

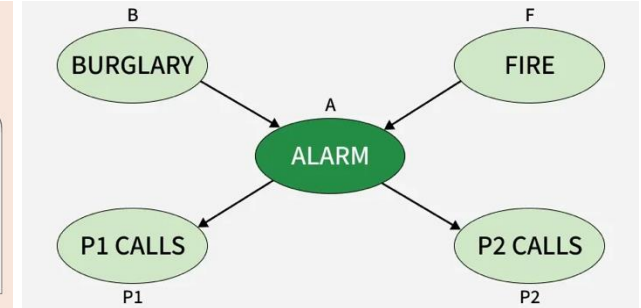
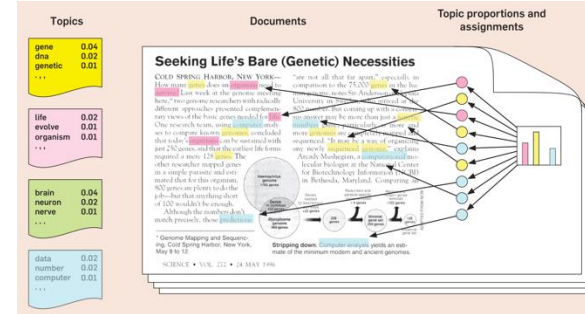
# Rethinking Probabilistic Circuit Parameter Learning

Anji Liu, Zilei Zoe Shao, Guy Van den Broeck  
AISTATS 2026

# Probabilistic Methods in the Era of LLMs

## (Probabilistic) Machine Learning in 2000

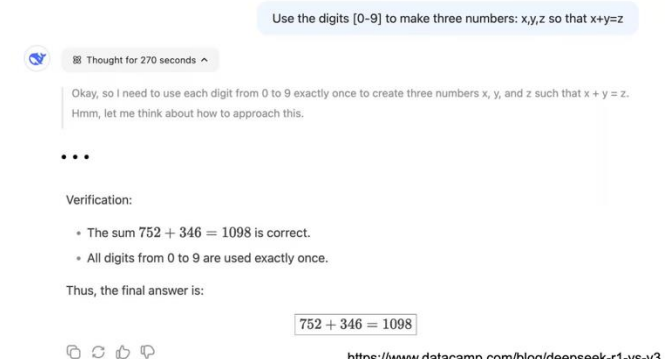
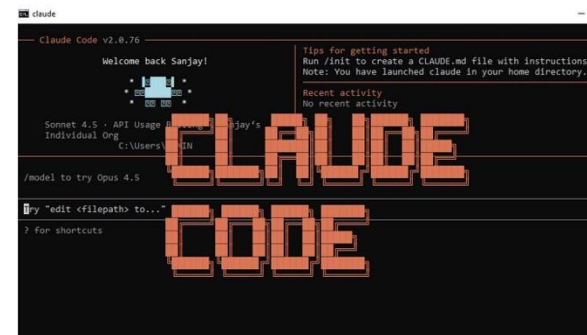
- Dominance of Probabilistic Graphical Models
- Tractable inference
- Incorporation of domain knowledge/constraints



 26 years

## (Probabilistic) Machine Learning in 2026

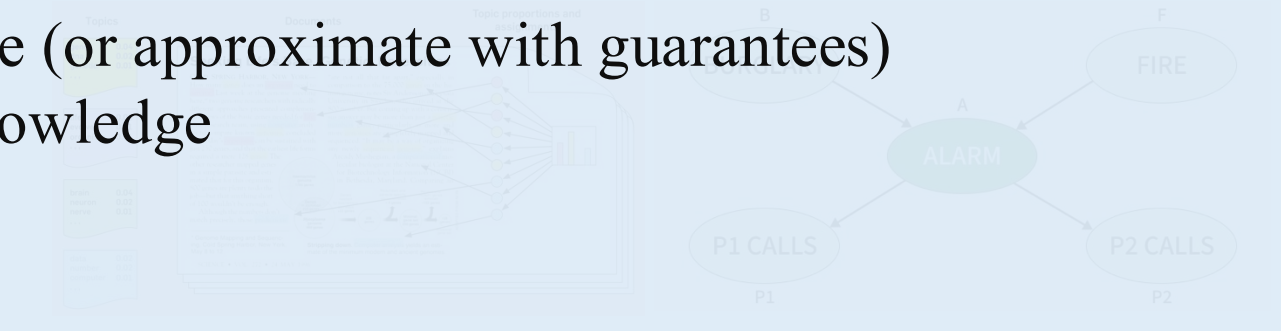
- Better model architectures
- Massive data
- Improved prompting techniques



<https://www.datacamp.com/blog/deepseek-r1-vs-v3>

# Probabilistic Methods in the Era of LLMs

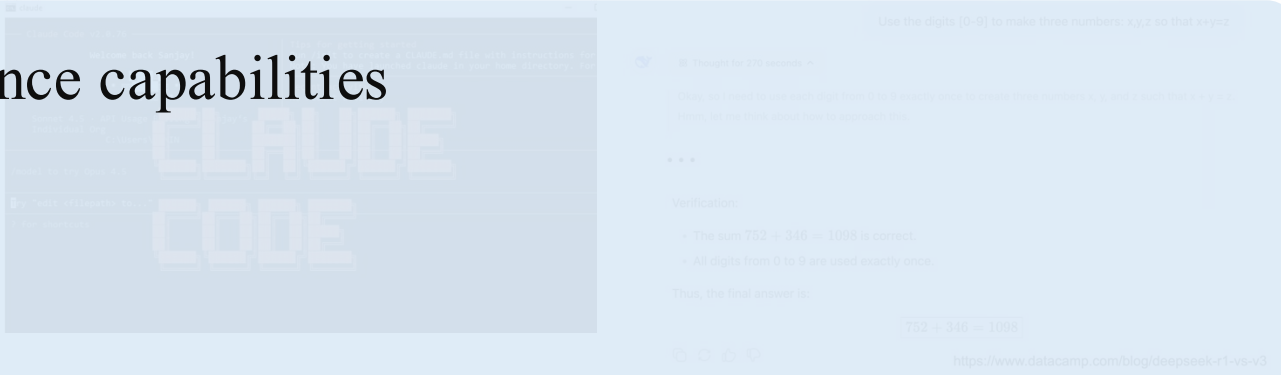
## (Probabilistic) Machine Learning in 2000

- Exact probabilistic inference (or approximate with guarantees)
  - Incorporation of domain knowledge
- 
- The collage includes a neural network diagram with layers labeled 'Input', 'Hidden', and 'Output'. It also features a bar chart with three bars of increasing height, and a flowchart with nodes labeled 'P1 CALLS', 'P2 CALLS', 'ALARM', and 'FIRE'.



26 years

## (Probabilistic) Machine Learning in 2026

- Limited probabilistic inference capabilities
  - Can only learn from data
- 
- The collage includes a screenshot of the 'CLAUDE CODE' interface showing a code editor. It also features a tweet from 'AI' discussing a math problem: 'Use the digits [0-9] to make three numbers: x,y,z so that x+y=z. How, do I need to use each digit from 0 to 9 exactly once to create three numbers x, y, and z such that x + y = z. Hint: let me think about how to approach this.' The tweet includes a verification section stating 'The sum 752 + 346 = 1098 is correct' and 'All digits from 0 to 9 are used exactly once.' The final answer is given as '752 + 346 = 1098'.

# Probabilistic Methods in the Era of LLMs

## (Probabilistic) Machine Learning in 2000

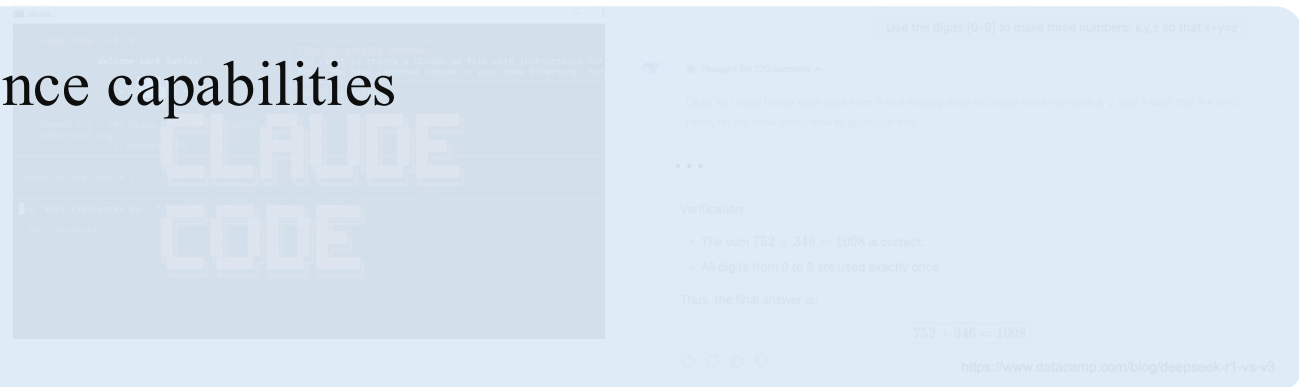
- Dominance of exact probabilistic inference (or approximate with guarantees)
  - Tractable inference
  - Incorporation of domain knowledge/constraints
- ... *but* ...
- Hard to scale



26 years

## (Probabilistic) Machine Learning in 2026

- Better model architectures
  - Massive data
  - Improved prompting techniques
- ... *but* ...
- Extremely scalable



# Probabilistic Graphical Models Are Hard to Scale

Structure learning is provably hard

**Large-Sample Learning of Bayesian Networks is NP-Hard**

**David Maxwell Chickering**

**David Heckerman**

**Christopher Meek**

*Microsoft Research*

*Redmond, WA 98052, USA*

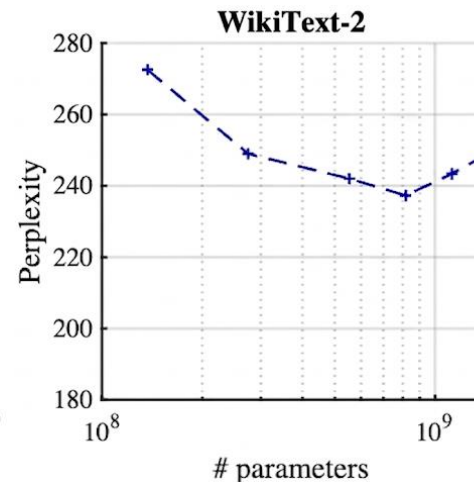
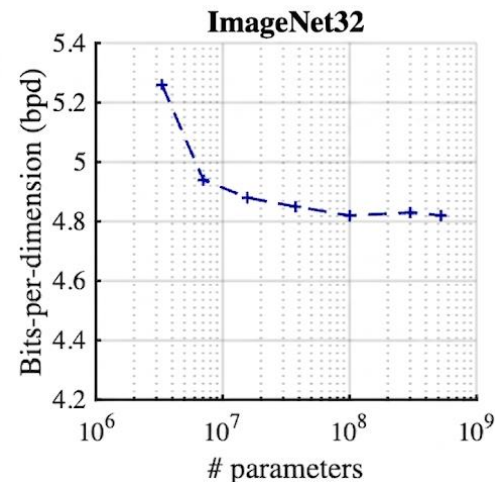
DMAX@MICROSOFT.COM

HECKERMA@MICROSOFT.COM

MEEK@MICROSOFT.COM

**Parameter learning given a fixed structure is empirically hard**

Log-likelihood saturates as the parameter count goes up (Liu et al.)

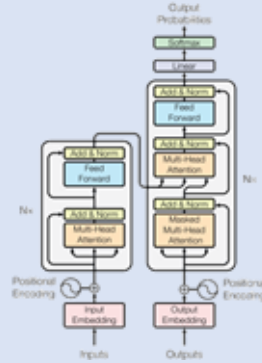


# A Pathway Towards Scaling Up Probabilistic Graphical Models

## Neural Networks

### Scalable Representation

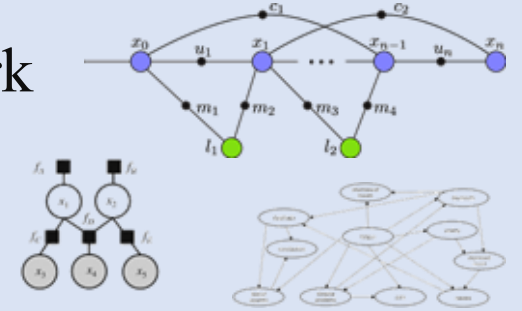
- Transformer
- Mamba
- Etc.



## Probabilistic Graphical Models

### Scalable Representation

- Bayesian Network
- Factor Graph
- Etc.

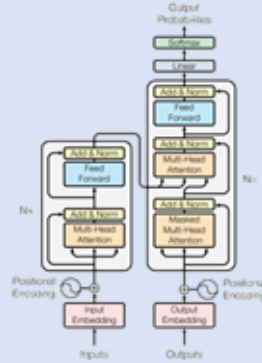


# A Pathway Towards Scaling Up Probabilistic Graphical Models

## Neural Networks

### Scalable Representation

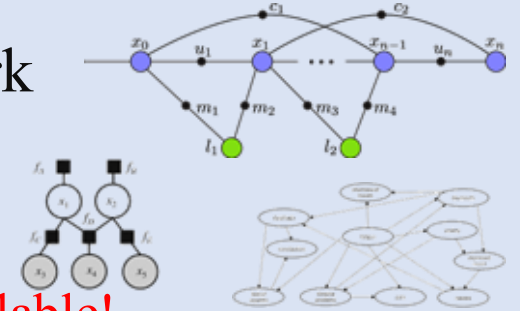
- Transformer
- Mamba
- Etc.



## Probabilistic Graphical Models

### Scalable Representation

- Bayesian Network
- Factor Graph
- Etc.



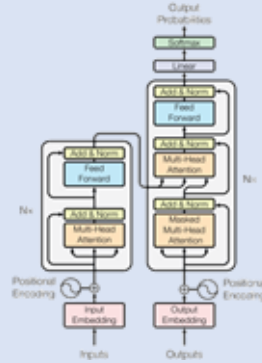
Not computationally scalable!

# A Pathway Towards Scaling Up Probabilistic Graphical Models

## Neural Networks

### Scalable Representation

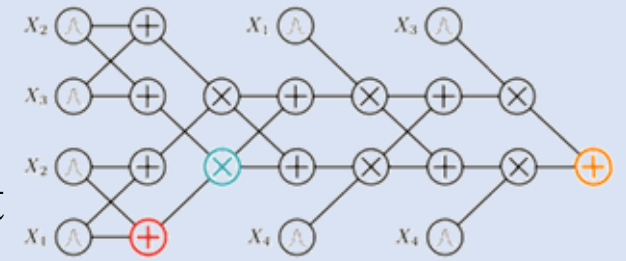
- Transformer
- Mamba
- Etc.



## Probabilistic Graphical Models

### Scalable Representation

- Probabilistic Circuit (or Sum-Product Network)



# Tractable Deep Generative Models: Probabilistic Circuits

*Encode complex distributions by starting with simple distributions and recursively combining them to represent more complex ones.*

# Tractable Deep Generative Models: Probabilistic Circuits

Dataset

$(x = \text{red square}, y = \text{green square})$

$(x = \text{blue square}, y = \text{yellow square})$

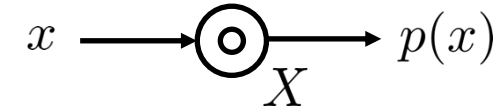
# Tractable Deep Generative Models: Probabilistic Circuits

Dataset

$(x = \text{red}, y = \text{green})$

$(x = \text{blue}, y = \text{yellow})$

- **Input nodes** encode univariate primitive distributions.



$$p(x) = \begin{cases} 1 & x = \text{red}, \\ 0 & \text{otherwise.} \end{cases}$$



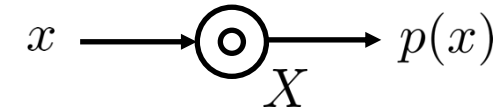
# Tractable Deep Generative Models: Probabilistic Circuits

Dataset

$(x = \text{red}, y = \text{green})$

$(x = \text{blue}, y = \text{yellow})$

- **Input nodes** encode univariate primitive distributions.



$$p(x) = \begin{cases} 1 & x = \text{red}, \\ 0 & \text{otherwise.} \end{cases} \quad p(y) = \begin{cases} 1 & y = \text{green}, \\ 0 & \text{otherwise.} \end{cases}$$

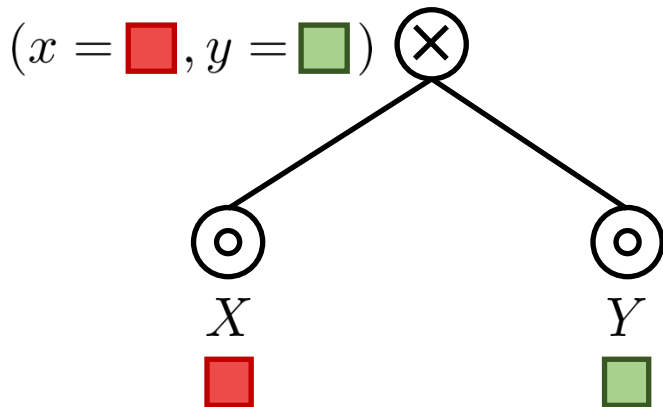


# Tractable Deep Generative Models: Probabilistic Circuits

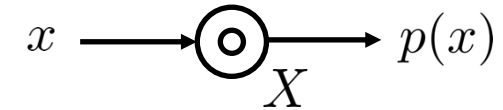
Dataset

$(x = \text{red}, y = \text{green})$

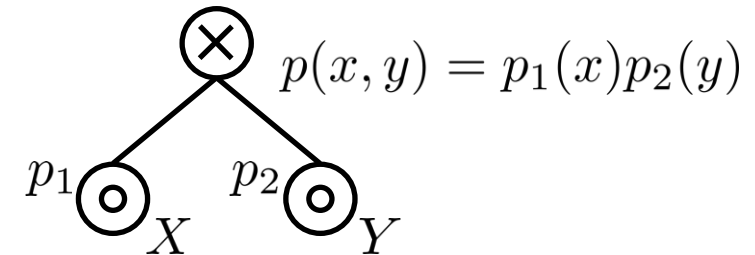
$(x = \text{blue}, y = \text{yellow})$



- **Input nodes** encode univariate primitive distributions.



- **Product nodes** build factorized distributions over their children.

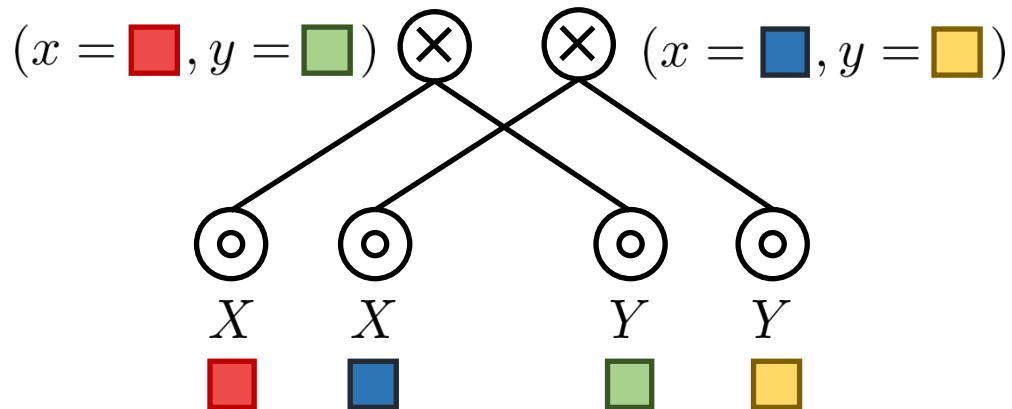


# Tractable Deep Generative Models: Probabilistic Circuits

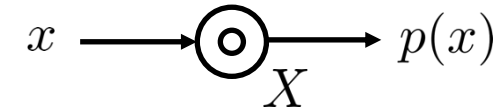
Dataset

$(x = \text{red}, y = \text{green})$

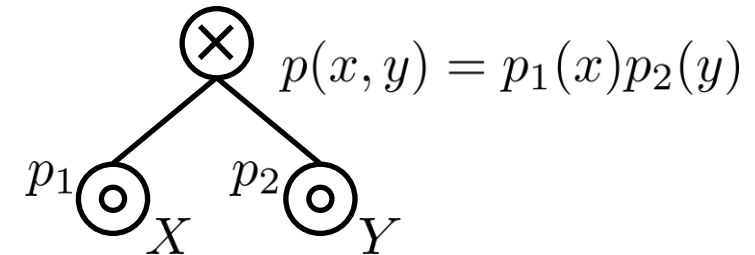
$(x = \text{blue}, y = \text{yellow})$



- **Input nodes** encode univariate primitive distributions.



- **Product nodes** build factorized distributions over their children.



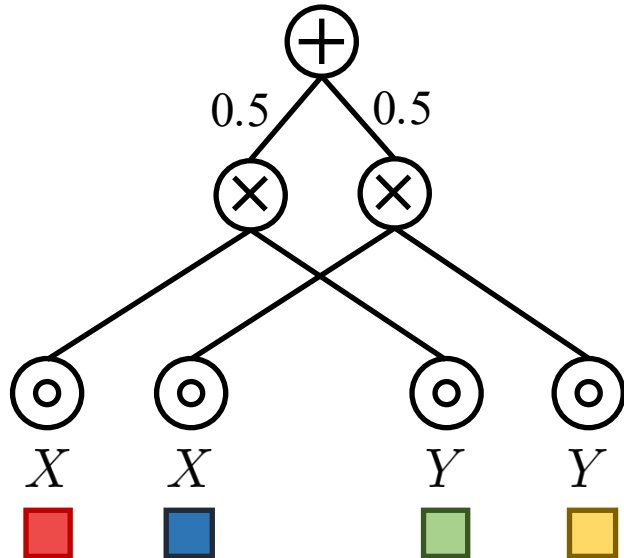
# Tractable Deep Generative Models: Probabilistic Circuits

Dataset

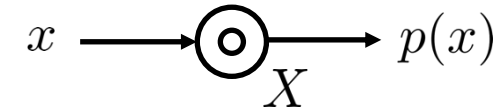
$(x = \text{red}, y = \text{green})$

$(x = \text{blue}, y = \text{yellow})$

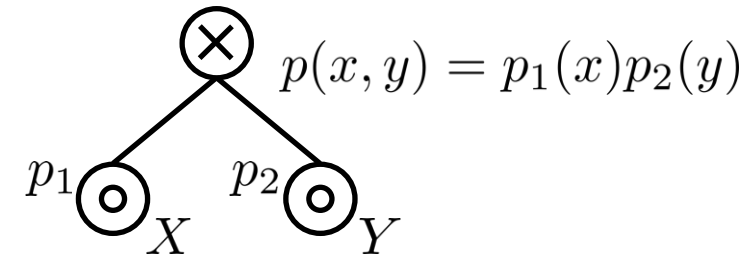
$$0.5 \cdot (x = \text{red}, y = \text{green}) + 0.5 \cdot (x = \text{blue}, y = \text{yellow})$$



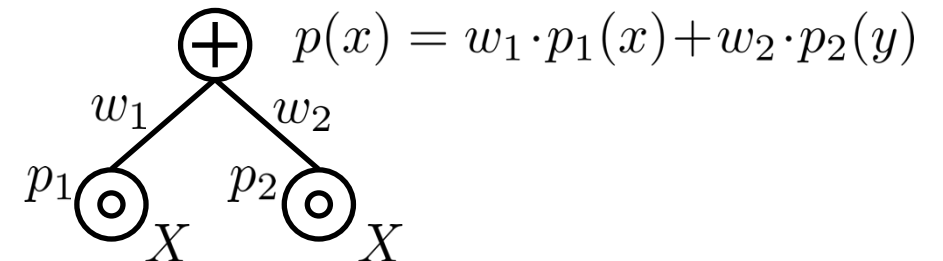
- **Input nodes** encode univariate primitive distributions.



- **Product nodes** build factorized distributions over their children.

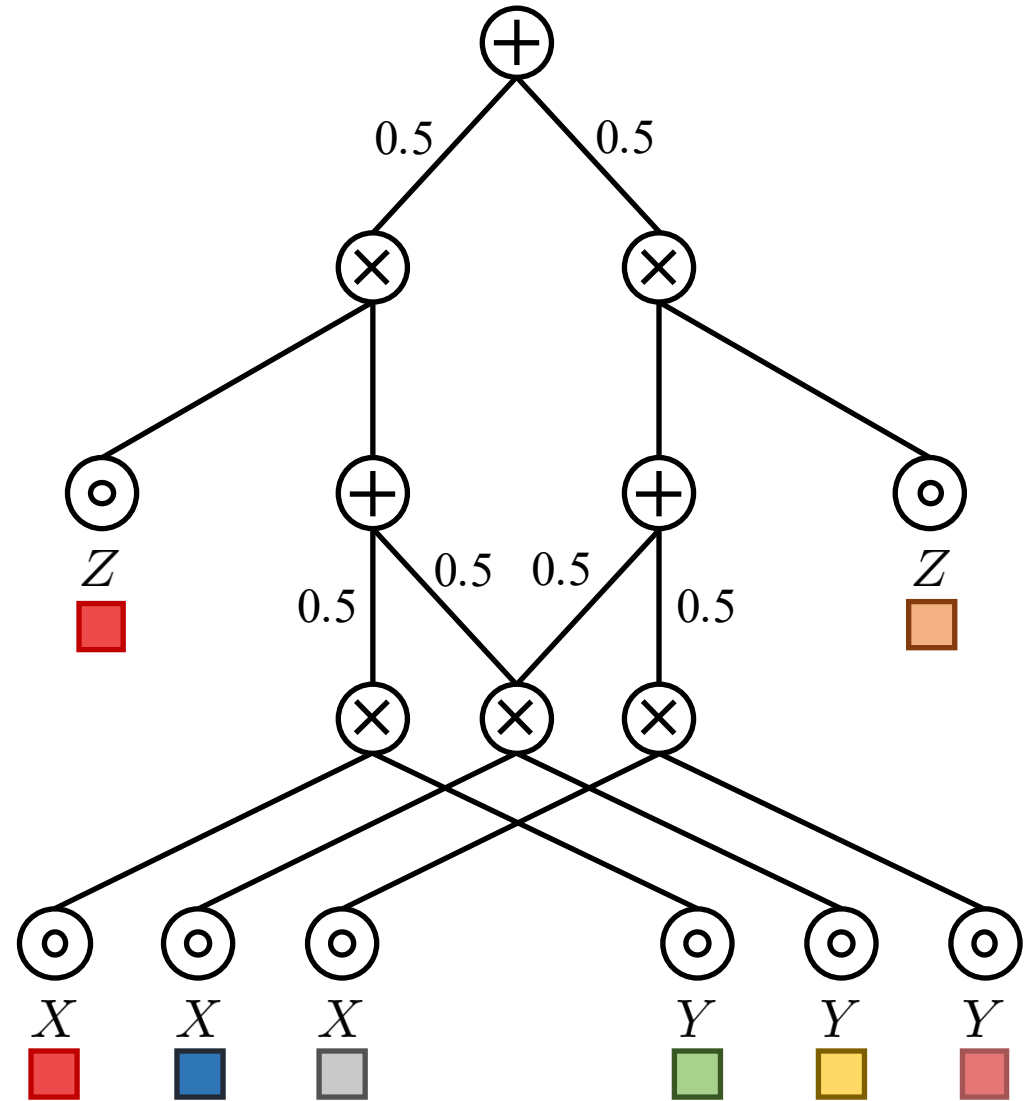


- **Sum nodes** construct weighted mixtures over their child distributions.



# Compute Likelihood

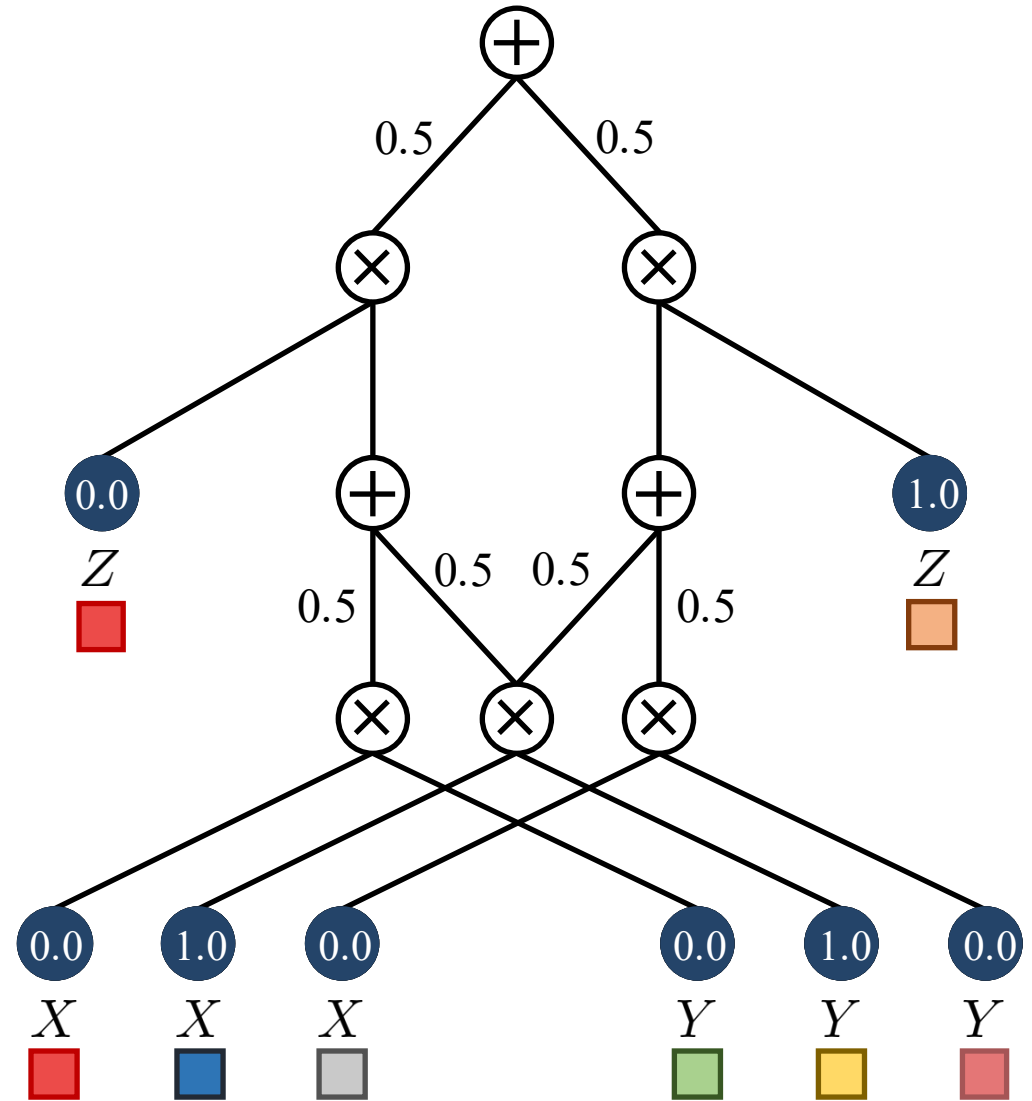
Compute  $p(x = \blacksquare, y = \blacksquare, z = \blacksquare)$



# Compute Likelihood

Compute  $p(x = \blacksquare, y = \blacksquare, z = \blacksquare)$

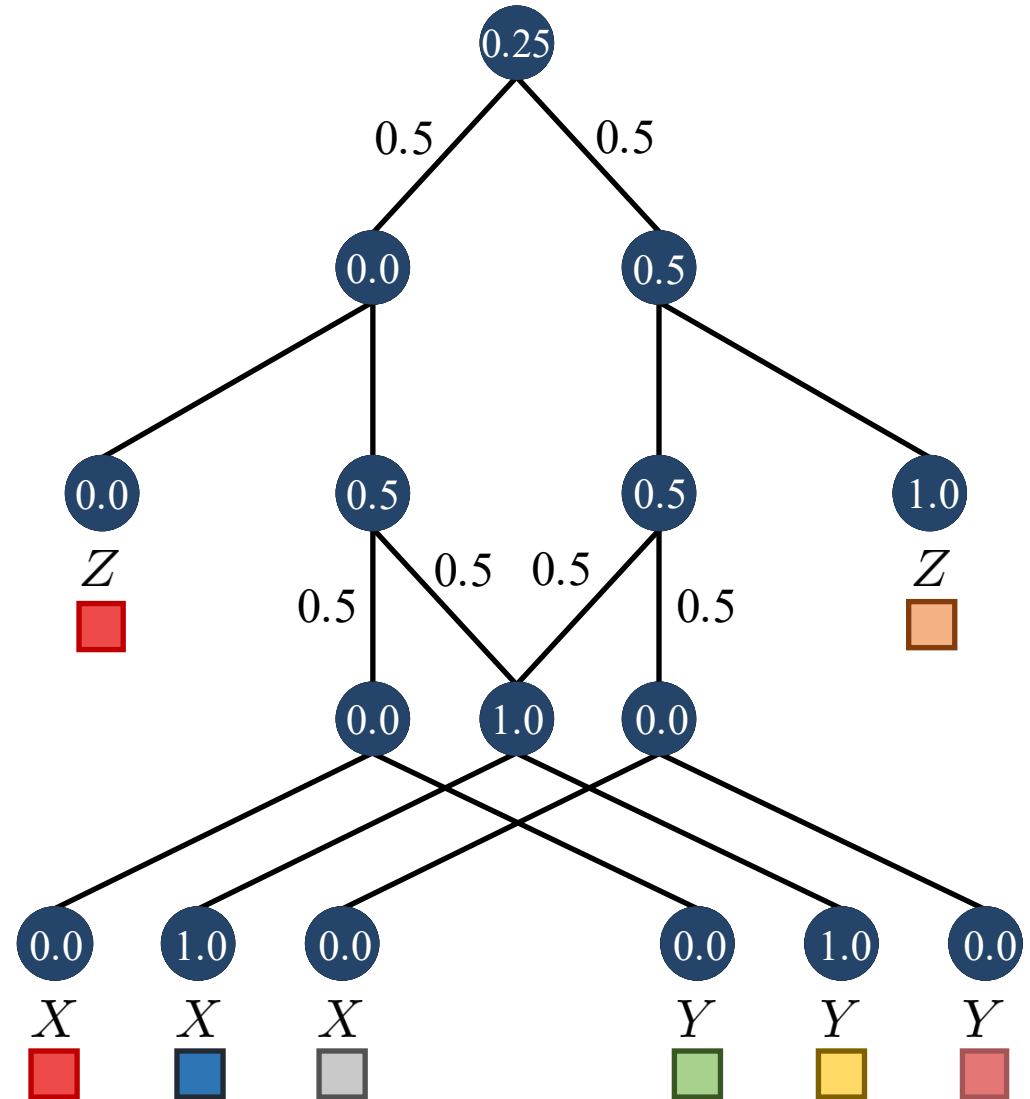
- Compute the likelihood of every **input node**.



# Compute Likelihood

Compute  $p(x = \blacksquare, y = \blacksquare, z = \blacksquare) = 0.25$

- Compute the likelihood of every **input node**.
- Compute the likelihood of every **sum/product node**.
- Readout likelihood from the **root node**.



# Computing Marginal

Compute  $p(x = \blacksquare) = \iint p(x = \blacksquare, y, z) dydz$

- Sum node  $\oplus_a$**

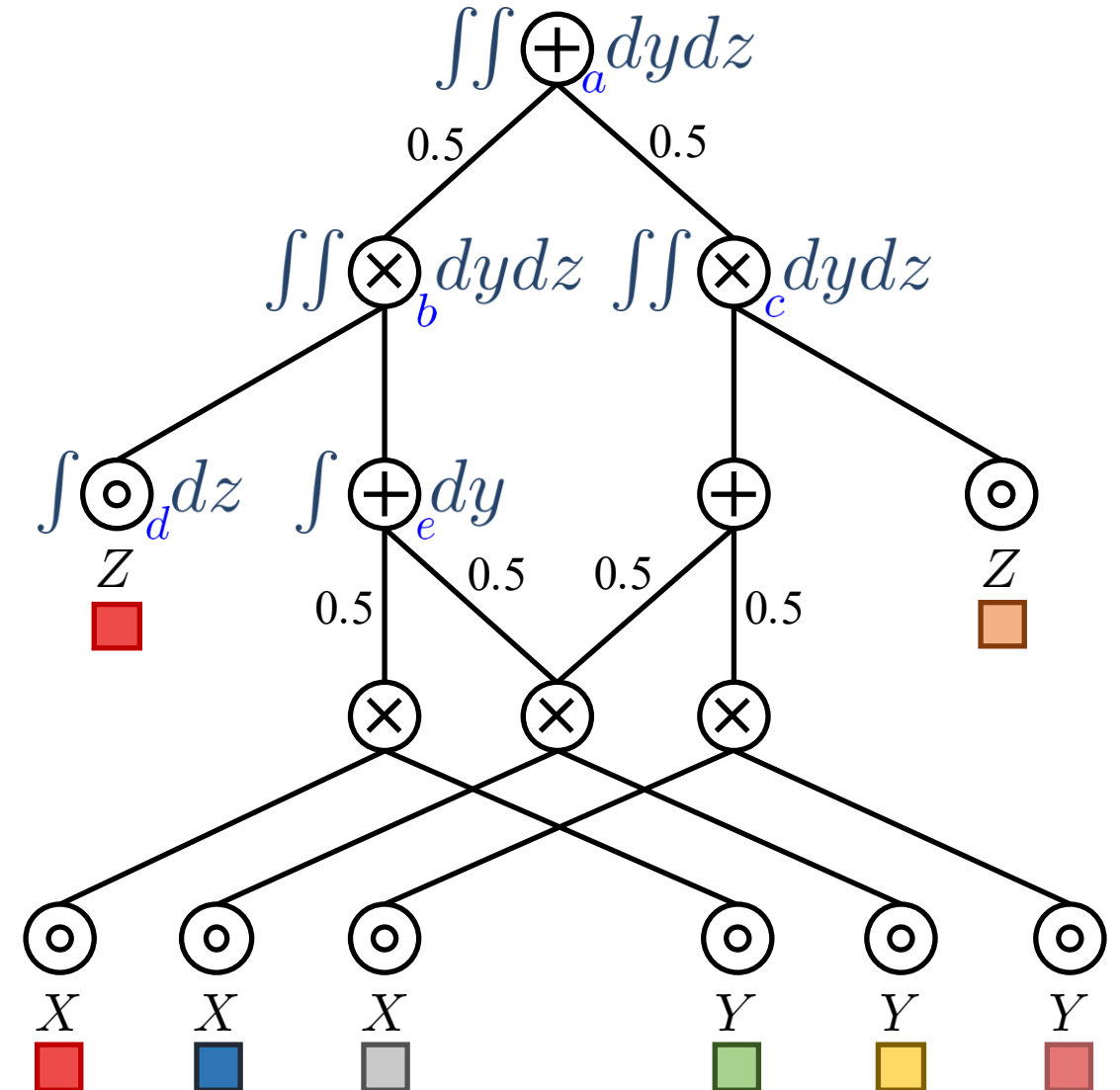
$$\begin{aligned} & \iint p_a(x = \blacksquare, y, z) dydz \\ &= \iint 0.5 \cdot p_b(x = \blacksquare, y, z) + 0.5 \cdot p_c(x = \blacksquare, y, z) dydz \\ &= 0.5 \underbrace{\iint p_b(x = \blacksquare, y, z) dydz}_{\iint \otimes_b dydz} + 0.5 \underbrace{\iint p_c(x = \blacksquare, y, z) dydz}_{\iint \otimes_c dydz} \end{aligned}$$

- Product node  $\otimes_b$**

$$\begin{aligned} & \iint p_b(x = \blacksquare, y, z) dydz \\ &= \iint p_d(z) \cdot p_e(x = \blacksquare, y) dydz \\ &= \underbrace{\int p_d(z) dz}_{\int \otimes_d dz} \cdot \underbrace{\int p_e(x = \blacksquare, y) dy}_{\int \otimes_e dy} \end{aligned}$$

- Input node  $\odot_d$**

$$\int p_d(z) = 1$$

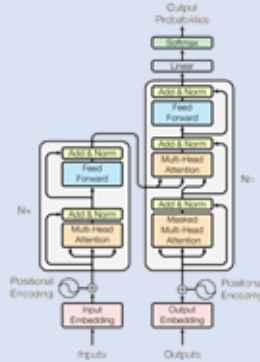


# A Pathway Towards Scaling Up Probabilistic Graphical Models

## Neural Networks

### Scalable Representation

- Transformer
- Mamba
- Etc.



### Efficient Systems

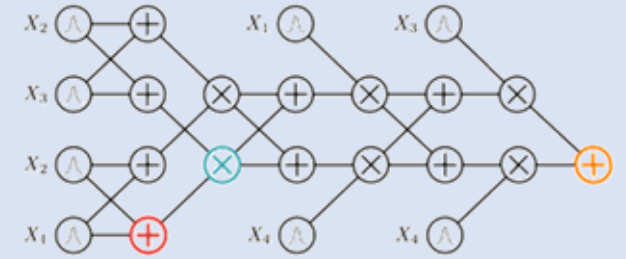
 PyTorch



## Probabilistic Graphical Models

### Scalable Representation

- Probabilistic Circuit



### Efficient Systems



 circuit

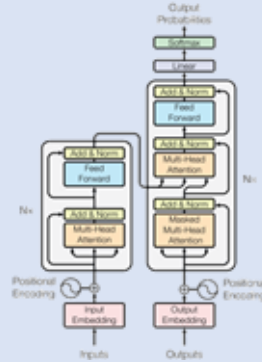
 FlowW

# A Pathway Towards Scaling Up Probabilistic Graphical Models

## Neural Networks

### Scalable Representation

- Transformer
- Mamba
- Etc.



### Efficient Systems

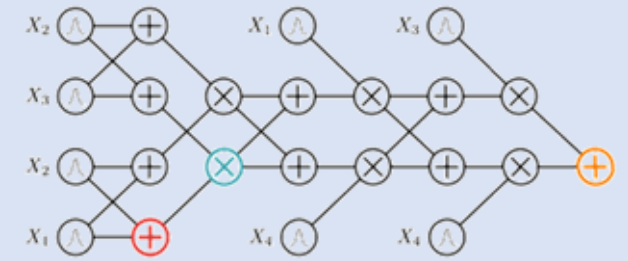
 PyTorch



## Probabilistic Graphical Models

### Scalable Representation

- Probabilistic Circuit



### Efficient Systems

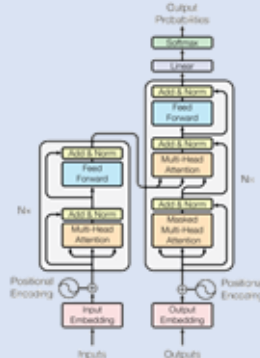
**PyJuice**  - 100x speedup vs. Google baselines.  
- Scaled to **billions of parameters**.

# A Pathway Towards Scaling Up Probabilistic Graphical Models

## Neural Networks

### Scalable Representation

- Transformer
- Mamba
- Etc.



### Efficient Systems

 PyTorch



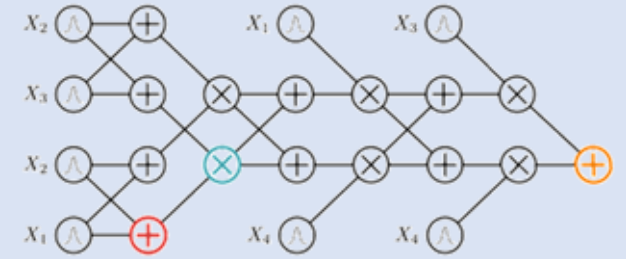
### Strong Optimizers

Adam, Muon, etc.

## Probabilistic Graphical Models

### Scalable Representation

- Probabilistic Circuit



### Efficient Systems

PyJuice  - 100x speedup vs. Google baselines.  
- Scaled to **billions of parameters**.

### Strong Optimizers

**This work.**

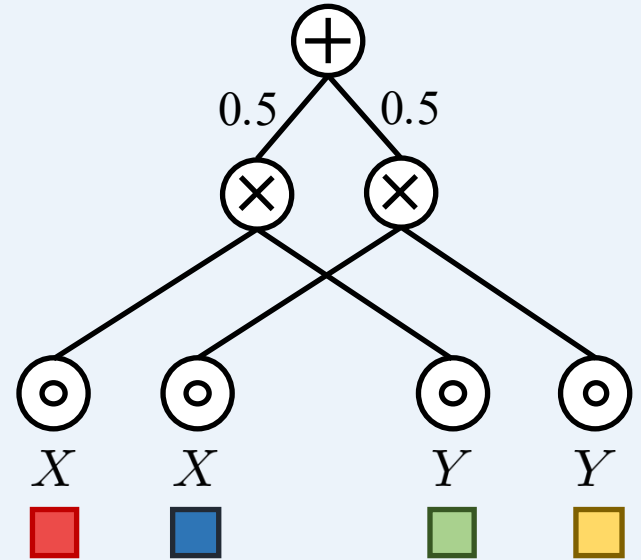
# Optimizing PCs with Expectation Maximization

PCs are latent variable models

$$p_{\phi}(\boldsymbol{x}) = \sum_{\boldsymbol{z}} p_{\phi}(\boldsymbol{x}, \boldsymbol{z})$$

**Example**

$$0.5 \cdot (x = \text{red}, y = \text{green}) + 0.5 \cdot (x = \text{blue}, y = \text{yellow})$$



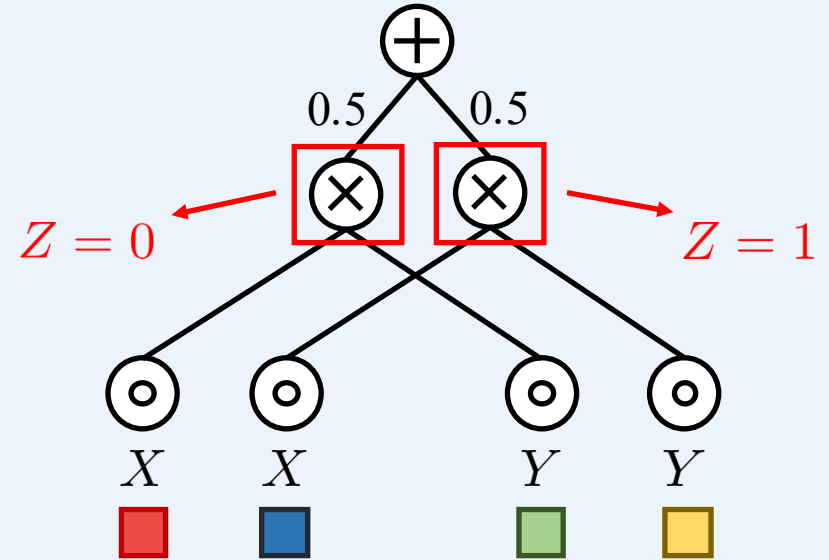
# Optimizing PCs with Expectation-Maximization

## PCs are latent variable models

$$p_{\phi}(\mathbf{x}) = \sum_{\mathbf{z}} p_{\phi}(\mathbf{x}, \mathbf{z})$$

## Example

$$0.5 \cdot (x = \text{red}, y = \text{green}) + 0.5 \cdot (x = \text{blue}, y = \text{yellow})$$



# Optimizing PCs with Expectation-Maximization

**PCs are latent variable models**

$$p_{\phi}(\mathbf{x}) = \sum_{\mathbf{z}} p_{\phi}(\mathbf{x}, \mathbf{z})$$

**Standard Expectation-Maximization**

Instead of optimizing the log-likelihood

$$\text{LL}(\phi) := \log p_{\phi}(\mathbf{x}) = \log \left( \sum_{\mathbf{z}} p_{\phi}(\mathbf{x}, \mathbf{z}) \right),$$

we maximize the following lower-bound

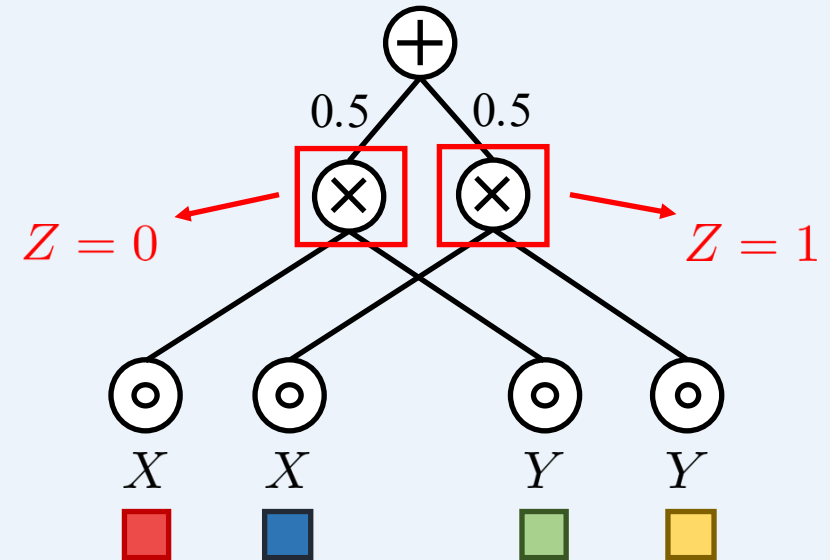
$$Q_{\phi}(\phi') := \sum_{\mathbf{z}} p_{\phi}(\mathbf{z}|\mathbf{x}) \cdot \log p_{\phi'}(\mathbf{x}, \mathbf{z}).$$

$\phi$  : the current parameters

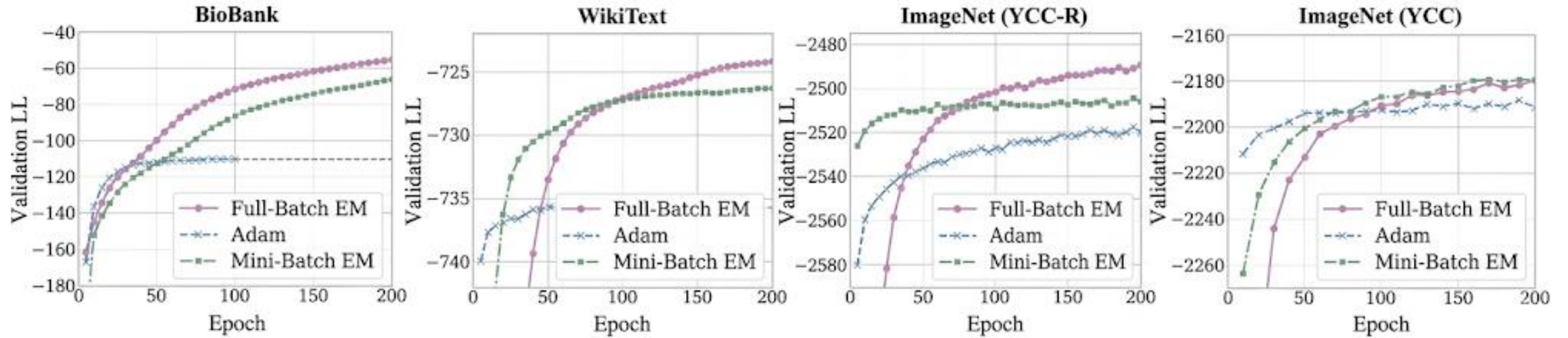
$\phi'$  : parameters to be optimized

**Example**

$$0.5 \cdot (x = \text{red}, y = \text{green}) + 0.5 \cdot (x = \text{blue}, y = \text{yellow})$$



# The Pros and Cons of Full-Batch Expectation-Maximization



## Pros:

- The best performance compared to baselines
- No need for hyperparameter tuning

## Cons

- Infrequent updates
- Hard to scale to large datasets
- Slow convergence

# The Existing Mini-Batch EM Algorithm

## Full-batch EM:

Given a dataset  $\mathcal{D}_{\text{full}}$ , solve for  $\phi'$  that maximizes

$$Q_{\phi}^{\mathcal{D}_{\text{full}}}(\phi') := \sum_{\mathbf{x} \in \mathcal{D}_{\text{full}}} \sum_{\mathbf{z}} p_{\phi}(\mathbf{z}|\mathbf{x}) \cdot \log p_{\phi'}(\mathbf{x}, \mathbf{z}).$$

# The Existing Mini-Batch EM Algorithm

## Full-batch EM:

Given a dataset  $\mathcal{D}_{\text{full}}$ , solve for  $\phi'$  that maximizes

$$Q_{\phi}^{\mathcal{D}_{\text{full}}}(\phi') := \sum_{\mathbf{x} \in \mathcal{D}_{\text{full}}} \sum_{\mathbf{z}} p_{\phi}(\mathbf{z}|\mathbf{x}) \cdot \log p_{\phi'}(\mathbf{x}, \mathbf{z}).$$

## Mini-batch EM:

Given a mini-batch of data  $\mathcal{D}_{\text{mini}}$ , solve for  $\phi'$  that maximizes

$$Q_{\phi}^{\mathcal{D}_{\text{mini}}}(\phi') := \sum_{\mathbf{x} \in \mathcal{D}_{\text{mini}}} \sum_{\mathbf{z}} p_{\phi}(\mathbf{z}|\mathbf{x}) \cdot \log p_{\phi'}(\mathbf{x}, \mathbf{z}).$$

# The Existing Mini-Batch EM Algorithm

## Full-batch EM:

Given a dataset  $\mathcal{D}_{\text{full}}$ , solve for  $\phi'$  that maximizes

$$Q_{\phi}^{\mathcal{D}_{\text{full}}}(\phi') := \sum_{\mathbf{x} \in \mathcal{D}_{\text{full}}} \sum_{\mathbf{z}} p_{\phi}(\mathbf{z}|\mathbf{x}) \cdot \log p_{\phi'}(\mathbf{x}, \mathbf{z}).$$

## Mini-batch EM:

Given a mini-batch of data  $\mathcal{D}_{\text{mini}}$ , solve for  $\phi'$  that maximizes

$$Q_{\phi}^{\mathcal{D}_{\text{mini}}}(\phi') := \sum_{\mathbf{x} \in \mathcal{D}_{\text{mini}}} \sum_{\mathbf{z}} p_{\phi}(\mathbf{z}|\mathbf{x}) \cdot \log p_{\phi'}(\mathbf{x}, \mathbf{z}).$$

However,  $\mathcal{D}_{\text{mini}}$  does not represent the entire dataset, so we should not “trust”  $\phi'$  entirely.

# The Existing Mini-Batch EM Algorithm

## Full-batch EM:

Given a dataset  $\mathcal{D}_{\text{full}}$ , solve for  $\phi'$  that maximizes

$$Q_{\phi}^{\mathcal{D}_{\text{full}}}(\phi') := \sum_{\mathbf{x} \in \mathcal{D}_{\text{full}}} \sum_{\mathbf{z}} p_{\phi}(\mathbf{z}|\mathbf{x}) \cdot \log p_{\phi'}(\mathbf{x}, \mathbf{z}).$$

## Mini-batch EM:

Given a mini-batch of data  $\mathcal{D}_{\text{mini}}$ , solve for  $\phi'$  that maximizes

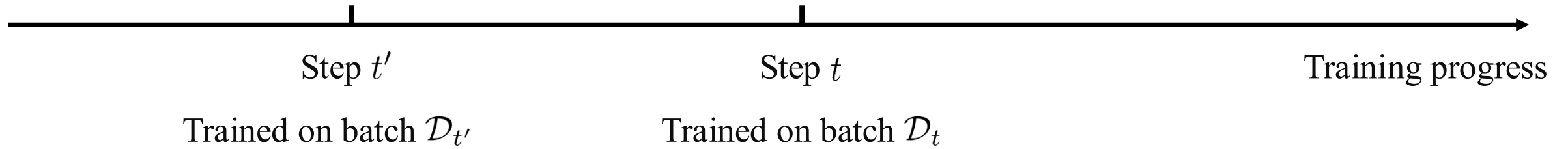
$$Q_{\phi}^{\mathcal{D}_{\text{mini}}}(\phi') := \sum_{\mathbf{x} \in \mathcal{D}_{\text{mini}}} \sum_{\mathbf{z}} p_{\phi}(\mathbf{z}|\mathbf{x}) \cdot \log p_{\phi'}(\mathbf{x}, \mathbf{z}).$$

However,  $\mathcal{D}_{\text{mini}}$  does not represent the entire dataset, so we should not “trust”  $\phi'$  entirely.

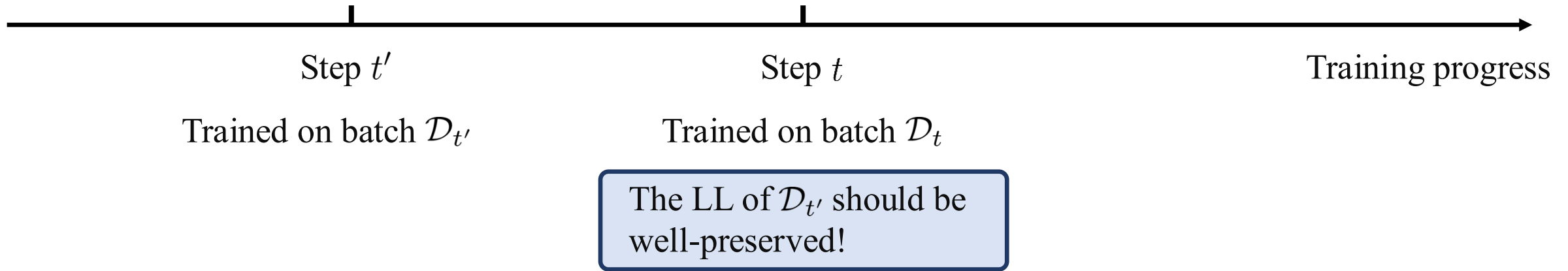
Instead, we update the parameters “incrementally”:

$$\phi_{\text{new}} := (1 - \alpha) \cdot \phi + \alpha \cdot \phi', \text{ where } \alpha \in (0, 1] \text{ is a learning rate.}$$

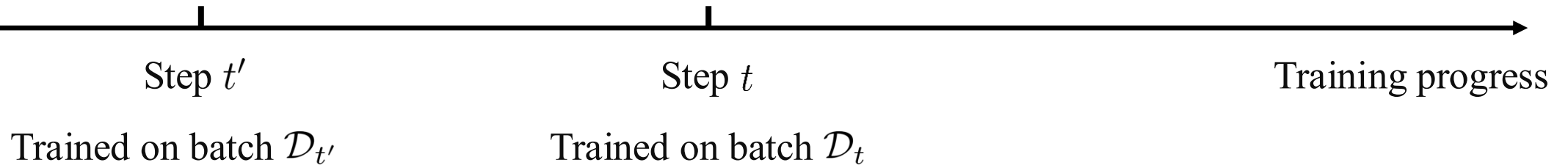
# The “Catastrophic Forgetting” Problem



# The “Catastrophic Forgetting” Problem

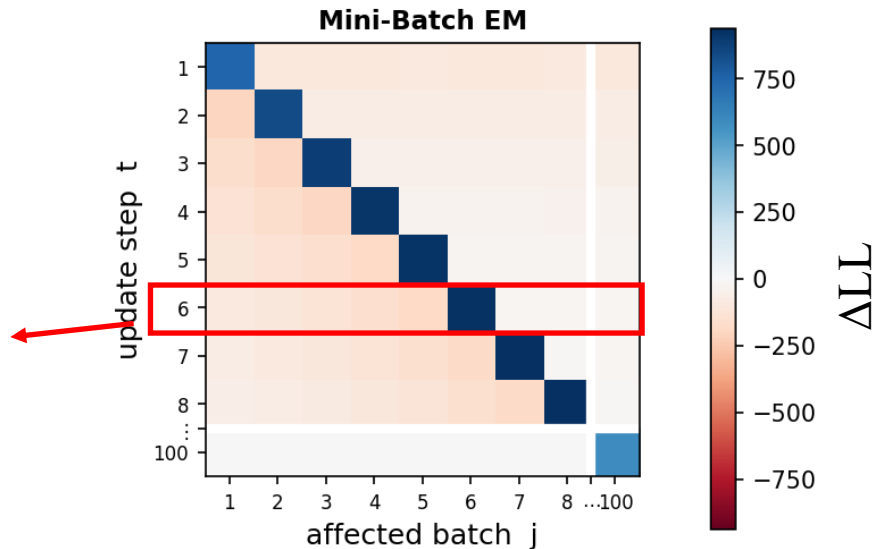


# The “Catastrophic Forgetting” Problem

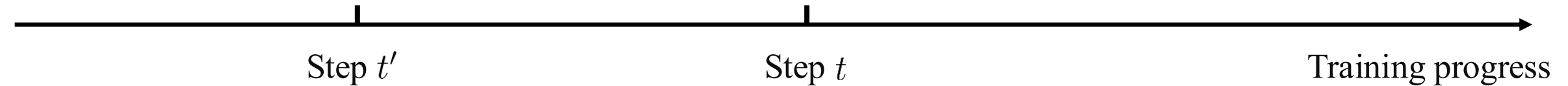


The LL of  $\mathcal{D}_{t'}$  should be well-preserved!

After training  
on batch #6 (in  
epoch 100)



# The “Catastrophic Forgetting” Problem

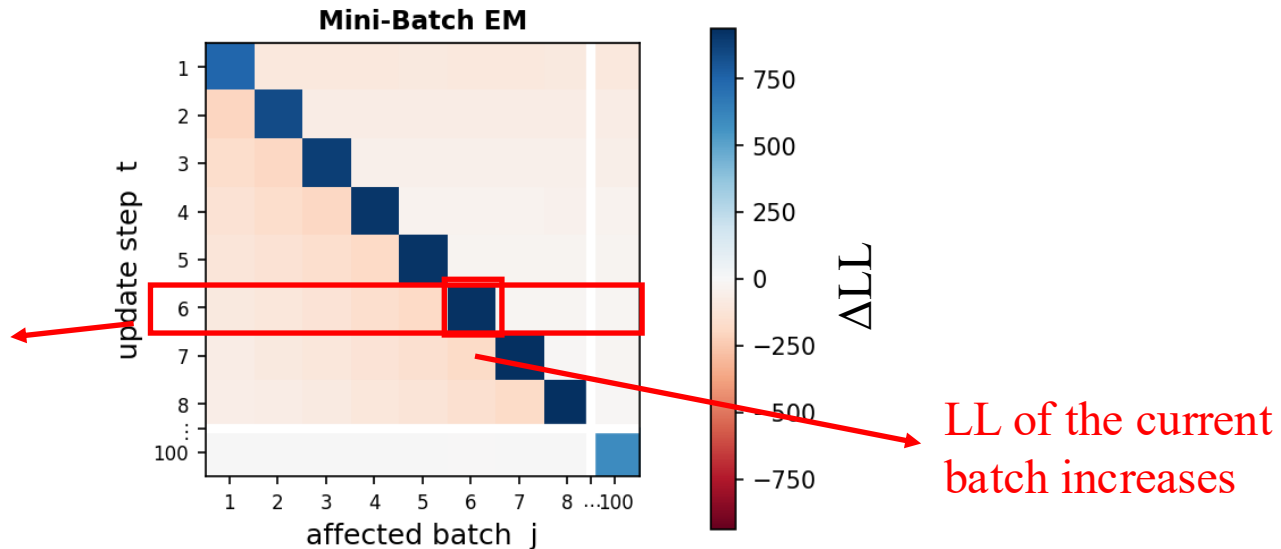


Trained on batch  $\mathcal{D}_{t'}$

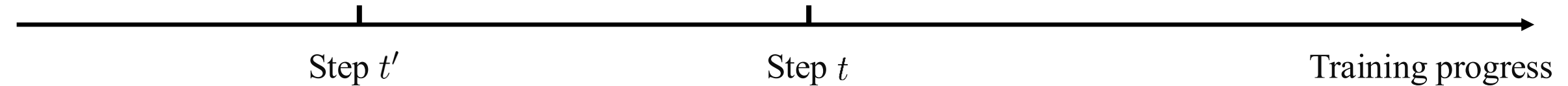
Trained on batch  $\mathcal{D}_t$

The LL of  $\mathcal{D}_{t'}$  should be well-preserved!

After training  
on batch #6 (in  
epoch 100)



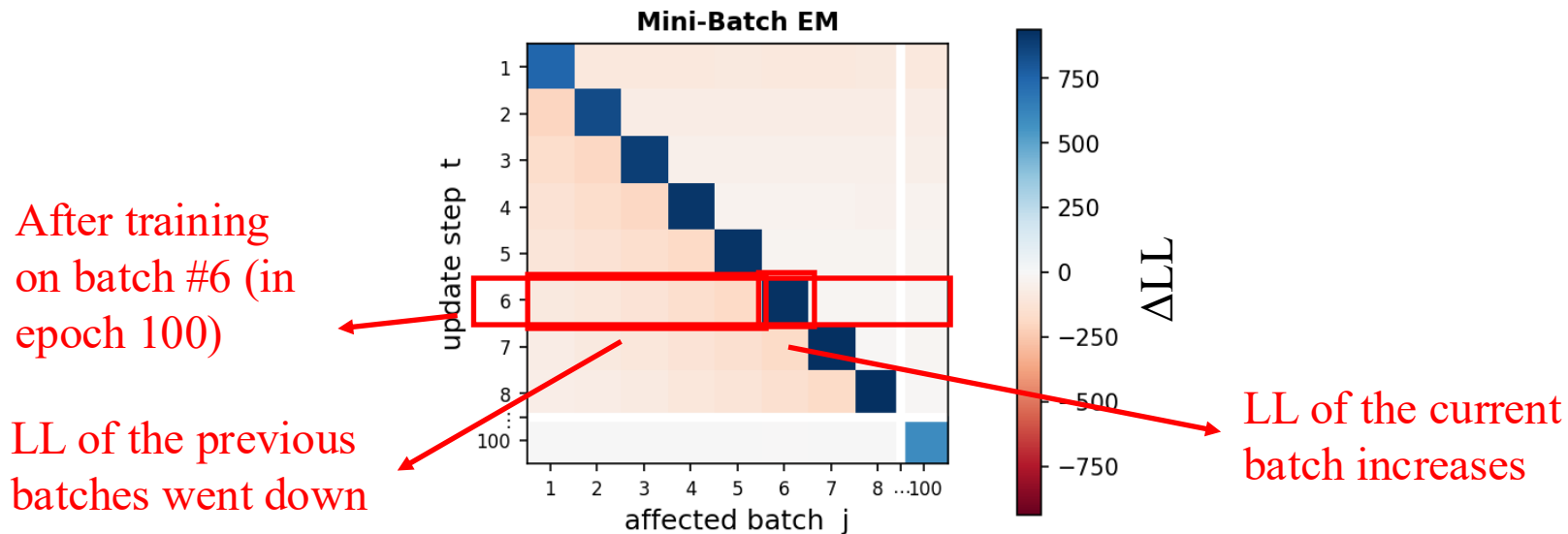
# The “Catastrophic Forgetting” Problem



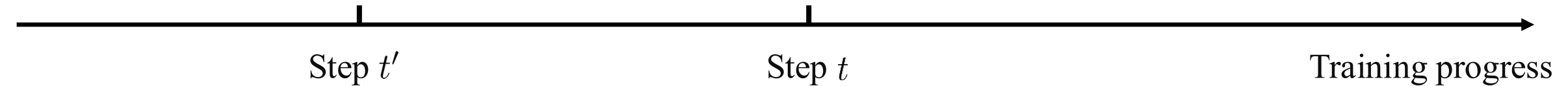
Trained on batch  $\mathcal{D}_{t'}$

Trained on batch  $\mathcal{D}_t$

The LL of  $\mathcal{D}_{t'}$  should be well-preserved!



# The “Catastrophic Forgetting” Problem



Trained on batch  $\mathcal{D}_{t'}$

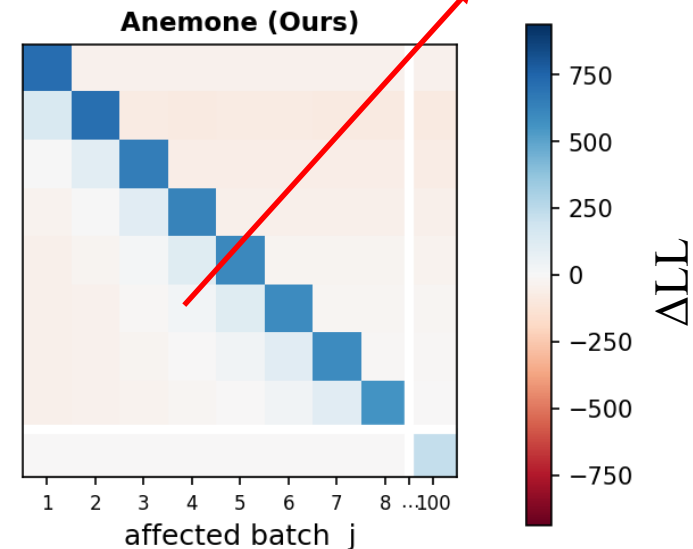
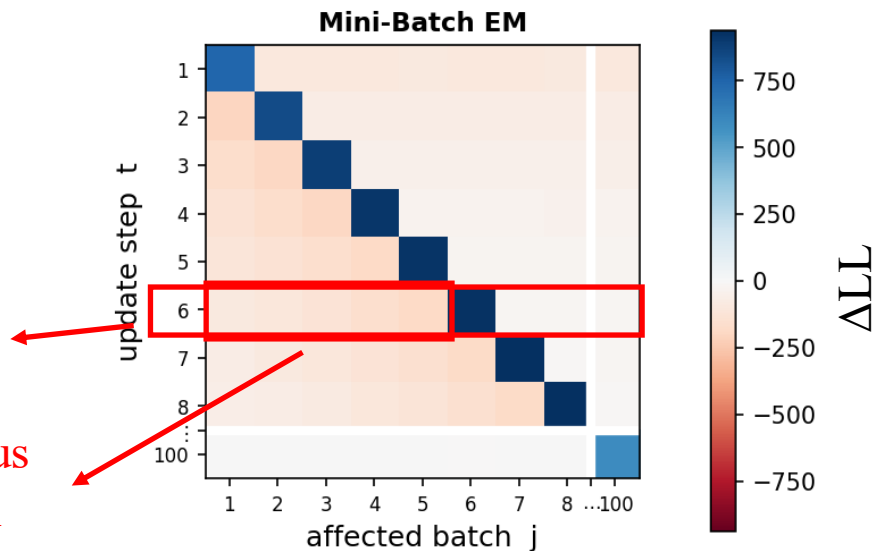
Trained on batch  $\mathcal{D}_t$

The LL of  $\mathcal{D}_{t'}$  should be well-preserved!

Our optimizer largely maintains the LL of previous batches

After training on batch #6 (in epoch 100)

LL of the previous batch went down



# Understanding the Catastrophic Forgetting Problem

## Reinterpreting full-batch EM\*

$$Q_{\phi}^{\mathcal{D}}(\phi') = \boxed{\frac{1}{|\mathcal{D}|} \sum_{\mathbf{x} \in \mathcal{D}} \log p_{\phi}(\mathbf{x}) + \left\langle \frac{\partial \log p_{\phi}(\mathbf{x})}{\partial \phi}, \phi' - \phi \right\rangle} - \boxed{\text{KL}_{\phi}(\phi')}$$

First-order Taylor  
expansion of the LL.

KL divergence: penalize  
changing the parameters  
too much.

\*Based on some results from the paper “Homeomorphic-Invariance of EM: Non-Asymptotic Convergence in KL Divergence for Exponential Families via Mirror Descent”. Kudos to the authors Frederik Kunstner, Raunak Kumar, and Mark Schmidt.

# Understanding the Catastrophic Forgetting Problem

## Reinterpreting full-batch EM\*

$$Q_{\phi}^{\mathcal{D}}(\phi') = \left[ \frac{1}{|\mathcal{D}|} \sum_{\mathbf{x} \in \mathcal{D}} \log p_{\phi}(\mathbf{x}) + \left\langle \frac{\partial \log p_{\phi}(\mathbf{x})}{\partial \phi}, \phi' - \phi \right\rangle \right] - \text{KL}_{\phi}(\phi')$$

Catastrophic forgetting happens because neither mini-batch EM nor Adam penalizes the parameters from “**changing too much**” properly.

\*Based on some results from the paper “Homeomorphic-Invariance of EM: Non-Asymptotic Convergence in KL Divergence for Exponential Families via Mirror Descent”. Kudos to the authors Frederik Kunstner, Raunak Kumar, and Mark Schmidt.

# Understanding the Catastrophic Forgetting Problem

## Reinterpreting full-batch EM\*

$$Q_{\phi}^{\mathcal{D}}(\phi') = \left[ \frac{1}{|\mathcal{D}|} \sum_{\mathbf{x} \in \mathcal{D}} \log p_{\phi}(\mathbf{x}) + \left\langle \frac{\partial \log p_{\phi}(\mathbf{x})}{\partial \phi}, \phi' - \phi \right\rangle \right] - \text{KL}_{\phi}(\phi')$$

Catastrophic forgetting happens because neither mini-batch EM nor Adam penalizes the parameters from “**changing too much**” properly.

## Gradient-based methods:

$$\text{maximize}_{\phi'} \frac{1}{|\mathcal{D}|} \sum_{\mathbf{x} \in \mathcal{D}} \log p_{\phi}(\mathbf{x}) + \left\langle \frac{\partial \log p_{\phi}(\mathbf{x})}{\partial \phi}, \phi' - \phi \right\rangle - \gamma \|\phi - \phi'\|_2^2$$

Small L2 distance in the parameter space could lead to large discrepancy in the distributions. Also vice versa.

# Understanding the Catastrophic Forgetting Problem

## Reinterpreting full-batch EM\*

$$Q_{\phi}^{\mathcal{D}}(\phi') = \frac{1}{|\mathcal{D}|} \sum_{\mathbf{x} \in \mathcal{D}} \log p_{\phi}(\mathbf{x}) + \left\langle \frac{\partial \log p_{\phi}(\mathbf{x})}{\partial \phi}, \phi' - \phi \right\rangle - \text{KL}_{\phi}(\phi')$$

Catastrophic forgetting happens because neither mini-batch EM nor Adam penalizes the parameters from “**changing too much**” properly.

## Mini-batch EM:

Similar to gradient-based methods, it only regularizes on the parameter space.

# The Anemone Algorithm

## Reinterpreting full-batch EM\*

$$Q_{\phi}^{\mathcal{D}}(\phi') = \frac{1}{|\mathcal{D}|} \sum_{\mathbf{x} \in \mathcal{D}} \log p_{\phi}(\mathbf{x}) + \left\langle \frac{\partial \log p_{\phi}(\mathbf{x})}{\partial \phi}, \phi' - \phi \right\rangle - \text{KL}_{\phi}(\phi')$$

Catastrophic forgetting happens because neither mini-batch EM nor Adam penalizes the parameters from “**changing too much**” properly.

## Anemone (our proposed method):



$$\text{maximize}_{\phi'} \frac{1}{|\mathcal{D}|} \sum_{\mathbf{x} \in \mathcal{D}} \log p_{\phi}(\mathbf{x}) + \left\langle \frac{\partial \log p_{\phi}(\mathbf{x})}{\partial \phi}, \phi' - \phi \right\rangle - \gamma \text{KL}_{\phi}(\phi')$$

# The Anemone Algorithm

**Anemone (our proposed method):**



$$\text{maximize}_{\phi'} \left[ \frac{1}{|\mathcal{D}|} \sum_{\mathbf{x} \in \mathcal{D}} \log p_{\phi}(\mathbf{x}) + \left\langle \frac{\partial \log p_{\phi}(\mathbf{x})}{\partial \phi}, \phi' - \phi \right\rangle \right] - \gamma \text{KL}_{\phi}(\phi')$$

Can be computed by autodiff!

Can also be efficiently computed!

**Lemma 1.** Assume the distributions defined by all nodes in a PC are normalized. For every  $\mathbf{x}$ , we have:

- (i)  $\partial \log p_{\phi}(\mathbf{x}) / \partial \phi = \partial \log \tilde{p}_{\phi}(\mathbf{x}) / \partial \phi - \text{TD}(\phi)$ ,
- (ii)  $\text{KL}_{\phi}(\phi') = -\langle \text{TD}(\phi), \phi' \rangle + C$ ,

where  $C$  is a constant term independent of  $\phi'$ .



# The Anemone Algorithm: An Intuition

Anemone can be viewed as EM with an **adaptive learning rate strategy (per parameter)**.

**Full-batch EM:**

$$\theta'_{n,c} = \hat{\mathbf{F}}_{\phi}^x(n, c) / Z$$

**Anemone:**

$$\theta'_{n,c} = (\theta_{n,c} + \boxed{\eta \cdot \text{rel}_{\phi}^x(n)} \cdot \boxed{\hat{\mathbf{F}}_{\phi}^x(n, c)}) / Z$$

↓  
The adaptive  
learning rate.

↓  
The target



# The Anemone Algorithm: An Intuition

Anemone can be viewed as EM with an **adaptive learning rate strategy (per parameter)**.

**Full-batch EM:**

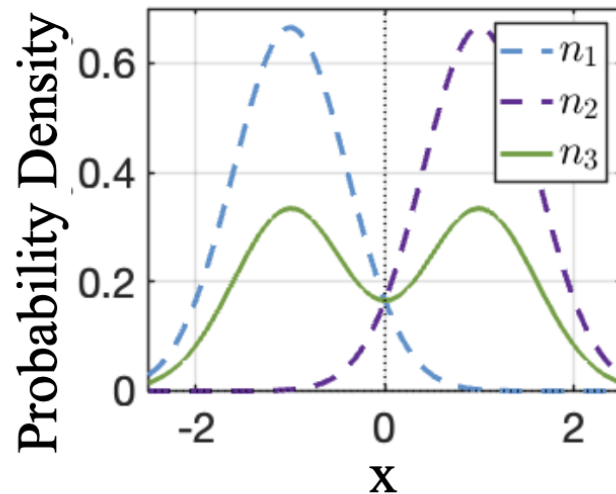
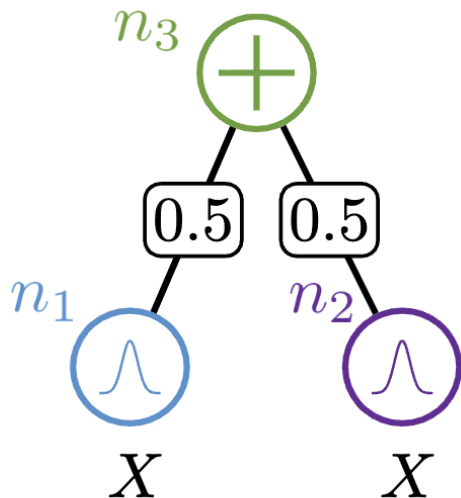
$$\theta'_{n,c} = \hat{F}_{\phi}^x(n, c) / Z$$

**Anemone:**

$$\theta'_{n,c} = (\theta_{n,c} + \eta \cdot \text{rel}_{\phi}^x(n) \cdot \hat{F}_{\phi}^x(n, c)) / Z$$

The adaptive  
learning rate.

The target



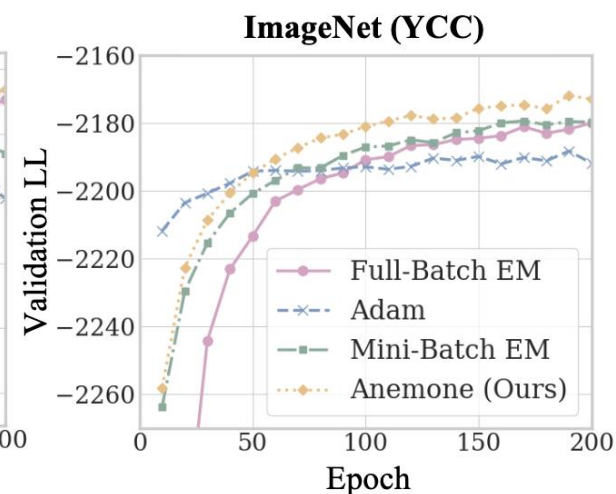
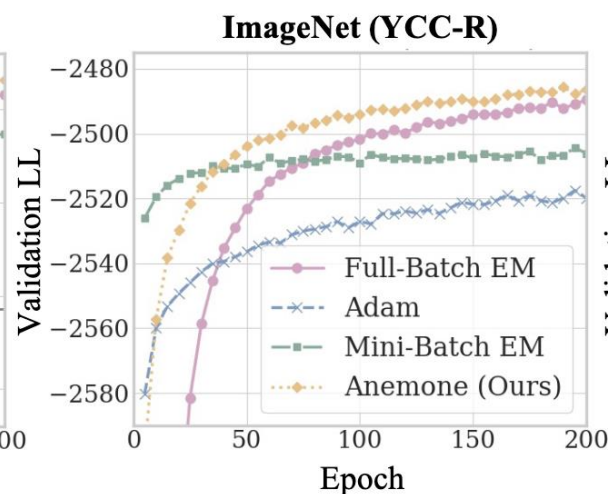
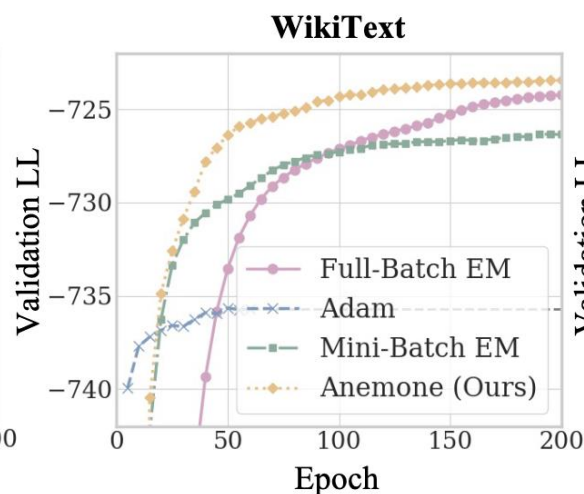
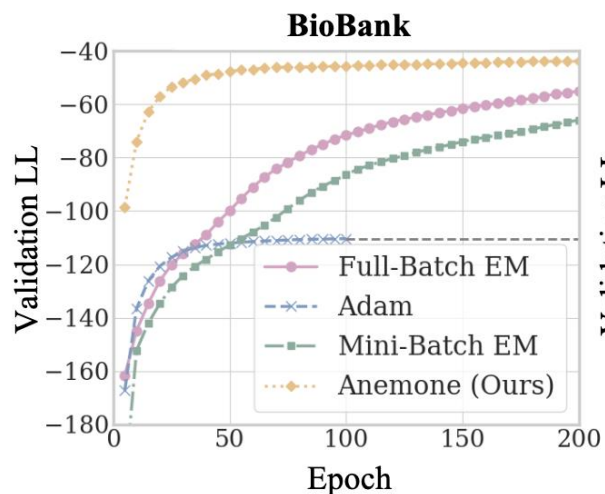
Given  $X = -1.5$ , the node  $n_1$  is more “important” for producing high LL.

Anemone will implicitly assign a much **higher adaptive learning rate** to parameters of  $n_1$ .



# The Anemone Algorithm: Results

- Better LL and faster convergence across benchmarks





# The Anemone Algorithm: Results

- Better LL and faster convergence across benchmarks
- Robust across data types and model architectures.

| Optimizer      | BioBank Chr6           |                        |                        |             |
|----------------|------------------------|------------------------|------------------------|-------------|
|                | PDHCLT 512             | PDHCLT 1024            | HCLT 512               | HCLT 1024   |
| Full-batch EM  | 45.9 $\pm$ 0.55        | 45.0 $\pm$ 0.34        | 55.2 $\pm$ 0.17        | 53.8        |
| Adam           | 113.7 $\pm$ 1.17       | OOM                    | 102.9 $\pm$ 0.22       | 100.4       |
| Mini-batch EM  | 49.6 $\pm$ 0.42        | 48.8 $\pm$ 1.41        | 55.0 $\pm$ 0.56        | 52.9        |
| Anemone (Ours) | <b>44.7</b> $\pm$ 0.48 | <b>42.3</b> $\pm$ 0.31 | <b>54.5</b> $\pm$ 0.16 | <b>52.1</b> |

| Optimizer      | ImageNet32 YCC-R       |                        |             |                  |                   |
|----------------|------------------------|------------------------|-------------|------------------|-------------------|
|                | PDHCLT 256             | HCLT 512               | HCLT 1024   | Monarch HCLT 256 | Monarch HCLT 1024 |
| Full-batch EM  | 2529 $\pm$ 0.48        | 2483 $\pm$ 3.11        | <b>2469</b> | 2526             | 2468              |
| Adam           | 2555 $\pm$ 1.96        | 2518 $\pm$ 0.29        | OOM         | —                | 2588              |
| Mini-batch EM  | <b>2527</b> $\pm$ 1.52 | 2506 $\pm$ 0.63        | 2470        | —                | 2496              |
| Anemone (Ours) | 2529 $\pm$ 1.87        | <b>2480</b> $\pm$ 2.46 | <b>2469</b> | <b>2516</b>      | <b>2465</b>       |

| Optimizer      | WikiText                |                         |                         |                   |
|----------------|-------------------------|-------------------------|-------------------------|-------------------|
|                | HMM 256                 | HMM 512                 | HMM 1024                | Monarch HCLT 1024 |
| Full-batch EM  | 723.3 $\pm$ 0.88        | 703.2 $\pm$ 1.28        | 683.4 $\pm$ 1.10        | 738.1             |
| Adam           | 735.5 $\pm$ 0.29        | OOM                     | OOM                     | OOM               |
| Mini-batch EM  | 724.6 $\pm$ 0.34        | 704.0 $\pm$ 1.11        | 683.6 $\pm$ 1.41        | 734.6             |
| Anemone (Ours) | <b>722.7</b> $\pm$ 0.48 | <b>702.5</b> $\pm$ 0.81 | <b>682.6</b> $\pm$ 0.33 | <b>734.0</b>      |

**Thank you!**  
**Please consider stopping by our  
poster #177 (today).**