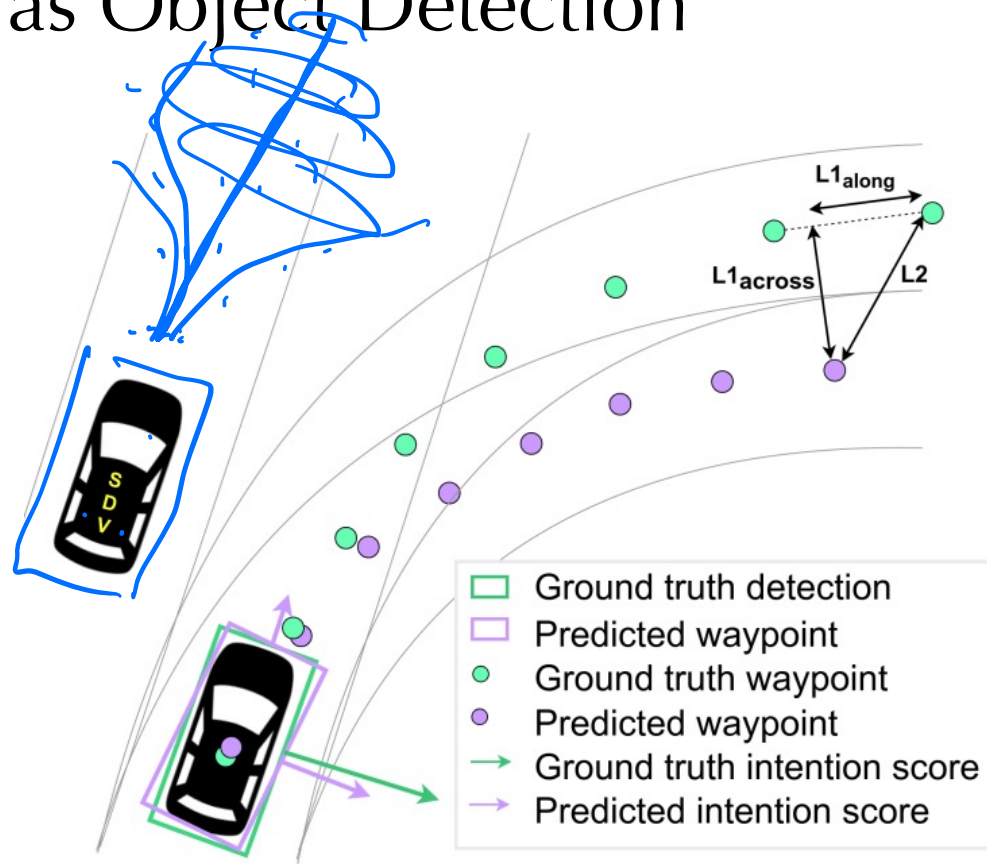
Trajectory Prediction as Object Detection

- We also need to predict external dynamic objects.
- How to capture multiple modes?



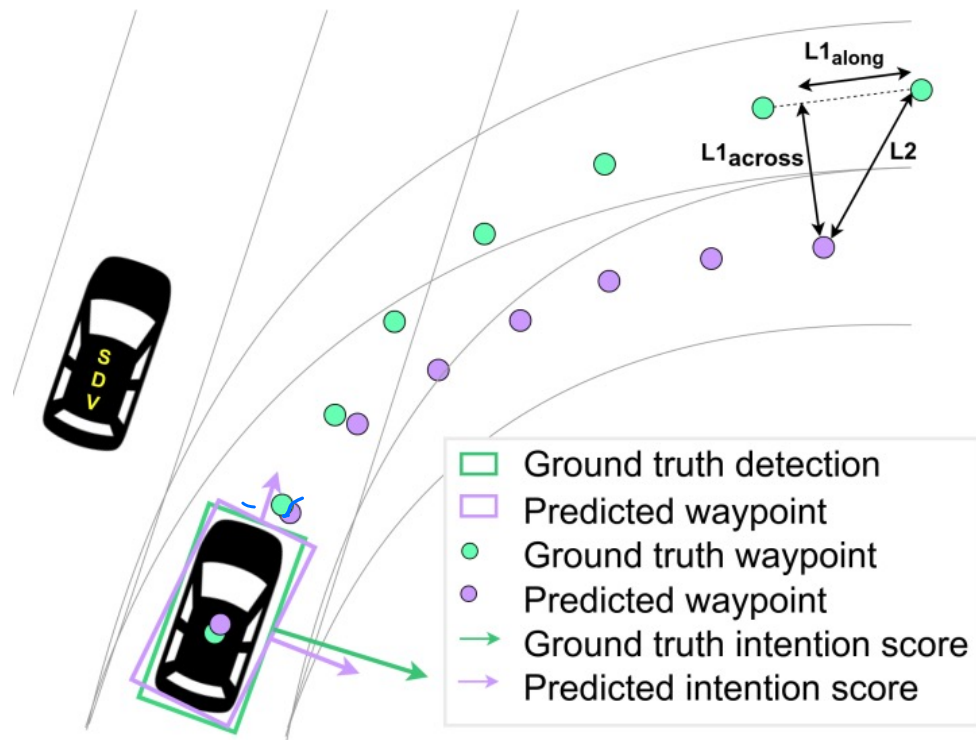
Trajectory Prediction as Object Detection

- We also need to predict external dynamic objects.
- How to capture multiple modes?
- Discrete intention prediction problem:
 - keep lane, turn left/right, left/right change lane, stopping, parked, etc.



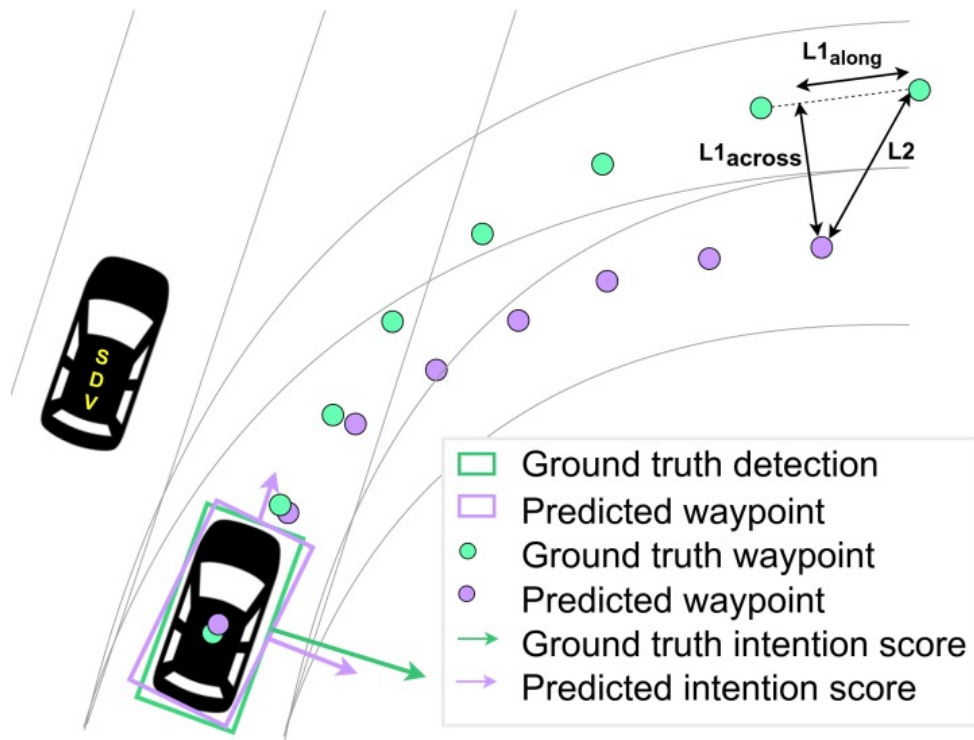
Trajectory Prediction as Object Detection

- We also need to predict external dynamic objects.
- How to capture multiple modes?
- Discrete intention prediction problem:
 - keep lane, turn left/right, left/right change lane, stopping, parked, etc.
- Output multiple trajectories



Trajectory Prediction as Object Detection

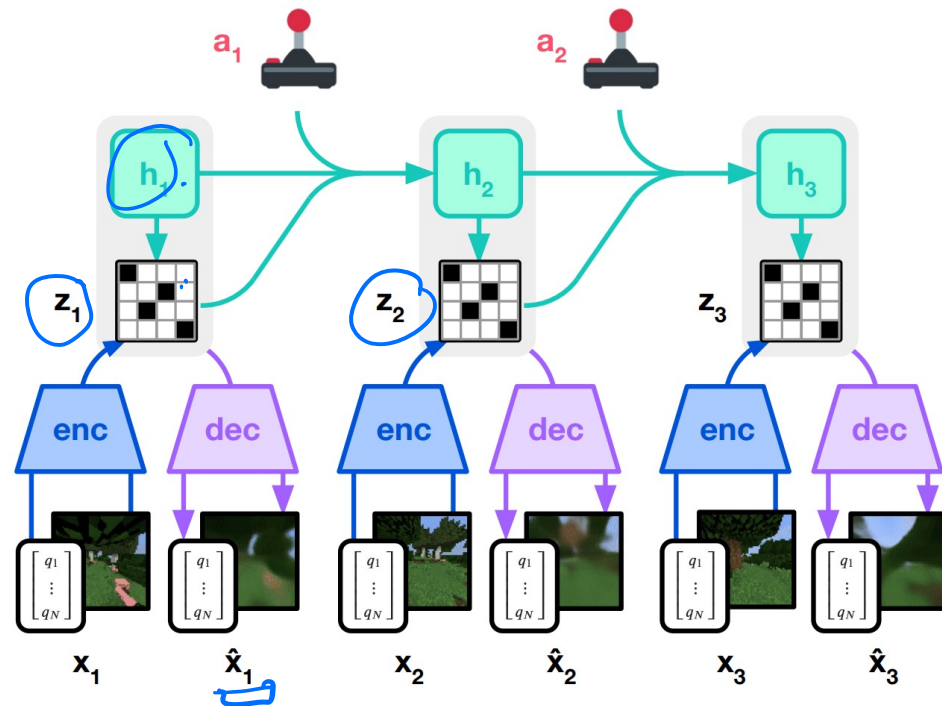
- We also need to predict external dynamic objects.
- How to capture multiple modes?
- Discrete intention prediction problem:
 - keep lane, turn left/right, left/right change lane, stopping, parked, etc.
- Output multiple trajectories
- Requires high-level action labels.



Latent Sequence World Model for RL

- Autoencoder to ensure the latent representations are meaningful.

$$z_t \sim q_\phi(z_t | h_t, x_t)$$
$$\hat{x}_t \sim p_\phi(\hat{x}_t | h_t, z_t)$$



Latent Sequence World Model for RL

- Autoencoder to ensure the latent representations are meaningful.

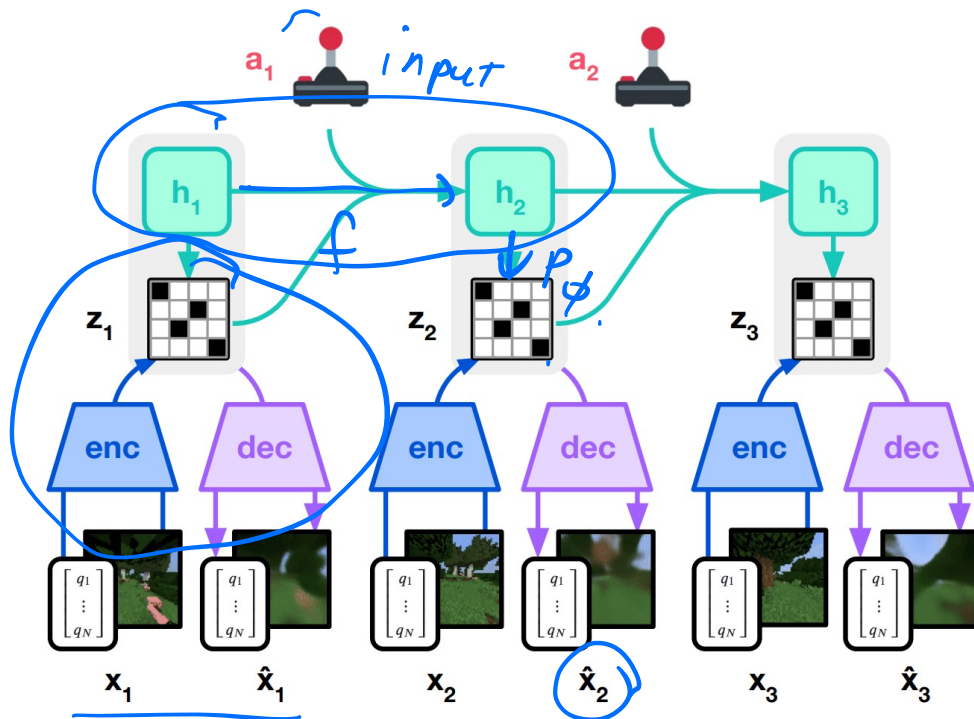
$$z_t \sim q_\phi(z_t \mid h_t, x_t)$$

$$\hat{x}_t \sim p_\phi(\hat{x}_t \mid h_t, z_t)$$

- Learn a sequence model to predict the latent conditioned on previous action.

$$h_t = f_\phi(h_{t-1}, z_{t-1}, a_{t-1})$$

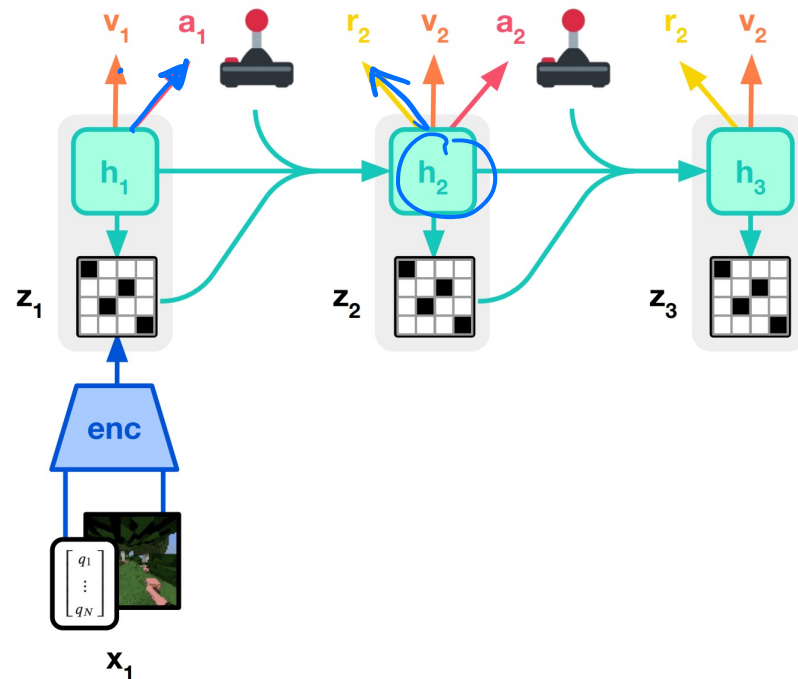
$$\hat{z}_t \sim p_\phi(\hat{z}_t \mid h_t)$$



Incorporating Rewards

- Predicting reward

$$\hat{r}_t \sim p_\phi(\hat{r}_t \mid h_t, z_t)$$



Incorporating Rewards

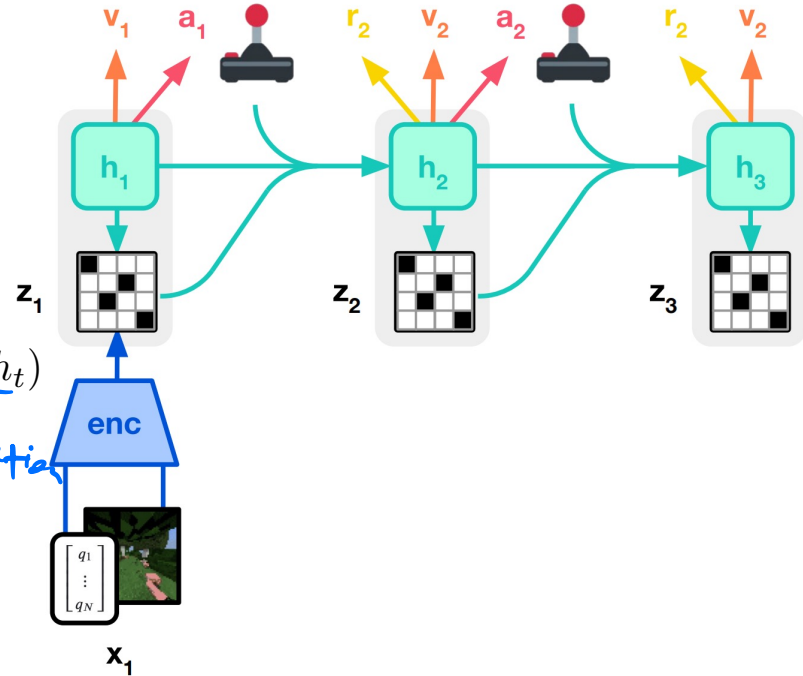
- Predicting reward

$$\hat{r}_t \sim p_\phi(\hat{r}_t \mid h_t, z_t)$$

- Reconstruction, reward, continue

$$\mathcal{L}_{\text{pred}}(\phi) = -\log p_\phi(\underbrace{x_t \mid z_t, h_t}_{\text{reConstruction}}) - \log p_\phi(\underbrace{r_t \mid z_t, h_t}_{\text{reward}}) - \log p_\phi(\underbrace{c_t \mid z_t, h_t}_{\text{Continuation}})$$

reConstruction . reward Continuation



Incorporating Rewards

- Predicting reward

$$\hat{r}_t \sim p_\phi(\hat{r}_t \mid h_t, z_t)$$

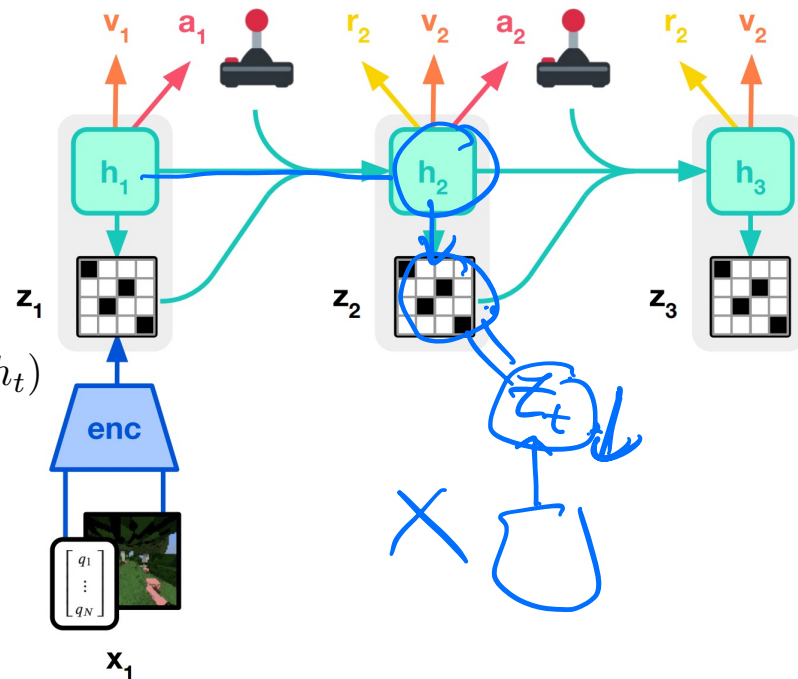
- Reconstruction, reward, continue

$$\mathcal{L}_{\text{pred}}(\phi) = -\log p_\phi(\underline{x}_t \mid z_t, h_t) - \log p_\phi(r_t \mid z_t, h_t) - \log p_\phi(c_t \mid z_t, h_t)$$

- Dynamics: Predicting future z

$$\mathcal{L}_{\text{dyn}}(\phi) = \max(1, \text{KL}[\text{sg}] q_\phi(z_t \mid h_t, x_t) \parallel p_\phi(z_t \mid h_t))$$

Stop-gradient



Incorporating Rewards

- Predicting reward

$$\hat{r}_t \sim p_\phi(\hat{r}_t \mid h_t, z_t)$$

- Reconstruction, reward, continue

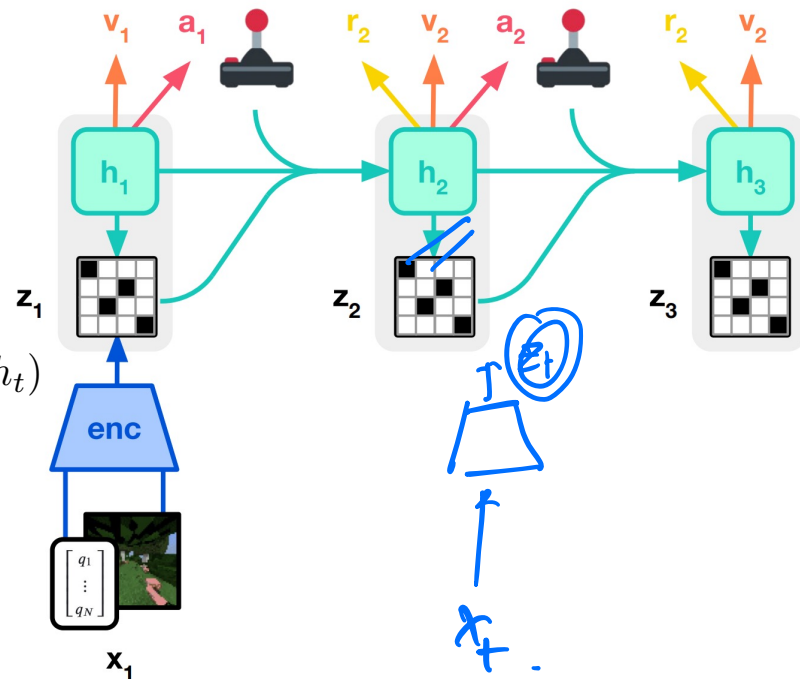
$$\mathcal{L}_{\text{pred}}(\phi) = -\log p_\phi(x_t \mid z_t, h_t) - \log p_\phi(r_t \mid z_t, h_t) - \log p_\phi(c_t \mid z_t, h_t)$$

- Dynamics: Predicting future z

$$\mathcal{L}_{\text{dyn}}(\phi) = \max(1, \text{KL}[\text{sg}(q_\phi(z_t \mid h_t, x_t)) \parallel p_\phi(z_t \mid h_t)])$$

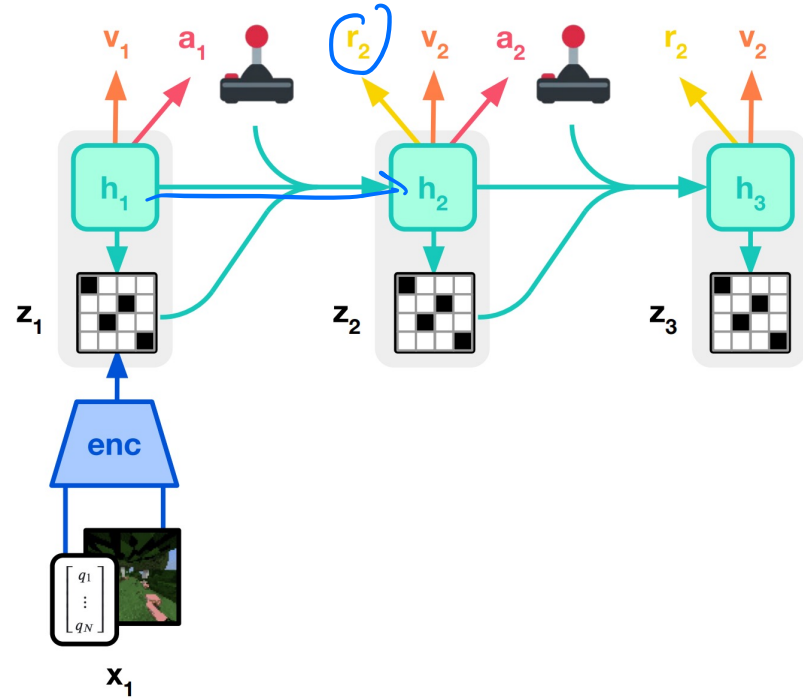
- Align representation

$$\mathcal{L}_{\text{rep}} = \max(1, \text{KL}[q_\phi(z_t \mid h_t, x_t) \parallel \text{sg}(p_\phi(z_t \mid h_t))])$$



Incorporating Rewards

- How to use WM in planning?
Predicting value by simulate a batch of trajectories.

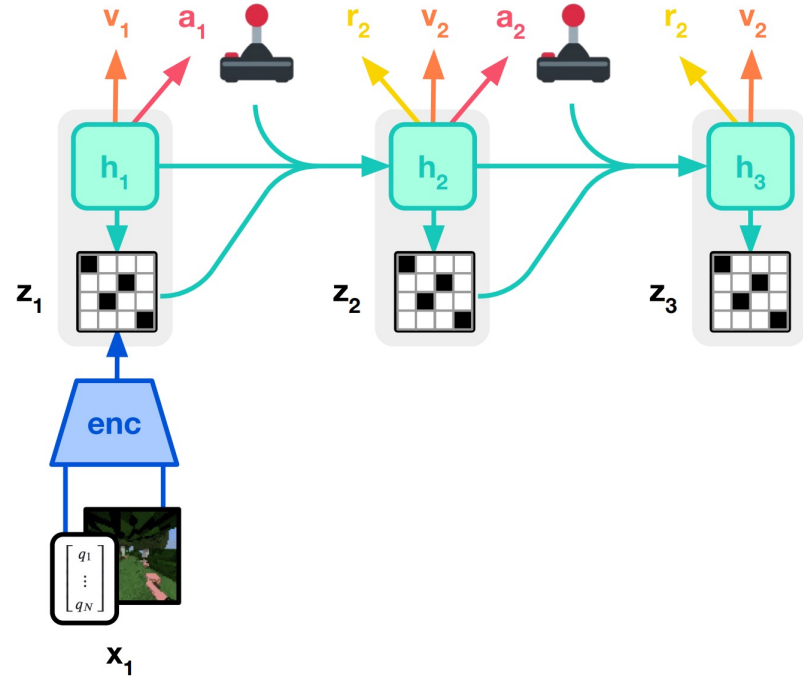


Incorporating Rewards

- How to use WM in planning?
Predicting value by simulate a batch of trajectories.

- Actor-Critic RL:

$$\underline{a_t} \sim \pi_{\theta}(a_t | s_t) \quad \overline{v_{\psi}(R_t | s_t)}$$

Incorporating Rewards

- How to use WM in planning?
Predicting value by simulate a batch of trajectories.

- Actor-Critic RL:

$$a_t \sim \pi_{\theta}(a_t | s_t) \quad v_{\psi}(R_t | s_t)$$

- Learning a critic:

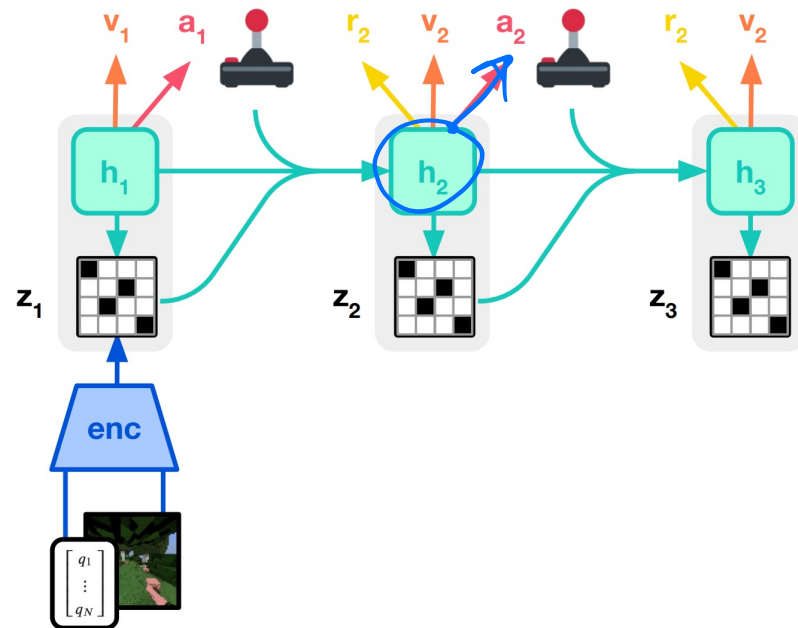
Categorical distribution

$$\mathcal{L}_{\phi} = - \sum_{t=1}^T \log p_{\phi}(R_t^{\lambda} | s_t) \quad s_t = \{h_t, z_t\}$$

Sum of discounted future rewards

↑ supervise critic.

$$R_t^{\lambda} = r_t + \gamma c_t [(1 - \lambda)v_t + \lambda R_{t+1}^{\lambda}] \quad R_T^{\lambda} = v_T$$



Actor Learning

- Learn a policy network $a_t \sim \pi_{\theta}(a_t \mid s_t)$

Actor Learning

- Learn a policy network $a_t \sim \pi_\theta(a_t | s_t)$
- REINFORCE algorithm [Williams 1992]

$$\mathcal{L}(\theta) = - \sum_{t=1}^T \text{sg}(\underbrace{(R_t^\lambda - v_\phi(s_t))}_{\text{Critic.}}) \underbrace{\frac{1}{\max(1, S)}}_{\text{Normalization}} \underbrace{\log \pi_\theta(a_t | s_t)}_{\text{Entropy Regularizer}} + \eta H[\pi_\theta(a_t | s_t)]$$

Actor Learning

- Learn a policy network $a_t \sim \pi_\theta(a_t \mid s_t)$
- REINFORCE algorithm [Williams 1992]

$$\mathcal{L}(\theta) = - \sum_{t=1}^T \overset{\text{Normalization}}{\text{sg}((R_t^\lambda - v_\phi(s_t)) / \max(1, S))} \log \pi_\theta(a_t \mid s_t) + \overset{\text{Entropy Regularizer}}{\eta H[\pi_\theta(a_t \mid s_t)]}$$

- Notes on policy gradient:

Actor Learning

- Learn a policy network $a_t \sim \pi_\theta(a_t \mid s_t)$
- REINFORCE algorithm [Williams 1992]

$$\mathcal{L}(\theta) = - \sum_{t=1}^T \text{sg}((R_t^\lambda - v_\phi(s_t)) / \overset{\text{Normalization}}{\max(1, S)}) \log \pi_\theta(a_t \mid s_t) + \overset{\text{Entropy Regularizer}}{\eta H[\pi_\theta(a_t \mid s_t)]}$$

- Notes on policy gradient:

$$\pi_\theta(\tau) \nabla_\theta \log \pi_\theta(\tau) = \pi_\theta(\tau) \frac{\nabla_\theta \pi_\theta(\tau)}{\pi_\theta(\tau)} = \nabla_\theta \pi_\theta(\tau).$$

Actor Learning

- Learn a policy network $a_t \sim \pi_\theta(a_t \mid s_t)$
- REINFORCE algorithm [Williams 1992]

$$\mathcal{L}(\theta) = - \sum_{t=1}^T \text{sg}((R_t^\lambda - v_\phi(s_t)) / \overset{\text{Normalization}}{\max(1, S)}) \log \pi_\theta(a_t \mid s_t) + \overset{\text{Entropy Regularizer}}{\eta H[\pi_\theta(a_t \mid s_t)]}$$

- Notes on policy gradient:

$$\pi_\theta(\tau) \nabla_\theta \log \pi_\theta(\tau) = \pi_\theta(\tau) \frac{\nabla_\theta \pi_\theta(\tau)}{\pi_\theta(\tau)} = \nabla_\theta \pi_\theta(\tau).$$

$$\pi_\theta(\tau) = \pi_\theta(s_1, a_1, \dots, s_T, a_T) = p(s_1) \prod_{t=1}^T \pi_\theta(a_t \mid s_t) p(s_{t+1} \mid s_t, a_t).$$

Actor Learning

- Learn a policy network $a_t \sim \pi_\theta(a_t | s_t)$
- REINFORCE algorithm [Williams 1992]

$$\mathcal{L}(\theta) = - \sum_{t=1}^T \text{sg}((R_t^\lambda - v_\phi(s_t)) / \overset{\text{Normalization}}{\max(1, S)}) \underbrace{\log \pi_\theta(a_t | s_t)}_{\text{Entropy Regularizer}} + \eta H[\pi_\theta(a_t | s_t)]$$

- Notes on policy gradient:

$$\underbrace{\pi_\theta(\tau)}_{\text{Policy}} \nabla_\theta \log \pi_\theta(\tau) = \underbrace{\pi_\theta(\tau)}_{\text{Policy}} \frac{\nabla_\theta \pi_\theta(\tau)}{\pi_\theta(\tau)} = \underbrace{\nabla_\theta \pi_\theta(\tau)}_{\text{Gradient of Policy}}.$$

$$\pi_\theta(\tau) = \pi_\theta(s_1, a_1, \dots, s_T, a_T) = p(s_1) \prod_{t=1}^T \pi_\theta(a_t | s_t) p(s_{t+1} | s_t, a_t).$$

$$\log \pi_\theta(\tau) = \log p(s_1) + \sum_{t=1}^T \log \pi_\theta(a_t | s_t) + \log p(s_{t+1} | s_t, a_t).$$

Actor Learning

- Learn a policy network $a_t \sim \pi_\theta(a_t \mid s_t)$
- REINFORCE algorithm [Williams 1992]

$$\mathcal{L}(\theta) = - \sum_{t=1}^T \text{sg}((R_t^\lambda - v_\phi(s_t)) / \overset{\text{Normalization}}{\max(1, S)}) \log \pi_\theta(a_t \mid s_t) + \overset{\text{Entropy Regularizer}}{\eta H[\pi_\theta(a_t \mid s_t)]}$$

- Notes on policy gradient:

$$\pi_\theta(\tau) \nabla_\theta \log \pi_\theta(\tau) = \pi_\theta(\tau) \frac{\nabla_\theta \pi_\theta(\tau)}{\pi_\theta(\tau)} = \nabla_\theta \pi_\theta(\tau).$$

$$\pi_\theta(\tau) = \pi_\theta(s_1, a_1, \dots, s_T, a_T) = p(s_1) \prod_{t=1}^T \pi_\theta(a_t \mid s_t) p(s_{t+1} \mid s_t, a_t).$$

$$\log \pi_\theta(\tau) = \log p(s_1) + \sum_{t=1}^T \log \pi_\theta(a_t \mid s_t) + \log p(s_{t+1} \mid s_t, a_t).$$

$$\mathcal{L}(\theta) = \mathbb{E}_{\tau \sim \pi_\theta(\tau)}[r(\tau)] = \int \pi_\theta r(\tau) d\tau.$$

Actor Learning

- Learn a policy network $a_t \sim \pi_\theta(a_t | s_t)$
- REINFORCE algorithm [Williams 1992]

$$\mathcal{L}(\theta) = - \sum_{t=1}^T \text{sg}((R_t^\lambda - v_\phi(s_t)) / \overset{\text{Normalization}}{\max(1, S)}) \log \pi_\theta(a_t | s_t) + \overset{\text{Entropy Regularizer}}{\eta H[\pi_\theta(a_t | s_t)]}$$

- Notes on policy gradient:

$$\pi_\theta(\tau) \nabla_\theta \log \pi_\theta(\tau) = \pi_\theta(\tau) \frac{\nabla_\theta \pi_\theta(\tau)}{\pi_\theta(\tau)} = \nabla_\theta \pi_\theta(\tau).$$

log. softmax over action

$$\pi_\theta(\tau) = \pi_\theta(s_1, a_1, \dots, s_T, a_T) = p(s_1) \prod_{t=1}^T \pi_\theta(a_t | s_t) p(s_{t+1} | s_t, a_t).$$

$$\log \pi_\theta(\tau) = \log p(s_1) + \sum_{t=1}^T \log \pi_\theta(a_t | s_t) + \log p(s_{t+1} | s_t, a_t).$$

$$\mathcal{L}(\theta) = \mathbb{E}_{\tau \sim \pi_\theta(\tau)}[r(\tau)] = \int \pi_\theta r(\tau) d\tau.$$

$$\nabla \mathcal{L}(\theta) = \int \nabla \pi_\theta r(\tau) d\tau = \int \pi_\theta(\tau) \nabla \log \pi_\theta r(\tau) d\tau = \mathbb{E}_{\tau \sim \pi} \log \pi_\theta r(\tau).$$

When Do We Need A Learned Actor?

- For low dimensional or discrete problems, we can directly take the argmax of the value function.

When Do We Need A Learned Actor?

- For low dimensional or discrete problems, we can directly take the argmax of the value function.
- For problems with a good model, we can roll out and sample many future trajectories. Evaluation can be done in real time with GPU. } *model-based.*

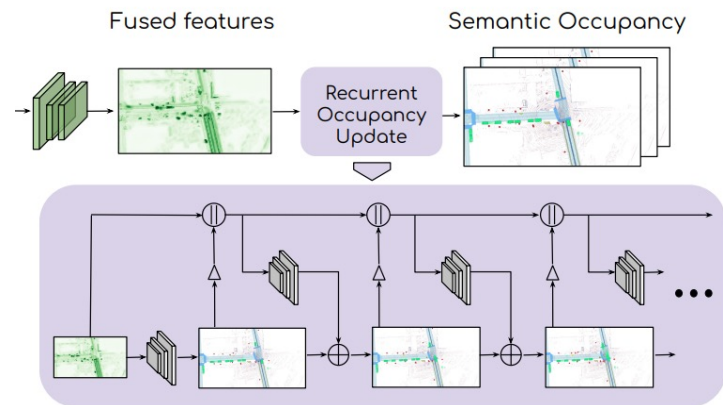
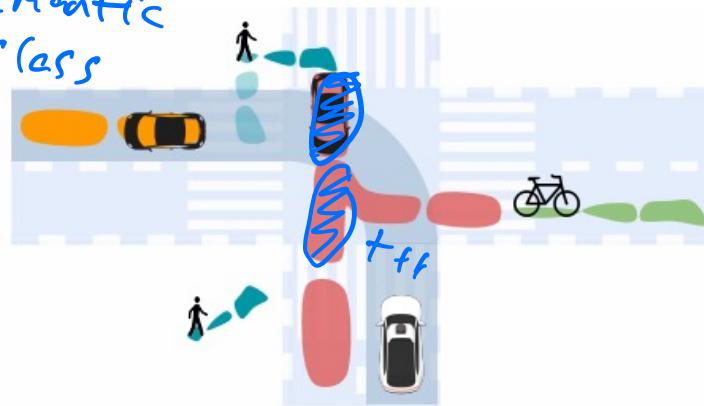
When Do We Need A Learned Actor?

- For low dimensional or discrete problems, we can directly take the argmax of the value function.
- For problems with a good model, we can roll out and sample many future trajectories. Evaluation can be done in real time with GPU.
- For general control problems, learning a separate actor can be a general solution without invoking domain knowledge.

Semantic Occupancy Volume Prediction

- 4-D volume: H, W, T, C (class)

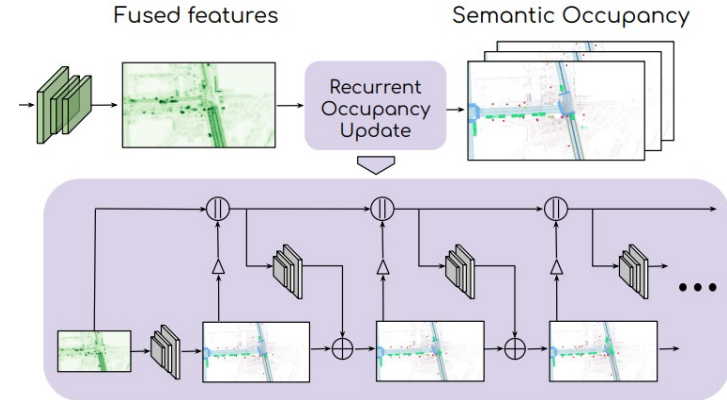
$o_{i,j}^{t,c}$ Semantic class



Semantic Occupancy Volume Prediction

- 4-D volume: H, W, T, C (class) $o_{i,j}^{t,c}$
- Doesn't grow with the increasing number of actors.

fixed dimensional.



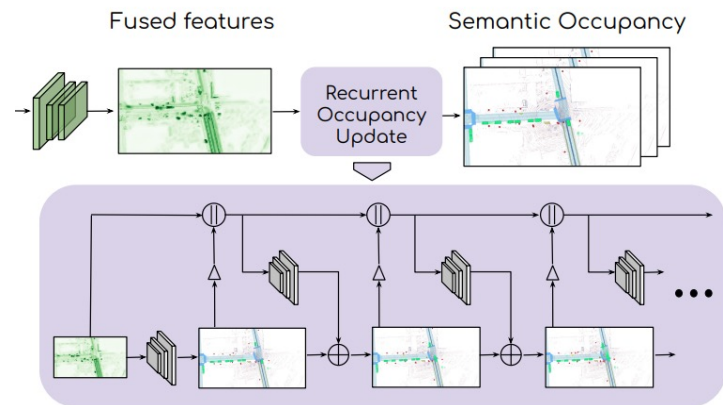
Semantic Occupancy Volume Prediction

- 4-D volume: H, W, T, C (class) $o_{i,j}^{t,c}$
- Doesn't grow with the increasing number of actors.
- Recurrent occupancy updates for further into the future.

$$l^{t,c} = l^{t-1,c} + \mathcal{U}_{\theta}^t(f_{\text{occ}}, l^{0:t-1,c})$$

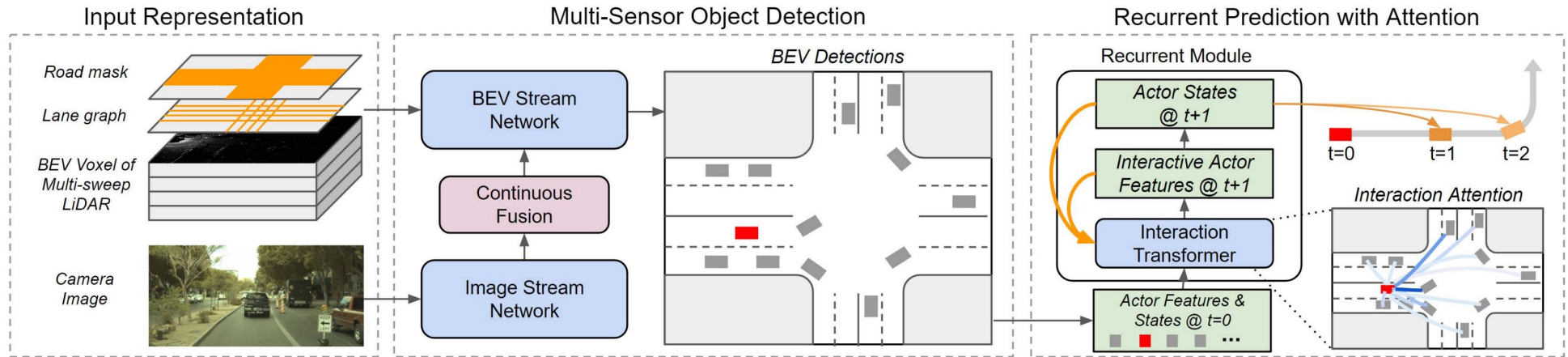
Handwritten annotations:

- $l^{t,c}$ → *occ next.*
- $l^{t-1,c}$ → *prev occ.*
- $\mathcal{U}_{\theta}^t$ → *network.*



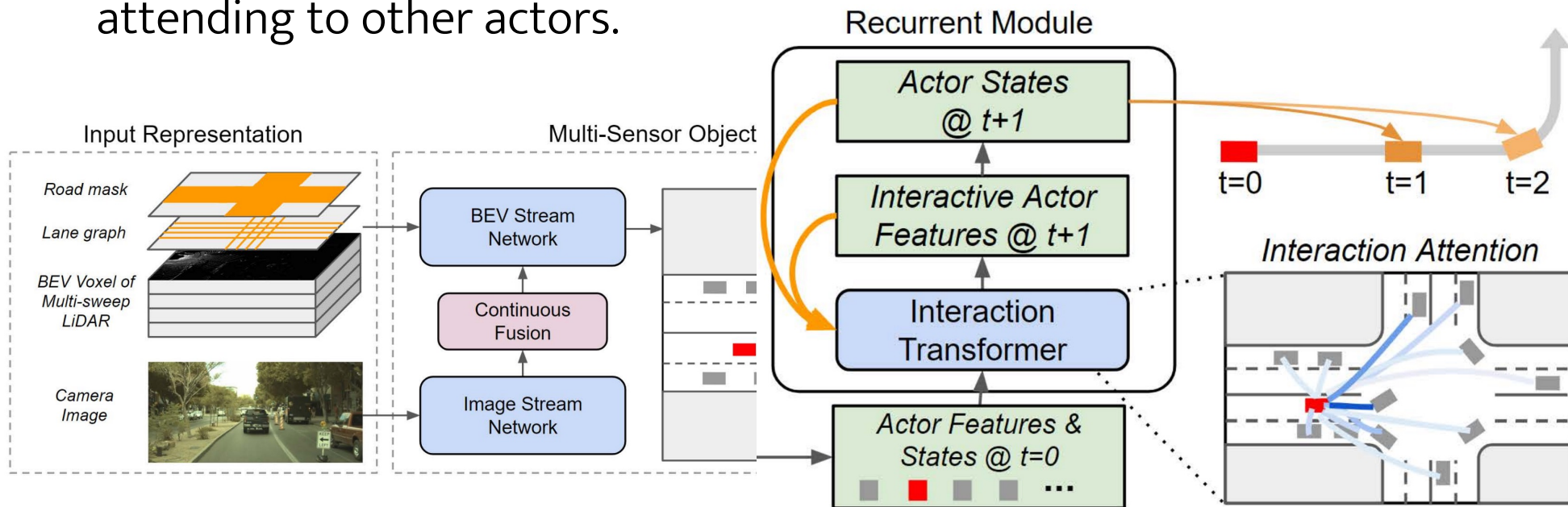
Multi-Agents Joint Prediction

- Joint predict future trajectories by attending to other actors.



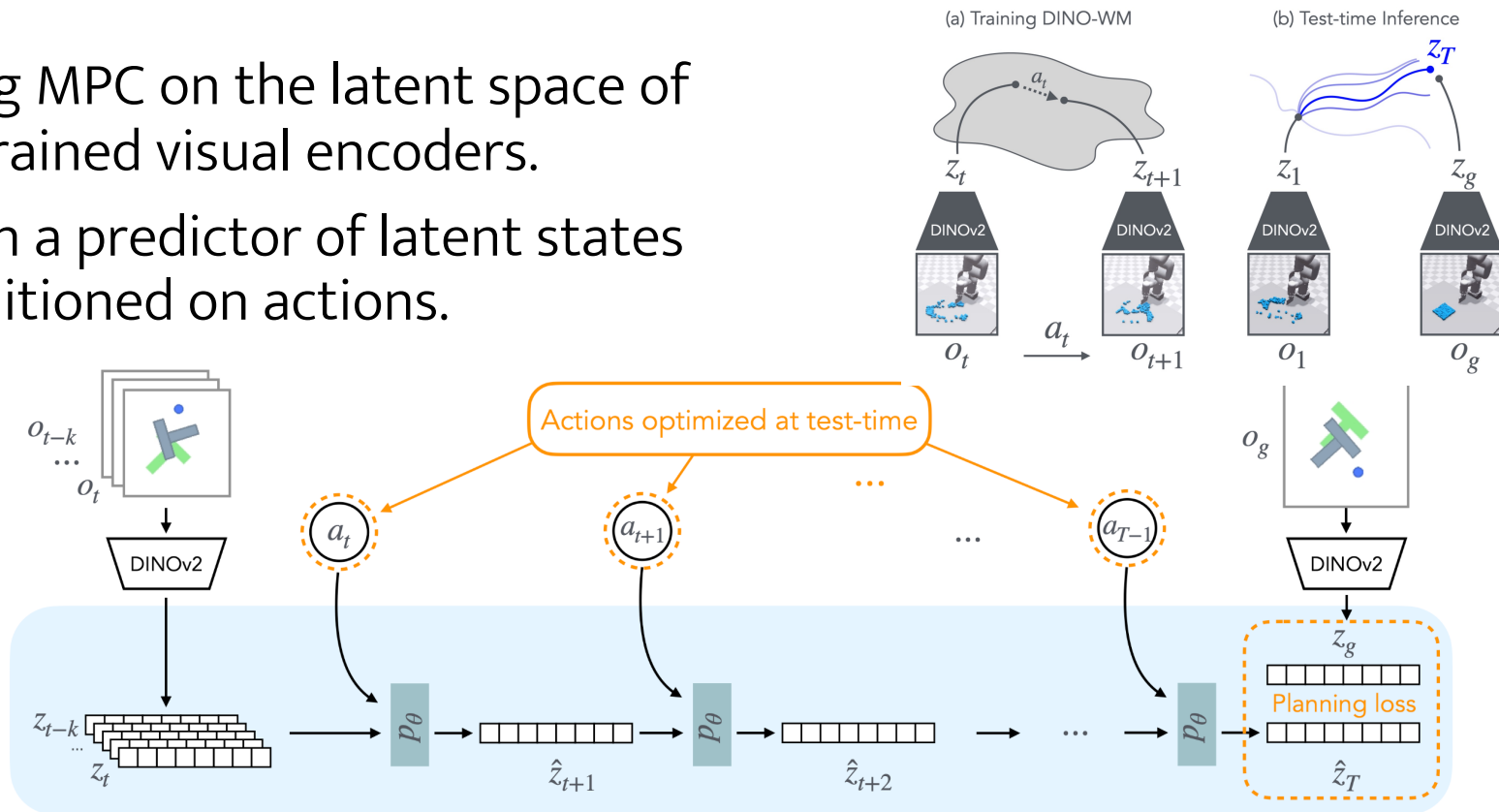
Multi-Agents Joint Prediction

- Joint predict future trajectories by attending to other actors.



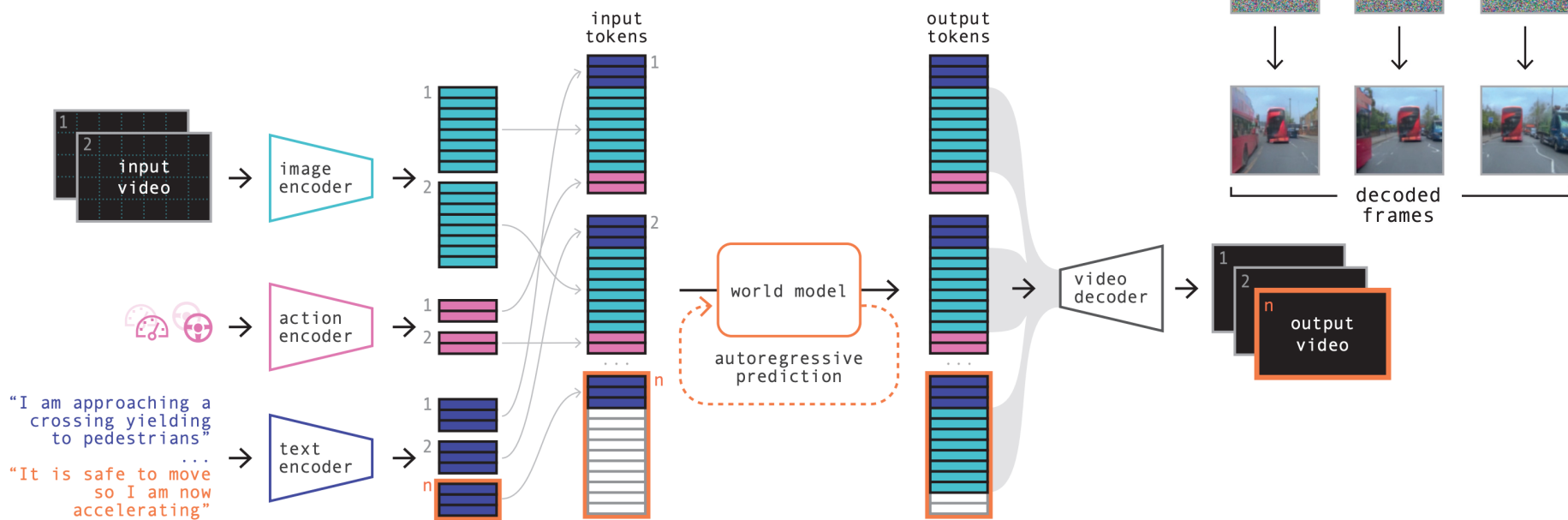
Latent Prediction + MPC

- Using MPC on the latent space of pretrained visual encoders.
- Learn a predictor of latent states conditioned on actions.

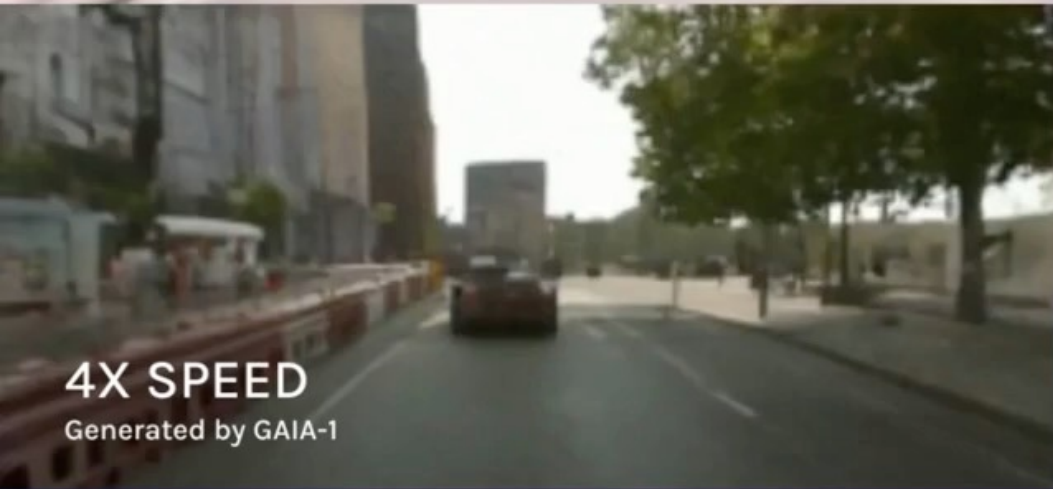


World Model in Video Prediction

- Text+action conditioned generation. Diffusion decoder.

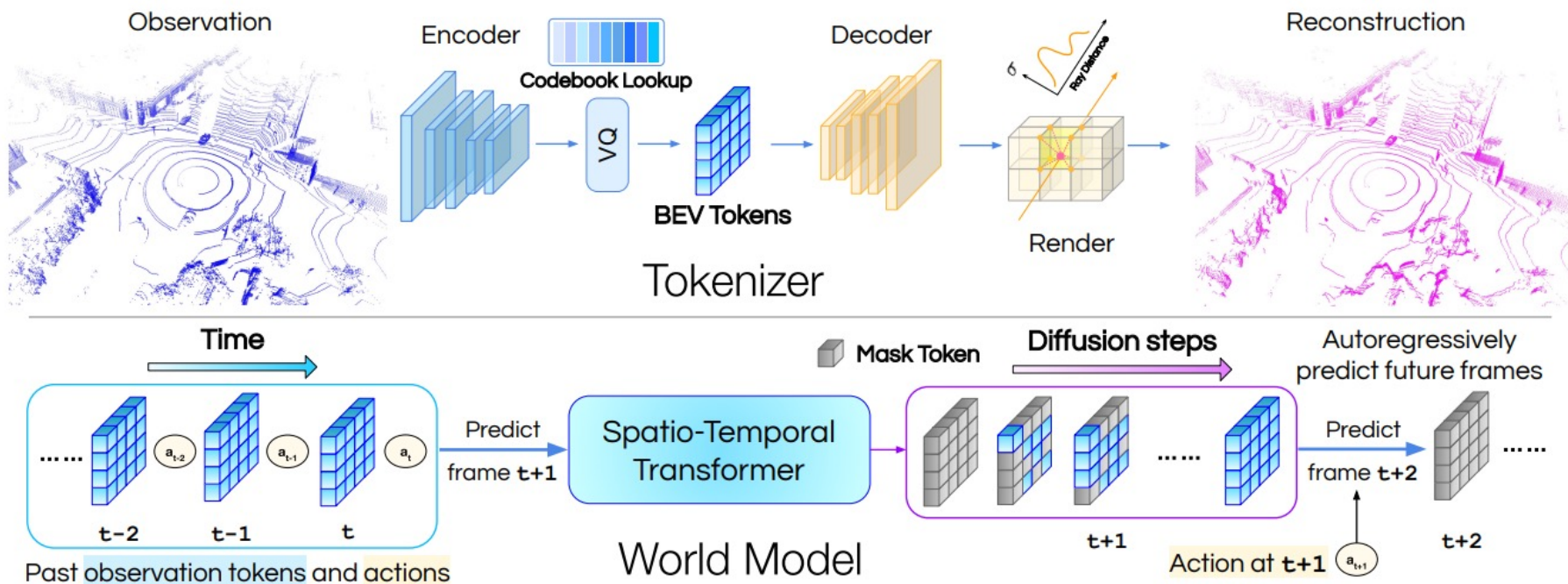


• X



World Model in 3D volume prediction

- Autoregressively predict future 3D point clouds.



Summary: World Models

Summary: World Models

- Explicit object representation
 - Traditional, light weight, instance-specific, hard to learn jointly

Summary: World Models

- Explicit object representation
 - Traditional, light weight, instance-specific, hard to learn jointly
- Differentiable occupancy, motion field
 - Relatively heavy, spatially grounded, end-to-end learnable

Summary: World Models

- Explicit object representation
 - Traditional, light weight, instance-specific, hard to learn jointly
- Differentiable occupancy, motion field
 - Relatively heavy, spatially grounded, end-to-end learnable
- Global latent, RNNs, graph landmarks
 - General-purpose, unstructured

Summary: World Models

- Explicit object representation
 - Traditional, light weight, instance-specific, hard to learn jointly
- Differentiable occupancy, motion field
 - Relatively heavy, spatially grounded, end-to-end learnable
- Global latent, RNNs, graph landmarks
 - General-purpose, unstructured
- Raw video/3D prediction
 - Expensive, good for simulation

Imitation Learning

- The explicit policy model, supervised learning (behavior cloning)

$$\hat{a} = f_{\theta}(x) \quad \mathcal{L} = \min_i \|a_i - \hat{a}\|_2^2 \quad \mathcal{L} = -\log \hat{a}_j$$

Imitation Learning

- The explicit policy model, supervised learning (behavior cloning)

$$\hat{a} = f_{\theta}(x) \quad \mathcal{L} = \min_i \|a_i - \hat{a}\|_2^2 \quad \mathcal{L} = -\log \hat{a}_j$$

- Energy-based (cost-based) approach

$$\tau^* = \operatorname{argmin}_{\tau} E(x, \tau)$$
$$p(\tau \mid x) = \frac{\exp(E(x, \tau))}{\int_{\tau} \exp(E(x, \tau))}$$

Imitation Learning

- The explicit policy model, supervised learning (behavior cloning)

$$\hat{a} = f_{\theta}(x) \quad \mathcal{L} = \min_i \|a_i - \hat{a}\|_2^2 \quad \mathcal{L} = -\log \hat{a}_j$$

- Energy-based (cost-based) approach

$$\tau^* = \operatorname{argmin}_{\tau} E(x, \tau)$$
$$p(\tau \mid x) = \frac{\exp(E(x, \tau))}{\int_{\tau} \exp(E(x, \tau))}$$

- Dataset Aggregation (Dagger)
 - Learned policy may deviate from experts
 - Need to collect more groundtruths

```
Initialize  $\mathcal{D} \leftarrow \emptyset$ .  
Initialize  $\hat{\pi}_1$  to any policy in  $\Pi$ .  
for  $i = 1$  to  $N$  do  
  Let  $\pi_i = \beta_i \pi^* + (1 - \beta_i) \hat{\pi}_i$ .  
  Sample  $T$ -step trajectories using  $\pi_i$ .  
  Get dataset  $\mathcal{D}_i = \{(s, \pi^*(s))\}$  of visited states by  $\pi_i$   
  and actions given by expert.  
  Aggregate datasets:  $\mathcal{D} \leftarrow \mathcal{D} \cup \mathcal{D}_i$ .  
  Train classifier  $\hat{\pi}_{i+1}$  on  $\mathcal{D}$ .  
end for  
Return best  $\hat{\pi}_i$  on validation.
```

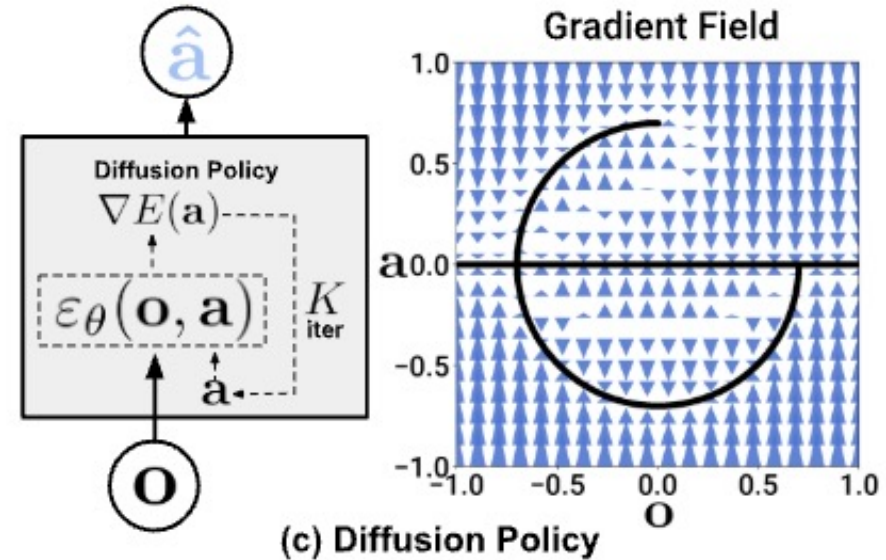
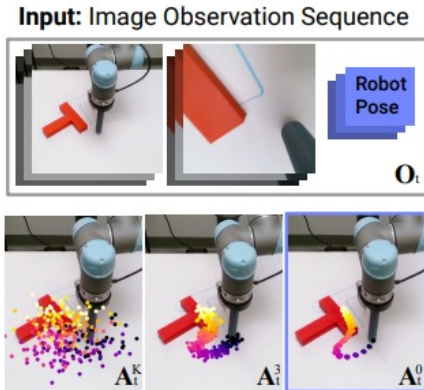
Algorithm 3.1: DAGGER Algorithm.

Direct Policy Learning from Diffusion

- Error prediction network is conditioned on observation features.

$$A_t^{k-1} = \alpha(A_t^k - \gamma \epsilon_\theta(O_t, A_t^k, k) + \mathcal{N}(0, \sigma^2 I)).$$

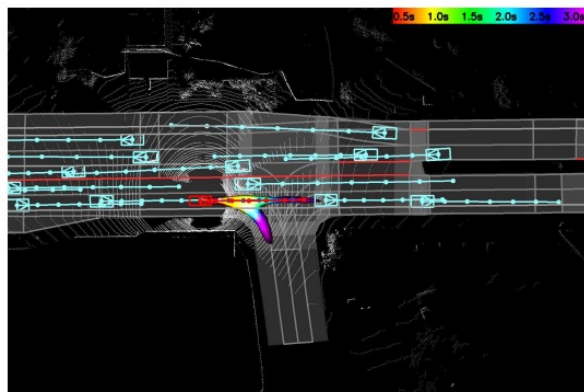
$$\mathcal{L} = MSE(\epsilon^k, \epsilon_\theta(O_t, A_t + \epsilon^k, k)).$$



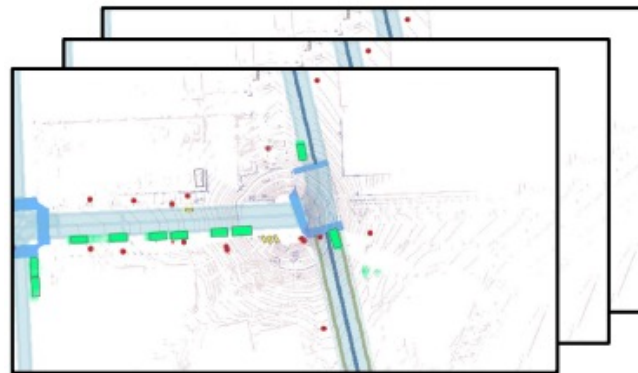
Cost/Value Volume Reasoning

- Interpretability (both costs and planner inputs)

Rasterization



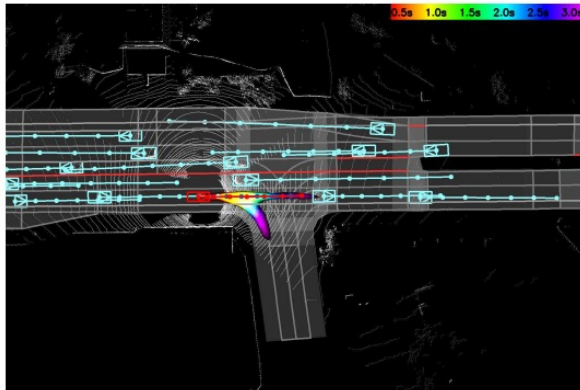
Semantic Occupancy



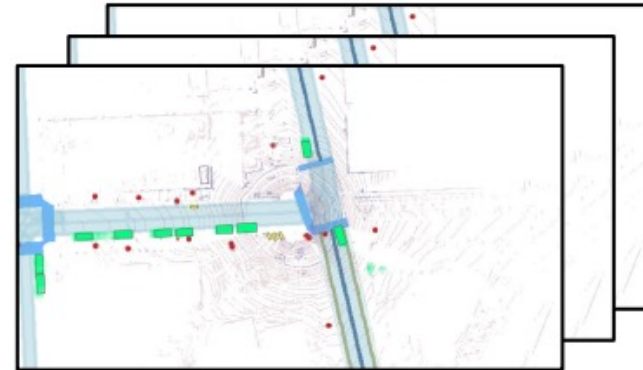
Cost/Value Volume Reasoning

- Interpretability (both costs and planner inputs)
- Use spatial geometry to form cost from explicit objects

Rasterization



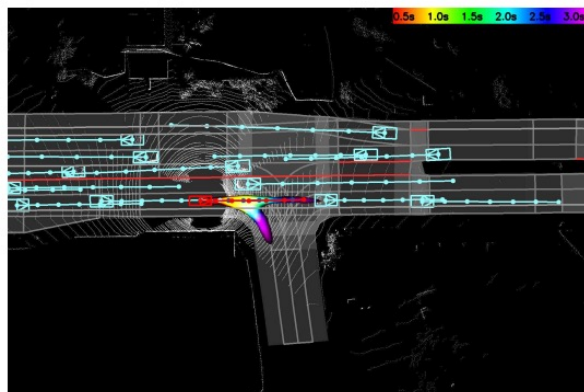
Semantic Occupancy



Cost/Value Volume Reasoning

- Interpretability (both costs and planner inputs)
- Use spatial geometry to form cost from explicit objects
- Predict spatial cost volume
 - Rasterize the scene for spatial inputs
 - Predict soft occupancy volumes (present and future)

Rasterization

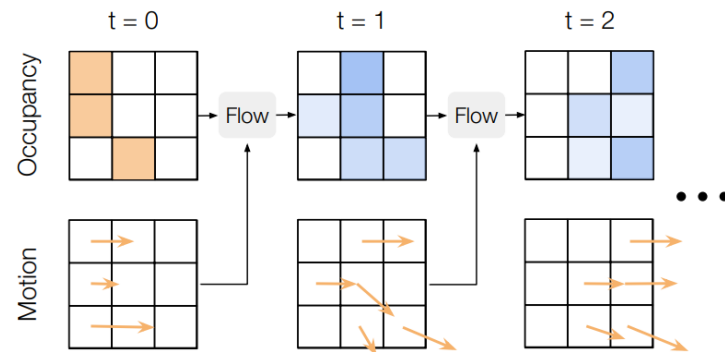
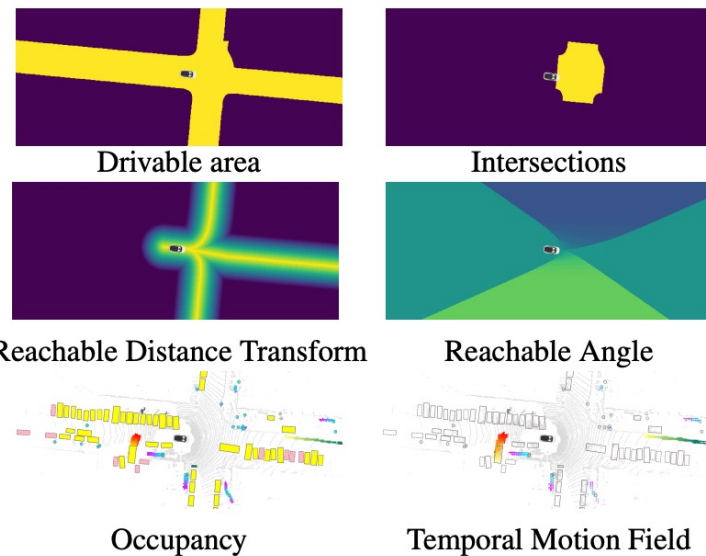
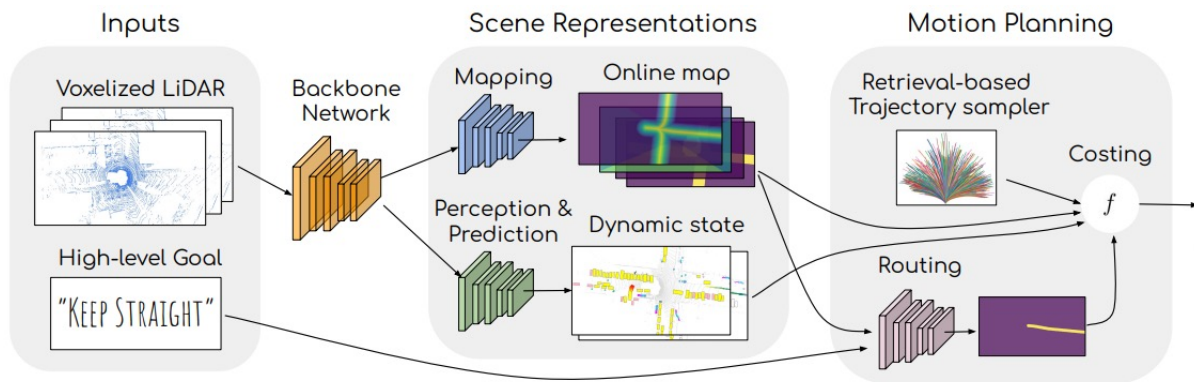


Semantic Occupancy



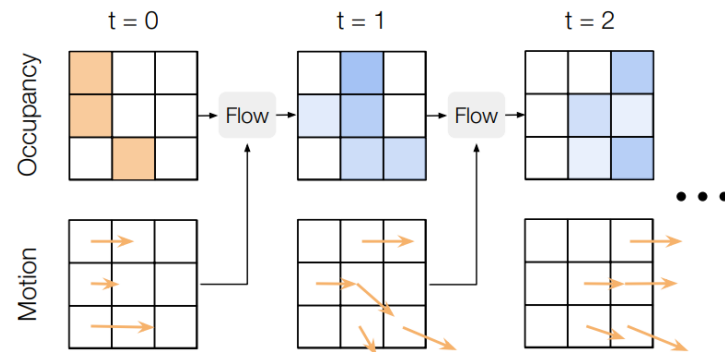
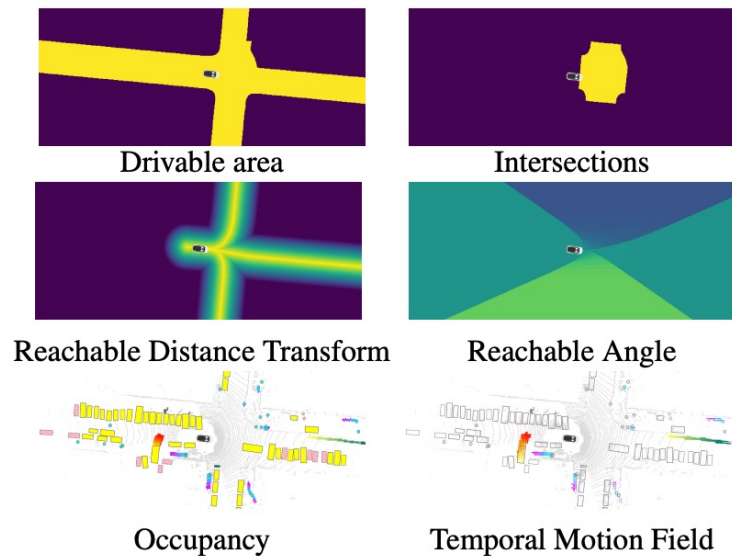
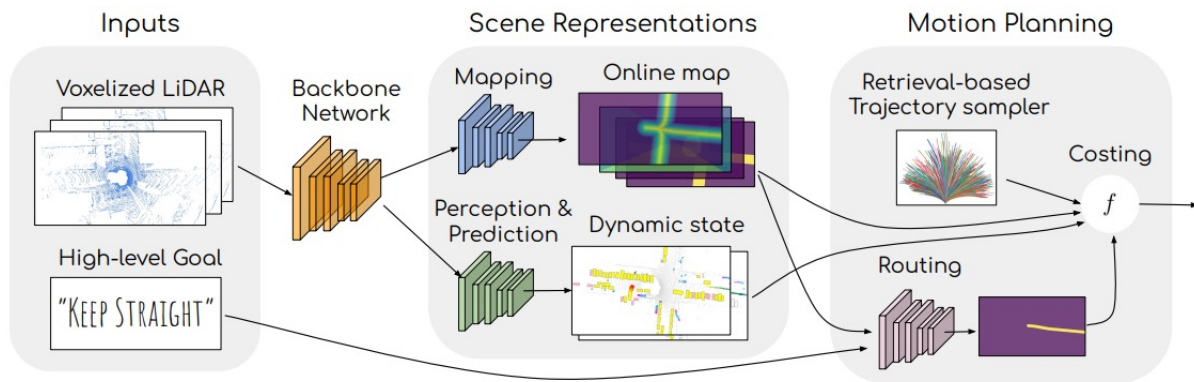
Learning Through Interpretable Predictions

- Semantic occupancy, motion field, mapping, etc. as intermediate predictions.



Learning Through Interpretable Predictions

- Semantic occupancy, motion field, mapping, etc. as intermediate predictions.
- Differentiable, supports end-to-end interpretable learning from perception to planning.



Max-Margin Planning with Explicit Cost Volume

- If we have an explicit cost volume, the cost of a trajectory can be directly queried.

Max-Margin Planning with Explicit Cost Volume

- If we have an explicit cost volume, the cost of a trajectory can be directly queried.
- We can use the max-margin objective to make the groundtruth trajectory have lower costs.

$$\operatorname{argmin}_{\theta} \sum_{\{(\hat{x}_i^t, \hat{y}_i^t)\}_{i=1 \dots N}} \max_i \sum_{t=1}^T C_{\theta}^t[x_t, y_t] - C_{\theta}^t[\hat{x}_i^t, \hat{y}_i^t] + d_i^t$$

Max-Margin Planning with Explicit Cost Volume

- If we have an explicit cost volume, the cost of a trajectory can be directly queried.
- We can use the max-margin objective to make the groundtruth trajectory have lower costs.
- Find the lowest cost trajectory among a batch of samples.

$$\operatorname{argmin}_{\theta} \sum_{\{(\hat{x}_i^t, \hat{y}_i^t)\}_{i=1 \dots N}} \max_i \sum_{t=1}^T C_{\theta}^t[x_t, y_t] - C_{\theta}^t[\hat{x}_i^t, \hat{y}_i^t] + d_i^t$$

Max-Margin Planning with Explicit Cost Volume

- If we have an explicit cost volume, the cost of a trajectory can be directly queried.
- We can use the max-margin objective to make the groundtruth trajectory have lower costs.
- Find the lowest cost trajectory among a batch of samples.
- Low-dimensional/known dynamics problems: External samplers

$$\operatorname{argmin}_{\theta} \sum_{\{(\hat{x}_i^t, \hat{y}_i^t)\}_{i=1 \dots N}} \max_i \sum_{t=1}^T C_{\theta}^t[x_t, y_t] - C_{\theta}^t[\hat{x}_i^t, \hat{y}_i^t] + d_i^t$$

Max-Margin Planning with Explicit Cost Volume

- If we have an explicit cost volume, the cost of a trajectory can be directly queried.
- We can use the max-margin objective to make the groundtruth trajectory have lower costs.
- Find the lowest cost trajectory among a batch of samples.
- Low-dimensional/known dynamics problems: External samplers
- In general, needs to perform optimization (e.g. DP)

$$\operatorname{argmin}_{\theta} \sum_{\{(\hat{x}_i^t, \hat{y}_i^t)\}_{i=1 \dots N}} \max_i \sum_{t=1}^T C_{\theta}^t[x_t, y_t] - C_{\theta}^t[\hat{x}_i^t, \hat{y}_i^t] + d_i^t$$

EBM Planning

- Energy-based framework also needs negative samples.

EBM Planning

- Energy-based framework also needs negative samples.
- “Pick” the groundtruth sample among others.

EBM Planning

- Energy-based framework also needs negative samples.
- “Pick” the groundtruth sample among others.
- If there isn’t an external sampler, we can either use autoregressive energy of sampling one dimension at a time, or gradient-based Langevin MCMC.

EBM Planning

- Energy-based framework also needs negative samples.
- “Pick” the groundtruth sample among others.
- If there isn’t an external sampler, we can either use autoregressive energy of sampling one dimension at a time, or gradient-based Langevin MCMC.

Langevin MCMC: $\tilde{\mathbf{y}}_i^k = \tilde{\mathbf{y}}_i^{k-1} - \lambda \left(\frac{1}{2} \nabla_{\mathbf{y}} E_{\theta}(\mathbf{x}_i, \mathbf{y}_i^{k-1}) + \omega^k \right), \omega^k \sim \mathcal{N}(0, \sigma).$

EBM Planning

- Energy-based framework also needs negative samples.
- “Pick” the groundtruth sample among others.
- If there isn’t an external sampler, we can either use autoregressive energy of sampling one dimension at a time, or gradient-based Langevin MCMC.

Langevin MCMC: $\tilde{\mathbf{y}}_i^k = \tilde{\mathbf{y}}_i^{k-1} - \lambda \left(\frac{1}{2} \nabla_{\mathbf{y}} E_{\theta}(\mathbf{x}_i, \mathbf{y}_i^{k-1}) + \omega^k \right), \omega^k \sim \mathcal{N}(0, \sigma).$

$$\text{Loss: } \mathcal{L} = \sum_i -\log(p_{\theta}(\mathbf{y}_i \mid \mathbf{x}, \{\tilde{\mathbf{y}}_i\}_j)) \quad p_{\theta}(\mathbf{y}_i \mid \mathbf{x}, \{\tilde{\mathbf{y}}_i\}_j) = \frac{e^{-E_{\theta}(\mathbf{x}_i, \mathbf{y}_i)}}{e^{-E_{\theta}(\mathbf{x}_i, \mathbf{y}_i)} + \sum_j e^{-E_{\theta}(\mathbf{x}_i + \tilde{\mathbf{y}}_{i,j})}}$$

Value Iteration Networks

- A network design for predicting cost volumes that are grounded from the classic value iteration algorithm.

Value Iteration Networks

- A network design for predicting cost volumes that are grounded from the classic value iteration algorithm.
- Classic VI algorithm:

Value Iteration Networks

- A network design for predicting cost volumes that are grounded from the classic value iteration algorithm.
- Classic VI algorithm:

$$V^*(s) = \max_{\pi} V^{\pi}(s)$$

Value Iteration Networks

- A network design for predicting cost volumes that are grounded from the classic value iteration algorithm.
- Classic VI algorithm:

$$V^*(s) = \max_{\pi} V^{\pi}(s)$$

$$V^{\pi}(s) = \mathbb{E}^{\pi} \sum_{t=0}^{\infty} \gamma^t r(s_t, a_t)$$

Value Iteration Networks

- A network design for predicting cost volumes that are grounded from the classic value iteration algorithm.
- Classic VI algorithm:

$$V^*(s) = \max_{\pi} V^{\pi}(s)$$

$$V^{\pi}(s) = \mathbb{E}^{\pi} \sum_{t=0}^{\infty} \gamma^t r(s_t, a_t)$$

$$Q_n(s, a) = R(s, a) + \gamma \sum_{s'} P(s'|s, a) V_n(s')$$

Value Iteration Networks

- A network design for predicting cost volumes that are grounded from the classic value iteration algorithm.
- Classic VI algorithm:

$$V^*(s) = \max_{\pi} V^{\pi}(s)$$

$$V^{\pi}(s) = \mathbb{E}^{\pi} \sum_{t=0}^{\infty} \gamma^t r(s_t, a_t)$$

$$Q_n(s, a) = R(s, a) + \gamma \sum_{s'} P(s'|s, a) V_n(s')$$

$$V_{n+1}(s) = \max_a Q_n(s, a)$$

Value Iteration Networks

- A network design for predicting cost volumes that are grounded from the classic value iteration algorithm.
- Classic VI algorithm:

$$V^*(s) = \max_{\pi} V^{\pi}(s)$$

$$V^{\pi}(s) = \mathbb{E}^{\pi} \sum_{t=0}^{\infty} \gamma^t r(s_t, a_t)$$

$$Q_n(s, a) = R(s, a) + \gamma \sum_{s'} P(s'|s, a) V_n(s')$$

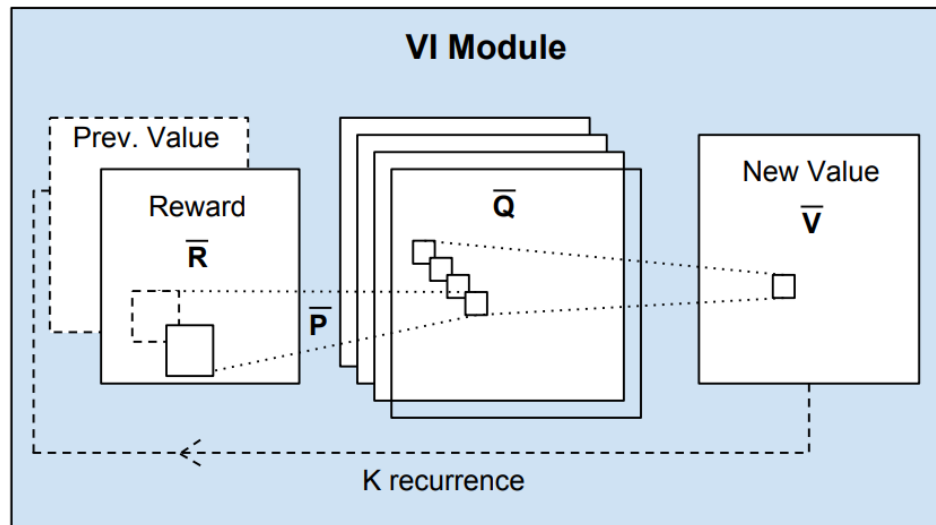
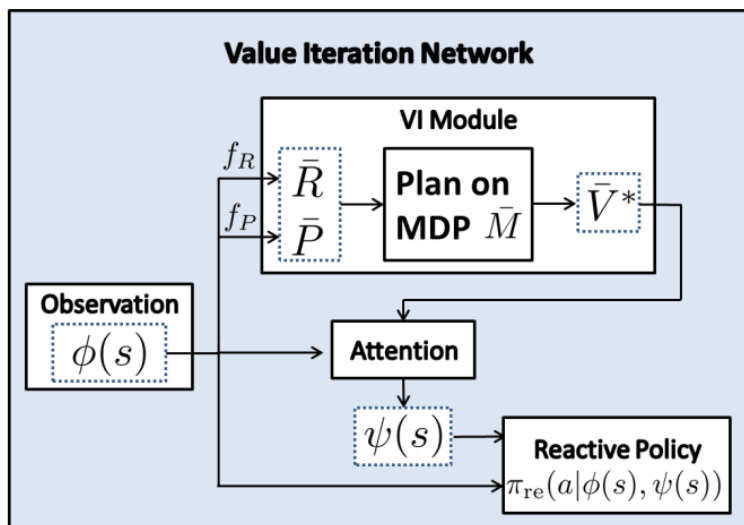
$$V_{n+1}(s) = \max_a Q_n(s, a)$$

$$\pi^*(s) = \operatorname{argmax}_a Q_{\infty}(s, a)$$

Value Iteration Networks

- Reward and previous value are fed into a CNN to generate Q of A channels. Transition matrix is convolutional kernel. Then Max-Pooling.

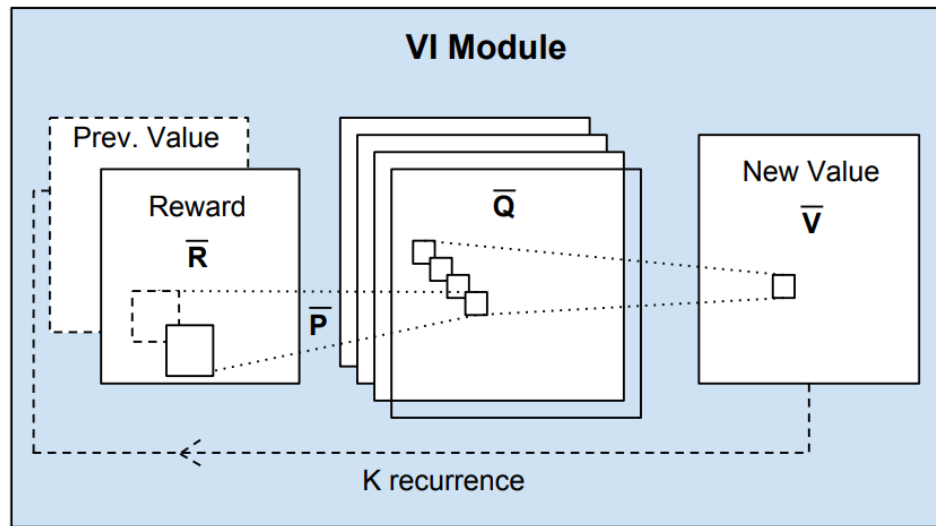
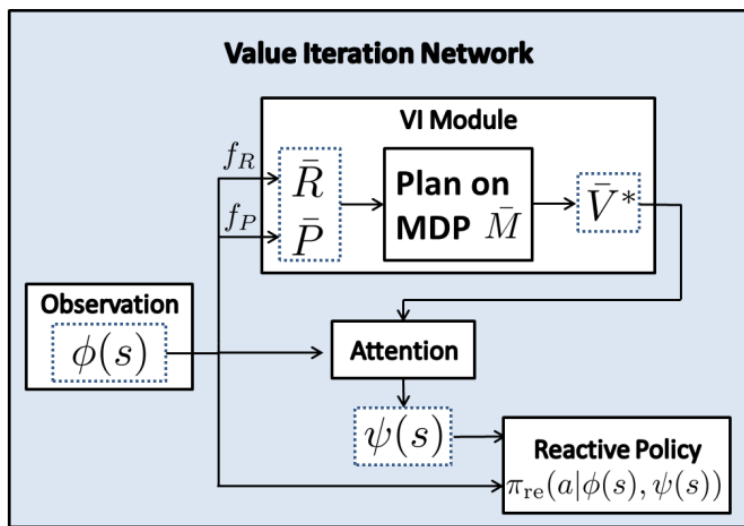
$$Q_n(s, a) = R(s, a) + \gamma \sum_{s'} P(s'|s, a) V_n(s') \quad V_{n+1}(s) = \max_a Q_n(s, a)$$



Value Iteration Networks

- Select the current state and choose an action from softmax.

$$\hat{a} \sim \text{softmax}_a(Q(s, a))$$



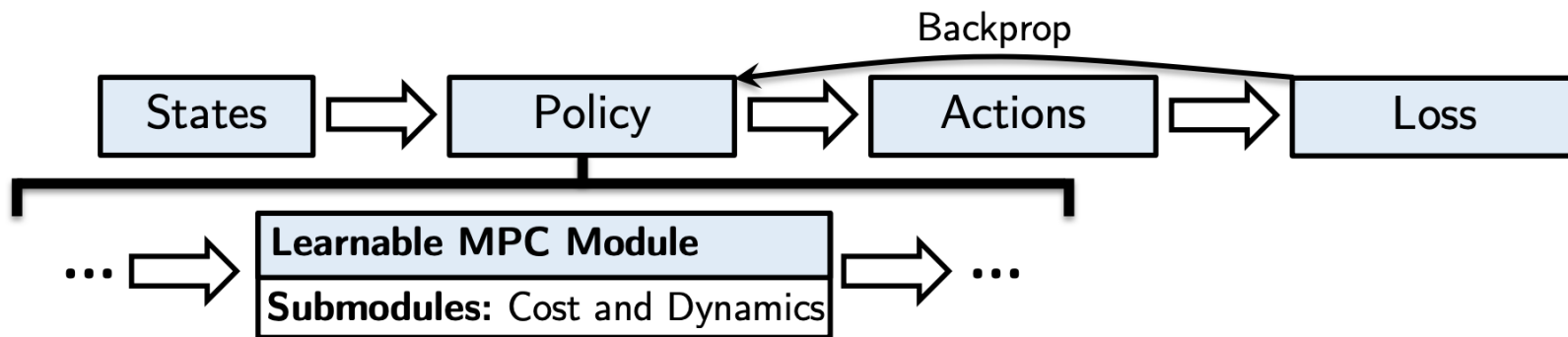
Value Iteration Networks

- A baseline would be to untie the weights through iterations, more like a feedforward CNN.
- Achieve more training data efficiency by imposing the structure.

Training data	VIN			VIN Untied Weights		
	Pred. loss	Succ. rate	Traj. diff.	Pred. loss	Succ. rate	Traj. diff.
20%	0.06	98.2%	0.106	0.09	91.9%	0.094
50%	0.05	99.4%	0.018	0.07	95.2%	0.078
100%	0.05	99.3%	0.089	0.05	95.6%	0.068

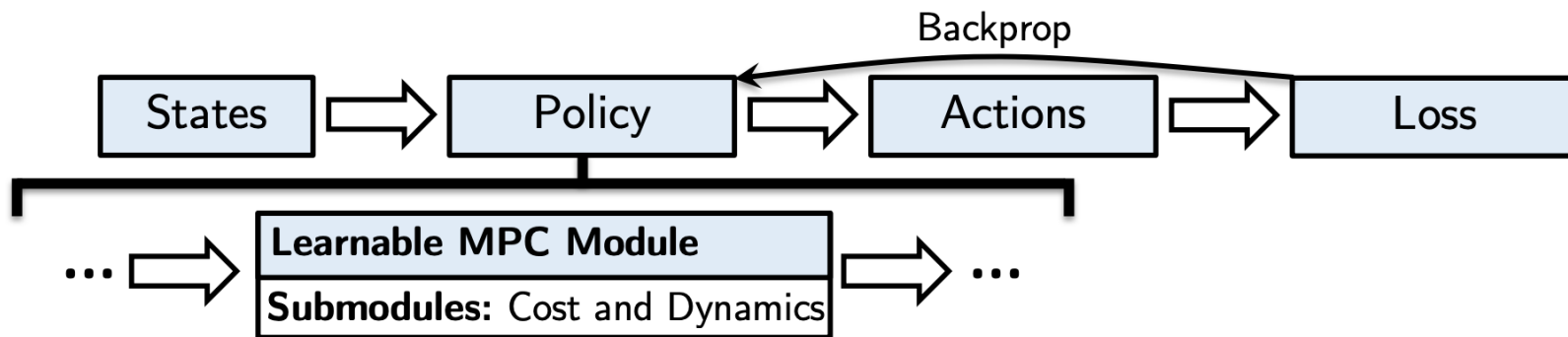
Backprop through Planning

- Treat planning as an end-to-end layer. Can be used for RL/Imitation.



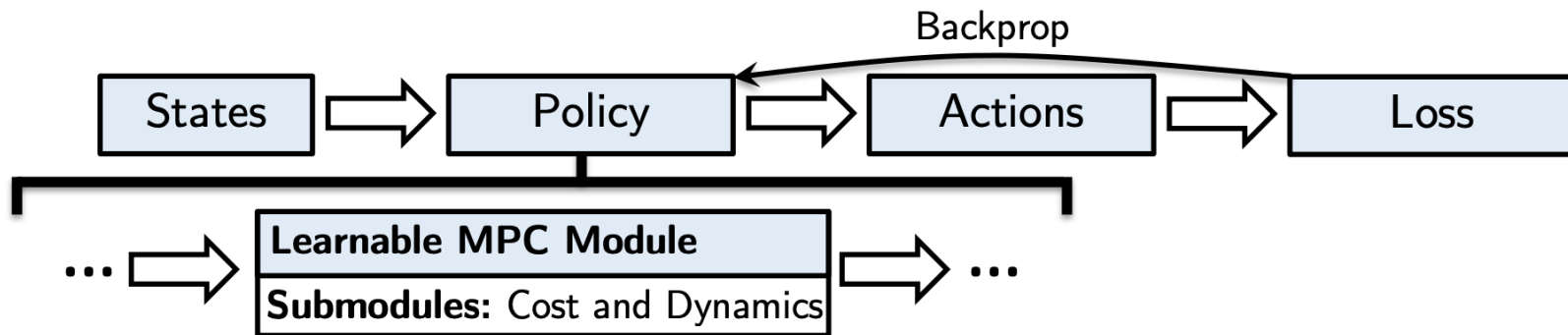
Backprop through Planning

- Treat planning as an end-to-end layer. Can be used for RL/Imitation.
- Option 1: Unrolling a finite number of steps



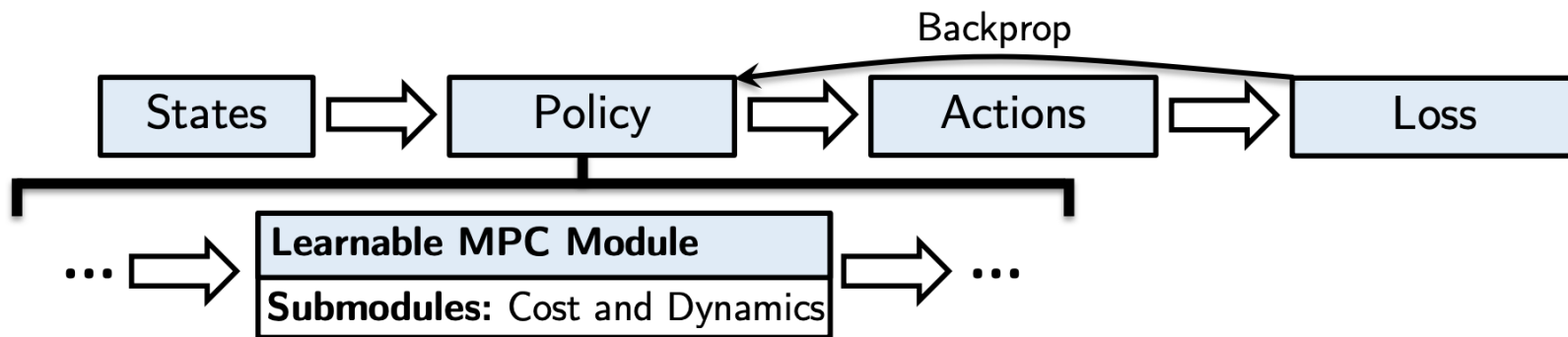
Backprop through Planning

- Treat planning as an end-to-end layer. Can be used for RL/Imitation.
- Option 1: Unrolling a finite number of steps
- Option 2: Solve till convergence, backprop for a finite step



Backprop through Planning

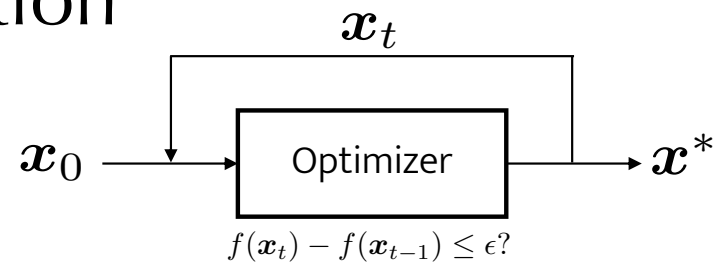
- Treat planning as an end-to-end layer. Can be used for RL/Imitation.
- Option 1: Unrolling a finite number of steps
- Option 2: Solve till convergence, backprop for a finite step
- Option 3: Converged at fixed point: Implicit differentiation



Implicit Differentiation

- Unconstrained case

$$\mathbf{x}^* = \operatorname{argmin}_{\mathbf{x}} f(\mathbf{x}; \boldsymbol{\theta}).$$

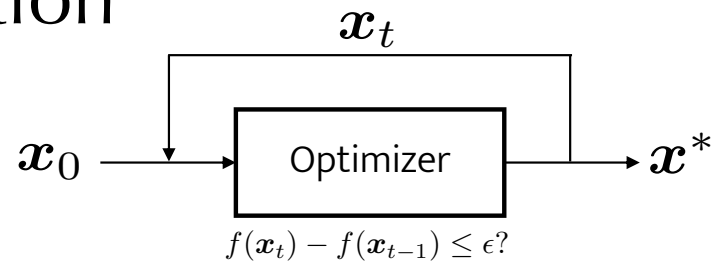


Implicit Differentiation

- Unconstrained case

$$\mathbf{x}^* = \operatorname{argmin}_{\mathbf{x}} f(\mathbf{x}; \boldsymbol{\theta}).$$

$$\mathbf{0} = \frac{\mathrm{d}}{\mathrm{d}\boldsymbol{\theta}} J_{f, \mathbf{x}^*}(\mathbf{x}^*; \boldsymbol{\theta})$$



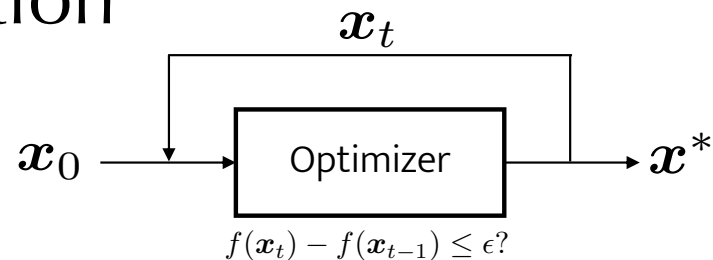
Implicit Differentiation

- Unconstrained case

$$\mathbf{x}^* = \operatorname{argmin}_{\mathbf{x}} f(\mathbf{x}; \boldsymbol{\theta}).$$

$$\mathbf{0} = \frac{\mathrm{d}}{\mathrm{d}\boldsymbol{\theta}} J_{f, \mathbf{x}^*}(\mathbf{x}^*; \boldsymbol{\theta})$$

$$\mathbf{0} = \frac{\partial}{\partial \mathbf{x}^*} J_{f, \mathbf{x}^*}(\mathbf{x}^*; \boldsymbol{\theta}) \frac{\partial \mathbf{x}^*}{\partial \boldsymbol{\theta}} + \frac{\partial}{\partial \boldsymbol{\theta}} J_{f, \mathbf{x}^*}(\mathbf{x}^*; \boldsymbol{\theta})$$



Implicit Differentiation

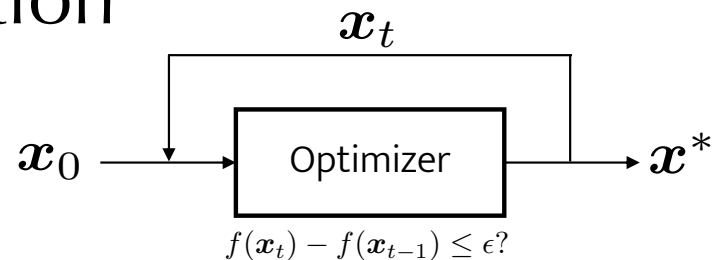
- Unconstrained case

$$\mathbf{x}^* = \operatorname{argmin}_{\mathbf{x}} f(\mathbf{x}; \boldsymbol{\theta}).$$

$$\mathbf{0} = \frac{\mathrm{d}}{\mathrm{d}\boldsymbol{\theta}} J_{f, \mathbf{x}^*}(\mathbf{x}^*; \boldsymbol{\theta})$$

$$\mathbf{0} = \frac{\partial}{\partial \mathbf{x}^*} J_{f, \mathbf{x}^*}(\mathbf{x}^*; \boldsymbol{\theta}) \frac{\partial \mathbf{x}^*}{\partial \boldsymbol{\theta}} + \frac{\partial}{\partial \boldsymbol{\theta}} J_{f, \mathbf{x}^*}(\mathbf{x}^*; \boldsymbol{\theta})$$

$$\mathbf{0} = H_{f, \mathbf{x}^*}(\mathbf{x}^*; \boldsymbol{\theta}) \frac{\partial \mathbf{x}^*}{\partial \boldsymbol{\theta}} + \frac{\partial}{\partial \boldsymbol{\theta}} J_{f, \mathbf{x}^*}(\mathbf{x}^*; \boldsymbol{\theta})$$



Implicit Differentiation

- Unconstrained case

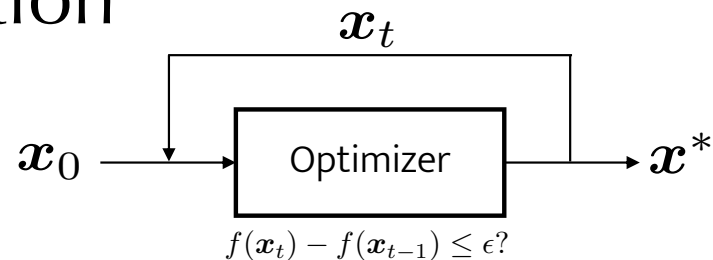
$$\mathbf{x}^* = \operatorname{argmin}_{\mathbf{x}} f(\mathbf{x}; \boldsymbol{\theta}).$$

$$\mathbf{0} = \frac{\mathrm{d}}{\mathrm{d}\boldsymbol{\theta}} J_{f, \mathbf{x}^*}(\mathbf{x}^*; \boldsymbol{\theta})$$

$$\mathbf{0} = \frac{\partial}{\partial \mathbf{x}^*} J_{f, \mathbf{x}^*}(\mathbf{x}^*; \boldsymbol{\theta}) \frac{\partial \mathbf{x}^*}{\partial \boldsymbol{\theta}} + \frac{\partial}{\partial \boldsymbol{\theta}} J_{f, \mathbf{x}^*}(\mathbf{x}^*; \boldsymbol{\theta})$$

$$\mathbf{0} = H_{f, \mathbf{x}^*}(\mathbf{x}^*; \boldsymbol{\theta}) \frac{\partial \mathbf{x}^*}{\partial \boldsymbol{\theta}} + \frac{\partial}{\partial \boldsymbol{\theta}} J_{f, \mathbf{x}^*}(\mathbf{x}^*; \boldsymbol{\theta})$$

$$\frac{\partial \mathbf{x}^*}{\partial \boldsymbol{\theta}} = H_{f, \mathbf{x}^*}(\mathbf{x}^*; \boldsymbol{\theta})^{-1} \frac{\partial}{\partial \boldsymbol{\theta}} J_{f, \mathbf{x}^*}(\mathbf{x}^*; \boldsymbol{\theta}).$$



Implicit Differentiation

- How to compute Hessian inverse vector product?

Implicit Differentiation

- How to compute Hessian inverse vector product?
- Conjugate gradient, solve $Ax = b$

Conjugate Gradient Method

$\mathbf{r}_0 := \mathbf{b} - \mathbf{A}\mathbf{x}_0$

if $\mathbf{r}_0$ is sufficiently small, then return $\mathbf{x}_0$ as the result

$\mathbf{p}_0 := \mathbf{r}_0$

$k := 0$

repeat

$$\alpha_k := \frac{\mathbf{r}_k^\top \mathbf{r}_k}{\mathbf{p}_k^\top \mathbf{A} \mathbf{p}_k}$$

$$\mathbf{x}_{k+1} := \mathbf{x}_k + \alpha_k \mathbf{p}_k$$

$$\mathbf{r}_{k+1} := \mathbf{r}_k - \alpha_k \mathbf{A} \mathbf{p}_k$$

if $\mathbf{r}_{k+1}$ is sufficiently small, then exit loop

$$\beta_k := \frac{\mathbf{r}_{k+1}^\top \mathbf{r}_{k+1}}{\mathbf{r}_k^\top \mathbf{r}_k}$$

$$\mathbf{p}_{k+1} := \mathbf{r}_{k+1} + \beta_k \mathbf{p}_k$$

$$k := k + 1$$

end repeat

return $\mathbf{x}_{k+1}$ as the result

Implicit Differentiation

- How to compute Hessian inverse vector product?
- Conjugate gradient, solve $Ax = b$
- Neumann series (finite truncation)

$$(I - A)^{-1} = \sum_{k=0}^{\infty} A^k.$$

- Same as backprop the last K steps (Option 2).
- Memory savings.

Conjugate Gradient Method

$\mathbf{r}_0 := \mathbf{b} - \mathbf{A}\mathbf{x}_0$

if $\mathbf{r}_0$ is sufficiently small, then return $\mathbf{x}_0$ as the result

$\mathbf{p}_0 := \mathbf{r}_0$

$k := 0$

repeat

$$\alpha_k := \frac{\mathbf{r}_k^\top \mathbf{r}_k}{\mathbf{p}_k^\top \mathbf{A} \mathbf{p}_k}$$

$$\mathbf{x}_{k+1} := \mathbf{x}_k + \alpha_k \mathbf{p}_k$$

$$\mathbf{r}_{k+1} := \mathbf{r}_k - \alpha_k \mathbf{A} \mathbf{p}_k$$

if $\mathbf{r}_{k+1}$ is sufficiently small, then exit loop

$$\beta_k := \frac{\mathbf{r}_{k+1}^\top \mathbf{r}_{k+1}}{\mathbf{r}_k^\top \mathbf{r}_k}$$

$$\mathbf{p}_{k+1} := \mathbf{r}_{k+1} + \beta_k \mathbf{p}_k$$

$$k := k + 1$$

end repeat

return $\mathbf{x}_{k+1}$ as the result

Differentiable LQR

- Now add linear equality constraints on the dynamics and initialization.

$$\tau_{1:T} = \{x_t, u_t\}_{1:T}$$

$$\operatorname{argmin}_{\tau_{1:T}} \sum_{t=1}^T \frac{1}{2} \tau_t^\top C_t \tau_t$$

$$\text{subject to } x_{t+1} = F_t \tau_t + f_t, x_1 = x_{\text{init}}.$$

Differentiable LQR

- Now add linear equality constraints on the dynamics and initialization.

$$\tau_{1:T} = \{x_t, u_t\}_{1:T}$$

- Chain rule: $\frac{\partial \ell}{\partial \theta} = \frac{\partial \ell}{\partial \tau_{1:T}^*} \frac{\partial \tau_{1:T}^*}{\partial \theta}.$

$$\operatorname{argmin}_{\tau_{1:T}} \sum_{t=1}^T \frac{1}{2} \tau_t^\top C_t \tau_t$$

$$\text{subject to } x_{t+1} = F_t \tau_t + f_t, x_1 = x_{\text{init}}.$$

Differentiable LQR

- Now add linear equality constraints on the dynamics and initialization.

$$\tau_{1:T} = \{x_t, u_t\}_{1:T}$$

- Chain rule: $\frac{\partial \ell}{\partial \theta} = \frac{\partial \ell}{\partial \tau_{1:T}^*} \frac{\partial \tau_{1:T}^*}{\partial \theta}.$

$$\operatorname{argmin}_{\tau_{1:T}} \sum_{t=1}^T \frac{1}{2} \tau_t^\top C_t \tau_t$$

subject to $x_{t+1} = F_t \tau_t + f_t, x_1 = x_{\text{init}}.$

General QP:

$$\mathbf{x}^* = \operatorname{argmin} \frac{1}{2} \mathbf{x}^\top Q \mathbf{x} + \mathbf{c}^\top \mathbf{x}$$

subject to $A\mathbf{x} = \mathbf{b}.$

Differentiable LQR

- Now add linear equality constraints on the dynamics and initialization.

$$\tau_{1:T} = \{x_t, u_t\}_{1:T}$$

- Chain rule: $\frac{\partial \ell}{\partial \theta} = \frac{\partial \ell}{\partial \tau_{1:T}^*} \frac{\partial \tau_{1:T}^*}{\partial \theta}.$
- KKT:

$$\operatorname{argmin}_{\tau_{1:T}} \sum_{t=1}^T \frac{1}{2} \tau_t^\top C_t \tau_t$$

subject to $x_{t+1} = F_t \tau_t + f_t, x_1 = x_{\text{init}}.$

General QP:

$$\mathbf{x}^* = \operatorname{argmin} \frac{1}{2} \mathbf{x}^\top Q \mathbf{x} + \mathbf{c}^\top \mathbf{x}$$

subject to $A\mathbf{x} = \mathbf{b}.$

Differentiable LQR

- Now add linear equality constraints on the dynamics and initialization.

$$\tau_{1:T} = \{x_t, u_t\}_{1:T}$$

- Chain rule: $\frac{\partial \ell}{\partial \theta} = \frac{\partial \ell}{\partial \tau_{1:T}^*} \frac{\partial \tau_{1:T}^*}{\partial \theta}$.
- KKT:

$$\begin{bmatrix} Q & A^\top \\ A & \mathbf{0} \end{bmatrix} \begin{bmatrix} \mathbf{x}^* \\ \lambda^* \end{bmatrix} = \begin{bmatrix} -\mathbf{c} \\ \mathbf{b} \end{bmatrix}$$

$$K \begin{bmatrix} \mathbf{x}^* \\ \lambda^* \end{bmatrix} = \mathbf{v}.$$

In classic LQR solver, the Riccati recursion solves this linear system.

$$\operatorname{argmin}_{\tau_{1:T}} \sum_{t=1}^T \frac{1}{2} \tau_t^\top C_t \tau_t$$

subject to $x_{t+1} = F_t \tau_t + f_t, x_1 = x_{\text{init}}$.

General QP:

$$\mathbf{x}^* = \operatorname{argmin} \frac{1}{2} \mathbf{x}^\top Q \mathbf{x} + \mathbf{c}^\top \mathbf{x}$$

subject to $A\mathbf{x} = \mathbf{b}$.

Differentiable LQR

- Apply differentiation

$$\frac{d}{d\theta} \left(K \begin{bmatrix} \mathbf{x}^* \\ \lambda^* \end{bmatrix} \right) = \frac{dv}{d\theta}.$$
$$\frac{dK}{d\theta} \begin{bmatrix} \mathbf{x}^* \\ \lambda^* \end{bmatrix} + K \begin{bmatrix} \frac{d\mathbf{x}^*}{d\theta} \\ \frac{d\lambda^*}{d\theta} \end{bmatrix} = \frac{dv}{d\theta}$$

Differentiable LQR

- Apply differentiation

$$\frac{d}{d\theta} \left(K \begin{bmatrix} \mathbf{x}^* \\ \lambda^* \end{bmatrix} \right) = \frac{dv}{d\theta}.$$

$$\frac{dK}{d\theta} \begin{bmatrix} \mathbf{x}^* \\ \lambda^* \end{bmatrix} + K \begin{bmatrix} \frac{d\mathbf{x}^*}{d\theta} \\ \frac{d\lambda^*}{d\theta} \end{bmatrix} = \frac{dv}{d\theta}$$

$$K \begin{bmatrix} \frac{d\mathbf{x}^*}{d\theta} \\ \frac{d\lambda^*}{d\theta} \end{bmatrix} = K \begin{bmatrix} \frac{d\mathbf{x}^*}{d\mathbf{c}} & \frac{d\mathbf{x}^*}{d\mathbf{b}} & \frac{d\mathbf{x}^*}{dQ} & \frac{d\mathbf{x}^*}{dA} \\ \frac{d\lambda^*}{d\mathbf{c}} & \frac{d\lambda^*}{d\mathbf{b}} & \frac{d\lambda^*}{dQ} & \frac{d\lambda^*}{dA} \end{bmatrix} = \begin{bmatrix} -I & 0 & -\mathbf{x}^* & -\lambda^* \\ 0 & I & 0 & -\mathbf{x}^* \end{bmatrix}$$

Differentiable LQR

- Apply differentiation

$$\frac{d}{d\theta} \left(K \begin{bmatrix} \mathbf{x}^* \\ \lambda^* \end{bmatrix} \right) = \frac{dv}{d\theta}.$$

$$\frac{dK}{d\theta} \begin{bmatrix} \mathbf{x}^* \\ \lambda^* \end{bmatrix} + K \begin{bmatrix} \frac{d\mathbf{x}^*}{d\theta} \\ \frac{d\lambda^*}{d\theta} \end{bmatrix} = \frac{dv}{d\theta}$$

$$K \begin{bmatrix} \frac{d\mathbf{x}^*}{d\theta} \\ \frac{d\lambda^*}{d\theta} \end{bmatrix} = K \begin{bmatrix} \frac{d\mathbf{x}^*}{d\mathbf{c}} & \frac{d\mathbf{x}^*}{d\mathbf{b}} & \frac{d\mathbf{x}^*}{dQ} & \frac{d\mathbf{x}^*}{dA} \\ \frac{d\lambda^*}{d\mathbf{c}} & \frac{d\lambda^*}{d\mathbf{b}} & \frac{d\lambda^*}{dQ} & \frac{d\lambda^*}{dA} \end{bmatrix} = \begin{bmatrix} -I & 0 & -x^* & -\lambda^* \\ 0 & I & 0 & -x^* \end{bmatrix} K \frac{\partial \ell}{\partial \mathbf{z}^*} \begin{bmatrix} \frac{d\mathbf{x}^*}{d\mathbf{c}} & \frac{d\mathbf{x}^*}{d\mathbf{b}} \\ \frac{d\lambda^*}{d\mathbf{c}} & \frac{d\lambda^*}{d\mathbf{b}} \end{bmatrix} = \begin{bmatrix} -\frac{\partial \ell}{\partial \mathbf{x}^*} \\ \mathbf{0} \end{bmatrix}$$

$$K d^* = \begin{bmatrix} -\frac{\partial \ell}{\partial \mathbf{x}^*} \\ \mathbf{0} \end{bmatrix}$$

Differentiable LQR

- Apply differentiation

$$\frac{d}{d\theta} \left(K \begin{bmatrix} \mathbf{x}^* \\ \lambda^* \end{bmatrix} \right) = \frac{dv}{d\theta}.$$

$$\frac{dK}{d\theta} \begin{bmatrix} \mathbf{x}^* \\ \lambda^* \end{bmatrix} + K \begin{bmatrix} \frac{d\mathbf{x}^*}{d\theta} \\ \frac{d\lambda^*}{d\theta} \end{bmatrix} = \frac{dv}{d\theta}$$

$$K \begin{bmatrix} \frac{d\mathbf{x}^*}{d\theta} \\ \frac{d\lambda^*}{d\theta} \end{bmatrix} = K \begin{bmatrix} \frac{d\mathbf{x}^*}{d\mathbf{c}} & \frac{d\mathbf{x}^*}{d\mathbf{b}} & \frac{d\mathbf{x}^*}{dQ} & \frac{d\mathbf{x}^*}{dA} \\ \frac{d\lambda^*}{d\mathbf{c}} & \frac{d\lambda^*}{d\mathbf{b}} & \frac{d\lambda^*}{dQ} & \frac{d\lambda^*}{dA} \end{bmatrix} = \begin{bmatrix} -I & 0 & -x^* & -\lambda^* \\ 0 & I & 0 & -x^* \end{bmatrix} K \frac{\partial \ell}{\partial \mathbf{z}^*} \begin{bmatrix} \frac{d\mathbf{x}^*}{d\mathbf{c}} & \frac{d\mathbf{x}^*}{d\mathbf{b}} \\ \frac{d\lambda^*}{d\mathbf{c}} & \frac{d\lambda^*}{d\mathbf{b}} \end{bmatrix} = \begin{bmatrix} -\frac{\partial \ell}{\partial \mathbf{x}^*} \\ \mathbf{0} \end{bmatrix}$$

$$K d^* = \begin{bmatrix} -\frac{\partial \ell}{\partial \mathbf{x}^*} \\ \mathbf{0} \end{bmatrix}$$

- Equivalent QP:

$$d^* = \underset{d}{\operatorname{argmin}} \frac{1}{2} d^\top Q d + \frac{\partial \ell}{\partial \mathbf{x}^*}^\top d,$$

subject to $Ad = \mathbf{0}$.

Differentiable LQR

- Apply differentiation

$$\frac{d}{d\theta} \left(K \begin{bmatrix} \mathbf{x}^* \\ \lambda^* \end{bmatrix} \right) = \frac{dv}{d\theta}.$$

$$\frac{dK}{d\theta} \begin{bmatrix} \mathbf{x}^* \\ \lambda^* \end{bmatrix} + K \begin{bmatrix} \frac{d\mathbf{x}^*}{d\theta} \\ \frac{d\lambda^*}{d\theta} \end{bmatrix} = \frac{dv}{d\theta}$$

$$K \begin{bmatrix} \frac{d\mathbf{x}^*}{d\theta} \\ \frac{d\lambda^*}{d\theta} \end{bmatrix} = K \begin{bmatrix} \frac{d\mathbf{x}^*}{d\mathbf{c}} & \frac{d\mathbf{x}^*}{d\mathbf{b}} & \frac{d\mathbf{x}^*}{dQ} & \frac{d\mathbf{x}^*}{dA} \\ \frac{d\lambda^*}{d\mathbf{c}} & \frac{d\lambda^*}{d\mathbf{b}} & \frac{d\lambda^*}{dQ} & \frac{d\lambda^*}{dA} \end{bmatrix} = \begin{bmatrix} -I & 0 & -\mathbf{x}^* & -\lambda^* \\ 0 & I & 0 & -\mathbf{x}^* \end{bmatrix} K \frac{\partial \ell}{\partial \mathbf{z}^*} \begin{bmatrix} \frac{d\mathbf{x}^*}{d\mathbf{c}} \\ \frac{d\lambda^*}{d\mathbf{c}} \end{bmatrix} \begin{bmatrix} \frac{d\mathbf{x}^*}{d\mathbf{b}} \\ \frac{d\lambda^*}{d\mathbf{b}} \end{bmatrix} = \begin{bmatrix} -\frac{\partial \ell}{\partial \mathbf{x}^*} \\ \mathbf{0} \end{bmatrix}$$

$$K d^* = \begin{bmatrix} -\frac{\partial \ell}{\partial \mathbf{x}^*} \\ \mathbf{0} \end{bmatrix}$$

- Equivalent QP:

$$d^* = \underset{d}{\operatorname{argmin}} \frac{1}{2} d^\top Q d + \frac{\partial \ell}{\partial \mathbf{x}^*}^\top d,$$

subject to $Ad = \mathbf{0}$.

$$\frac{\partial \ell}{\partial Q} = \frac{1}{2} (d_x^* \otimes \mathbf{x}^* + \mathbf{x}^* \otimes d_x^*)$$

$$\frac{\partial \ell}{\partial A} = d_\lambda^* \otimes \mathbf{x}^* + \lambda^* \otimes d_x^*.$$

Differentiable LQR

- The backward pass can also be formulated as a LQR problem.
- Swap c to $\nabla_{\tau^*} \ell$ and f to 0.

Module 1 Differentiable LQR

(The LQR algorithm is defined in [appendix A](#))

Input: Initial state x_{init}

Parameters: $\theta = \{C, c, F, f\}$

Forward Pass:

- 1: $\tau_{1:T}^* = \text{LQR}_T(x_{\text{init}}; C, c, F, f)$ ▷ Solve (2)
- 2: Compute $\lambda_{1:T}^*$ with (7)

Backward Pass:

- 1: $d_{\tau_{1:T}}^* = \text{LQR}_T(0; C, \nabla_{\tau^*} \ell, F, 0)$ ▷ Solve (9), ideally reusing the factorizations from the forward pass
- 2: Compute $d_{\lambda_{1:T}}^*$ with (7)
- 3: Compute the derivatives of ℓ with respect to C, c, F, f , and x_{init} with (8)

Differentiable MPC

- What about general MPC?

$$\begin{aligned} & \underset{x_{1:T} \in \mathcal{X}, u_{1:T} \in \mathcal{U}}{\operatorname{argmin}} \sum_{t=1}^T C_t(x_t, u_t) \\ & \text{subject to } x_{t+1} = f(x_t, u_t), x_1 = x_{\text{init}}. \end{aligned}$$

Differentiable MPC

- What about general MPC?

$$\operatorname{argmin}_{x_{1:T} \in \mathcal{X}, u_{1:T} \in \mathcal{U}} \sum_{t=1}^T C_t(x_t, u_t)$$

subject to $x_{t+1} = f(x_t, u_t), x_1 = x_{\text{init}}$.

- Use Taylor expansion to approximate.

$$\tilde{C}_{\theta,t}^i = C_{\theta,t}(\tau_t^i) + p_t^{i\top} (\tau_t - \tau_t^i) + \frac{1}{2} (\tau_t - \tau_t^i)^\top H_t^i (\tau_t - \tau_t^i).$$

Differentiable MPC

- What about general MPC?

$$\begin{aligned} & \underset{x_{1:T} \in \mathcal{X}, u_{1:T} \in \mathcal{U}}{\operatorname{argmin}} \sum_{t=1}^T C_t(x_t, u_t) \\ & \text{subject to } x_{t+1} = f(x_t, u_t), x_1 = x_{\text{init}}. \end{aligned}$$

- Use Taylor expansion to approximate.

$$\tilde{C}_{\theta,t}^i = C_{\theta,t}(\tau_t^i) + p_t^{i\top} (\tau_t - \tau_t^i) + \frac{1}{2} (\tau_t - \tau_t^i)^\top H_t^i (\tau_t - \tau_t^i).$$

- Fixed point iteration.

$$\tau^{i+1} = \underset{\tau}{\operatorname{argmin}} \sum_t^T \tilde{C}_t^i(\tau_t^i).$$

Differentiable MPC

- What about general MPC?

$$\begin{aligned} & \underset{x_{1:T} \in \mathcal{X}, u_{1:T} \in \mathcal{U}}{\operatorname{argmin}} \sum_{t=1}^T C_t(x_t, u_t) \\ & \text{subject to } x_{t+1} = f(x_t, u_t), x_1 = x_{\text{init}}. \end{aligned}$$

- Use Taylor expansion to approximate.

$$\tilde{C}_{\theta,t}^i = C_{\theta,t}(\tau_t^i) + p_t^{i\top} (\tau_t - \tau_t^i) + \frac{1}{2} (\tau_t - \tau_t^i)^\top H_t^i (\tau_t - \tau_t^i).$$

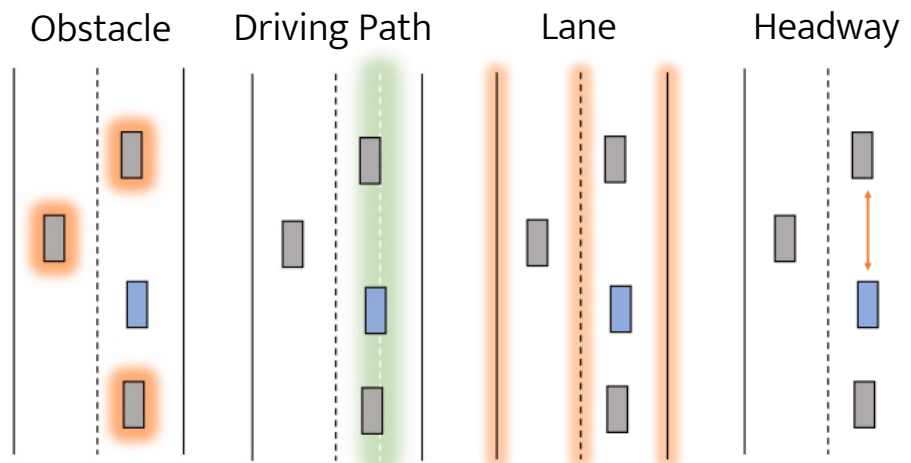
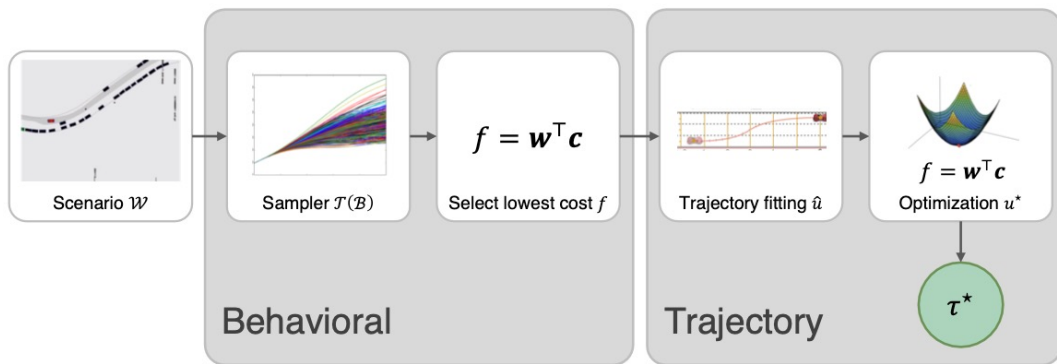
- Fixed point iteration.

$$\tau^{i+1} = \underset{\tau}{\operatorname{argmin}} \sum_t^T \tilde{C}_t^i(\tau_t^i).$$

- Backward only depends on final quadratic approximation.

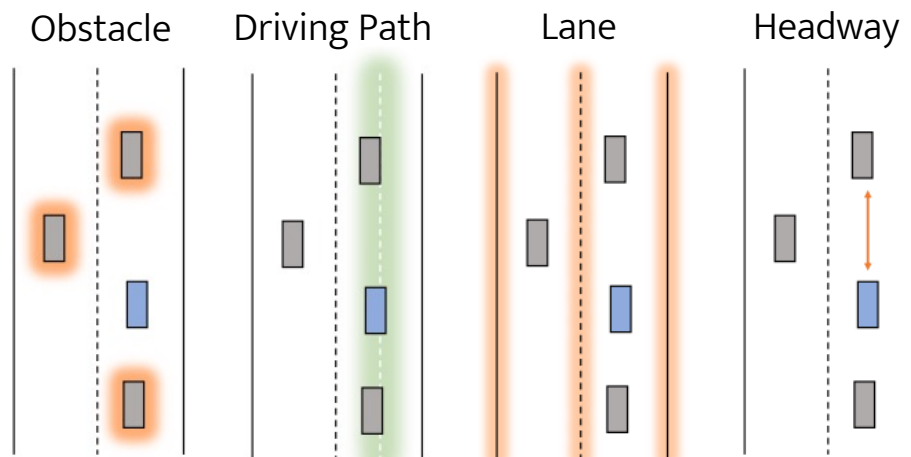
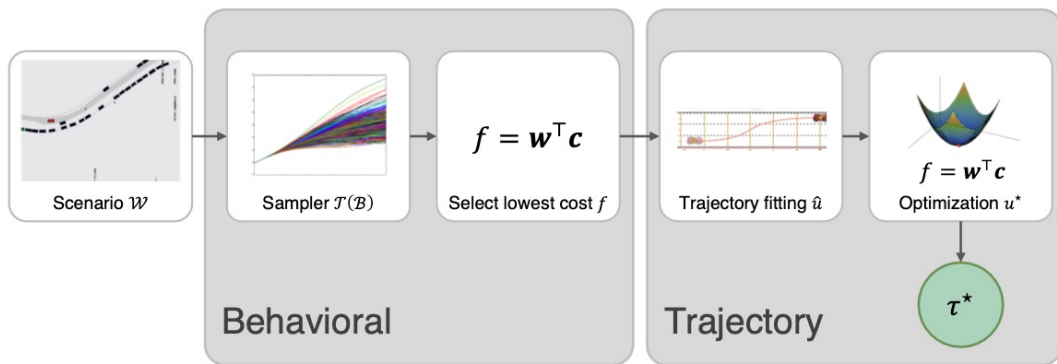
Behavioral vs. Trajectory Planning

- Gradient-based optimization provides a locally optimized trajectory.



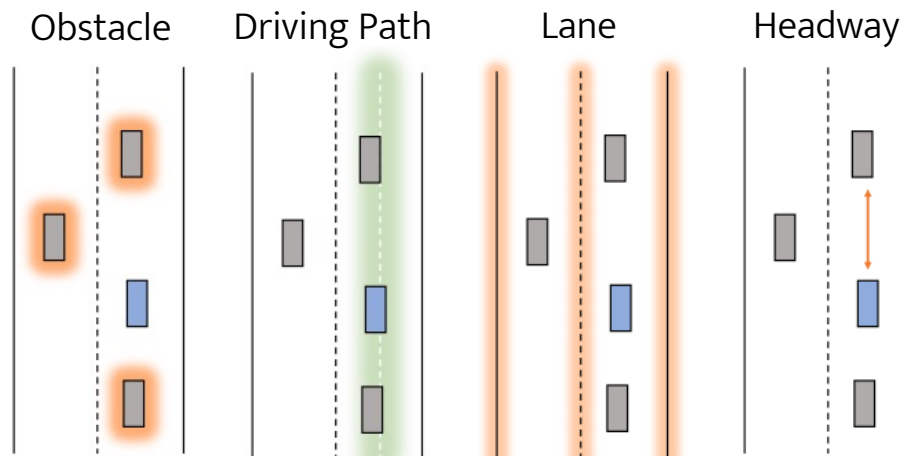
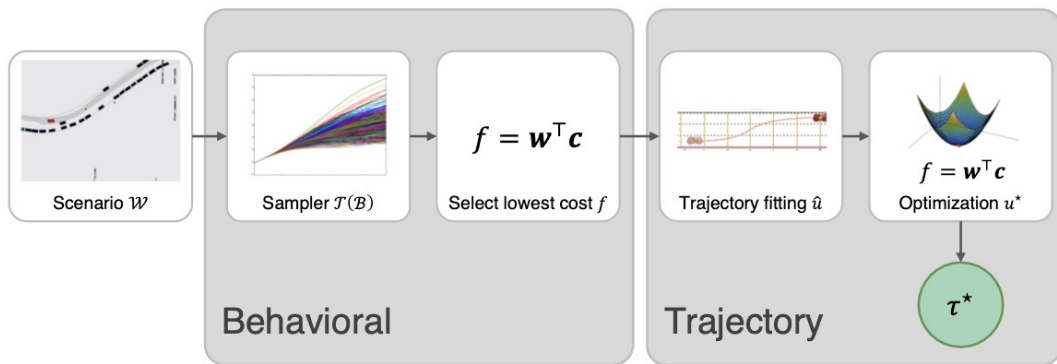
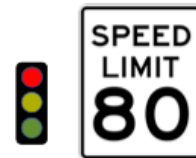
Behavioral vs. Trajectory Planning

- Gradient-based optimization provides a locally optimized trajectory.
- Samples may be needed for reasoning global structure.



Behavioral vs. Trajectory Planning

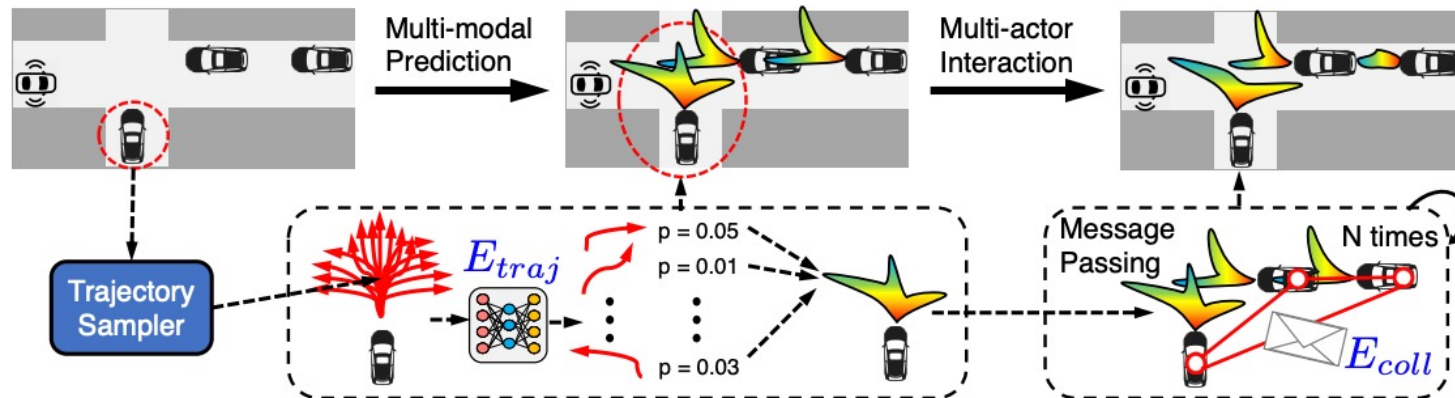
- Gradient-based optimization provides a locally optimized trajectory.
- Samples may be needed for reasoning global structure.
- Can learn together using the same learned costs.



Planning with Social Reasoning

- Jointly reason the future trajectories of multiple agents as an energy-based graphical model.

$$p(\mathbf{s}_1, \dots, \mathbf{s}_N \mid \mathbf{X}) = \frac{1}{Z} \exp(-E_\theta(\mathbf{s}_1, \dots, \mathbf{s}_N) \mid \mathbf{X})$$



Planning with Social Reasoning

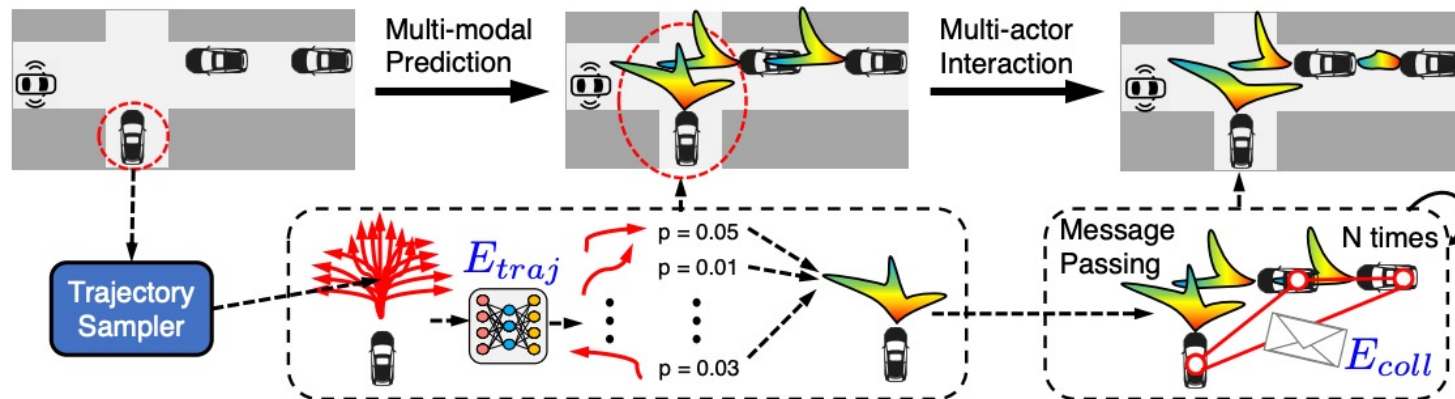
- Jointly reason the future trajectories of multiple agents as an energy-based graphical model.

$$p(\mathbf{s}_1, \dots, \mathbf{s}_N \mid \mathbf{X}) = \frac{1}{Z} \exp(-E_\theta(\mathbf{s}_1, \dots, \mathbf{s}_N) \mid \mathbf{X})$$

- Trajectory Goodness + Collision.

$$\sum_i E_\theta(s_i \mid X) + \sum_{i \neq j} E(s_i, s_j)$$

$$E(s_i, s_j) = \gamma \text{ if } s_i \text{ collides } s_j$$



Planning with Social Reasoning

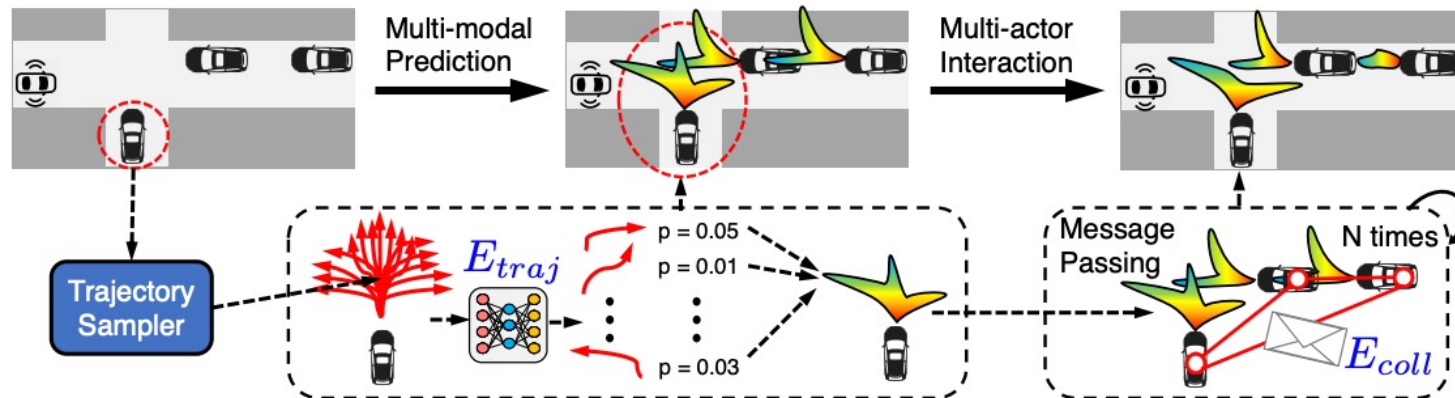
- Jointly reason the future trajectories of multiple agents as an energy-based graphical model.

$$p(\mathbf{s}_1, \dots, \mathbf{s}_N \mid \mathbf{X}) = \frac{1}{Z} \exp(-E_\theta(\mathbf{s}_1, \dots, \mathbf{s}_N) \mid \mathbf{X})$$

- Trajectory Goodness + Collision.
- Batch of trajectory samples.

$$\sum_i E_\theta(s_i \mid X) + \sum_{i \neq j} E(s_i, s_j)$$

$$E(s_i, s_j) = \gamma \text{ if } s_i \text{ collides } s_j$$



Planning with Social Reasoning

- Jointly reason the future trajectories of multiple agents as an energy-based graphical model.

$$p(\mathbf{s}_1, \dots, \mathbf{s}_N \mid \mathbf{X}) = \frac{1}{Z} \exp(-E_\theta(\mathbf{s}_1, \dots, \mathbf{s}_N) \mid \mathbf{X})$$

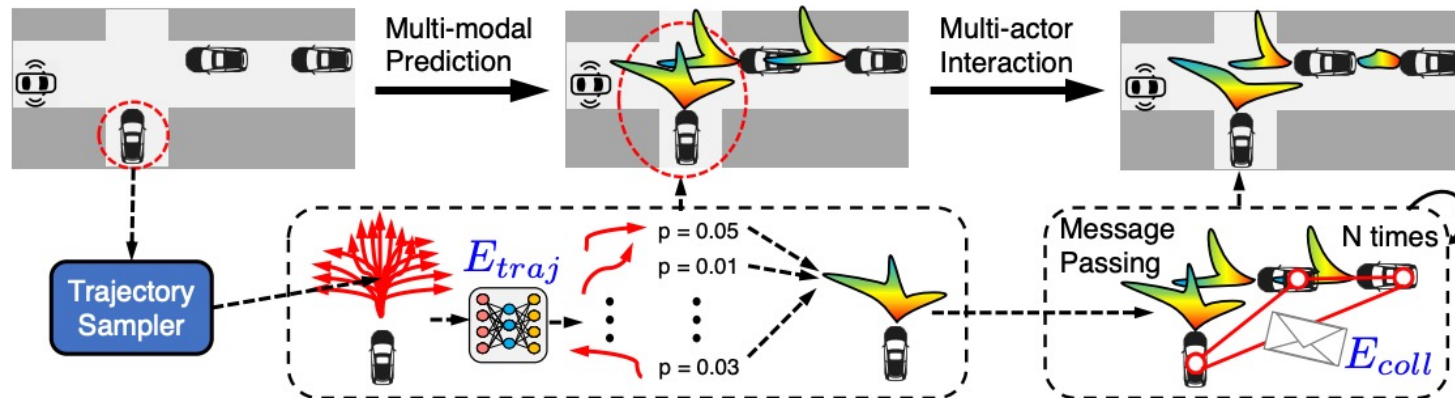
- Trajectory Goodness + Collision.

$$\sum_i E_\theta(s_i \mid X) + \sum_{i \neq j} E(s_i, s_j)$$

- Batch of trajectory samples.

$$E(s_i, s_j) = \gamma \text{ if } s_i \text{ collides } s_j$$

- Classification of groundtruth trajectory.



Summary: End-to-End Planning

- Direct Policy Prediction
 - Condition perception features into the model
 - Use of diffusion models

Summary: End-to-End Planning

- Direct Policy Prediction
 - Condition perception features into the model
 - Use of diffusion models
- Cost Learning (IRL) from Experts
 - Max-margin, max-entropy/EBM
 - Need negative samples
 - Can be combined with efficient external samplers
 - Cost volume prediction: parametric + non-parametric

Summary: End-to-End Planning

- Direct Policy Prediction
 - Condition perception features into the model
 - Use of diffusion models
- Cost Learning (IRL) from Experts
 - Max-margin, max-entropy/EBM
 - Need negative samples
 - Can be combined with efficient external samplers
 - Cost volume prediction: parametric + non-parametric
- Differentiable Planner
 - Backprop through local optimization
 - Can be memory efficient, implicit differentiation