

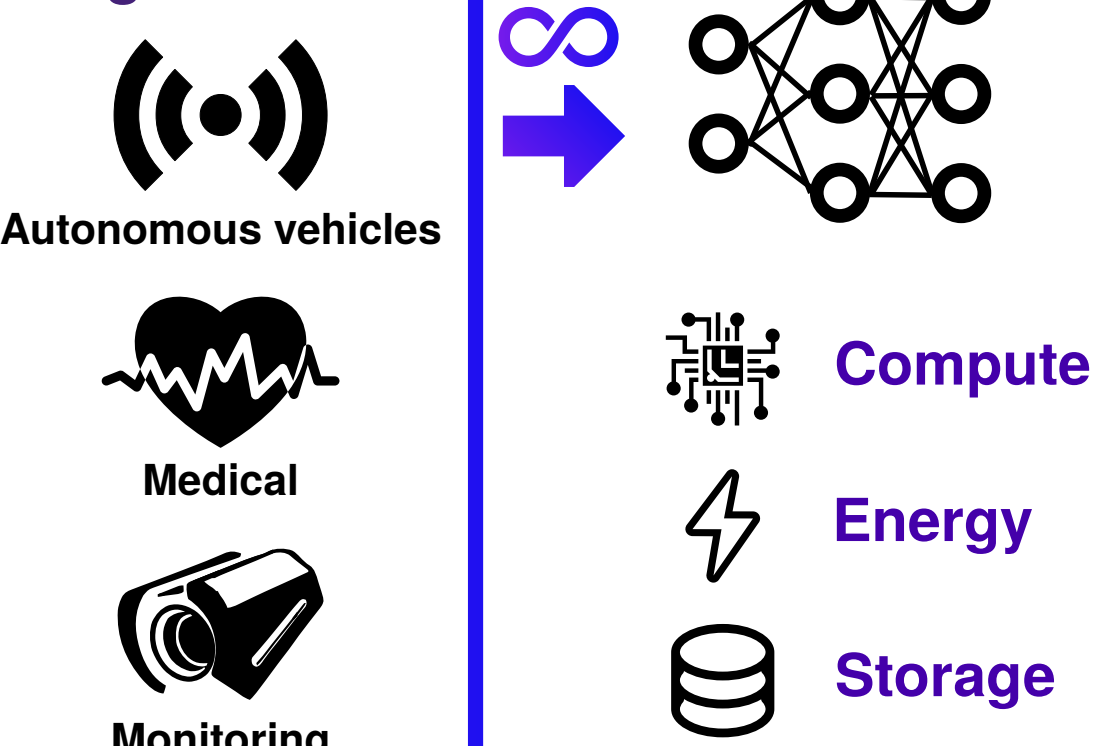


1. CONTEXT: AI AT THE EXTREME EDGE

Endless streams from sensors

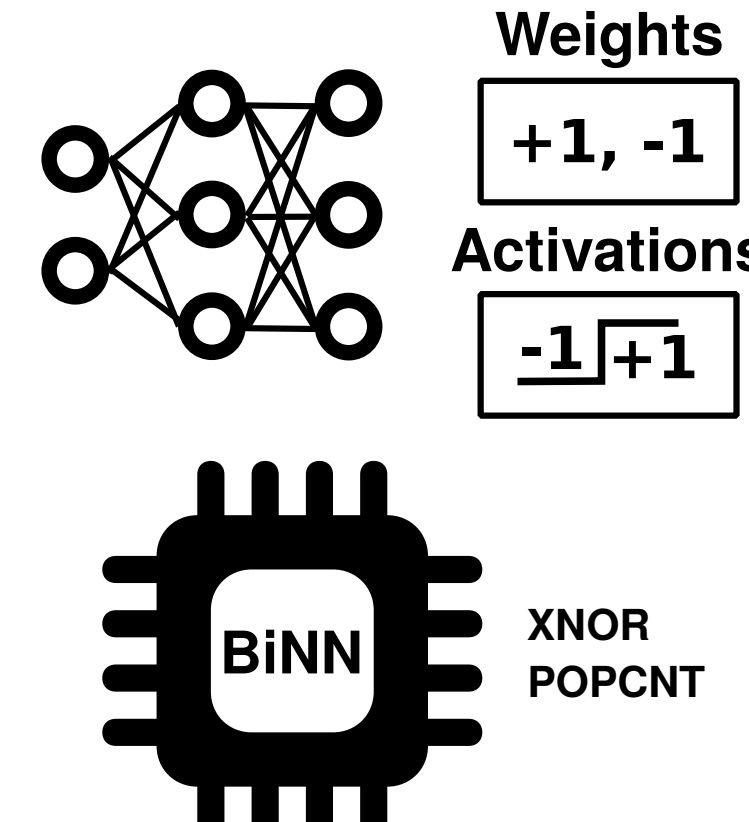
At the edge, deep neural networks face a core tension: AI must run **inference continuously**

Edge devices



Binary Neural Networks

Ultra-**fast**, ultra-**efficient** inference (bitwise MACs)



Why is training at the edge challenging?

Streaming and non stationary data

- Input data distribution changes over time $\Rightarrow$ Catastrophic forgetting and catastrophic remembering
- Epistemic uncertainty collapses over time due to weight saturation $\Rightarrow$ Vanishing uncertainties
- Imbalanced data with rare, pathological events $\Rightarrow$ Heavy missclassification on potentially critical classes

Our contributions

- Binary Metaplasticity from Uncertainty (BiMU):** Three forces for **stable**, **plastic** and **forgetful** learning in Binary Bayesian Neural Networks
- Maintaining uncertainty via metaplasticity:** synapses learn as a function of **variance** and **curvature**
- Continual active learning:** Retaining uncertainty over time and using **disagreement** to **select informative training examples**

2. MODELING: BINARIZED BAYESIAN NEURAL NETWORKS

Each weight a Bernoulli random variable over $\{-1, 1\}$. We hold a mean-field posterior q over weights.

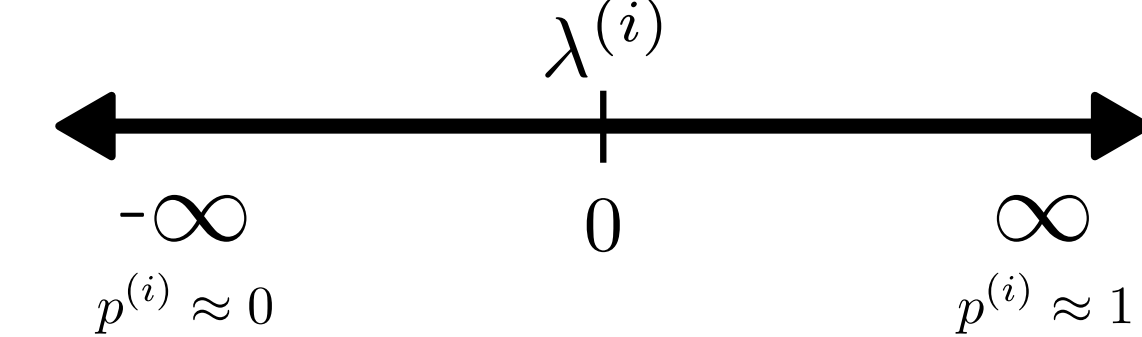
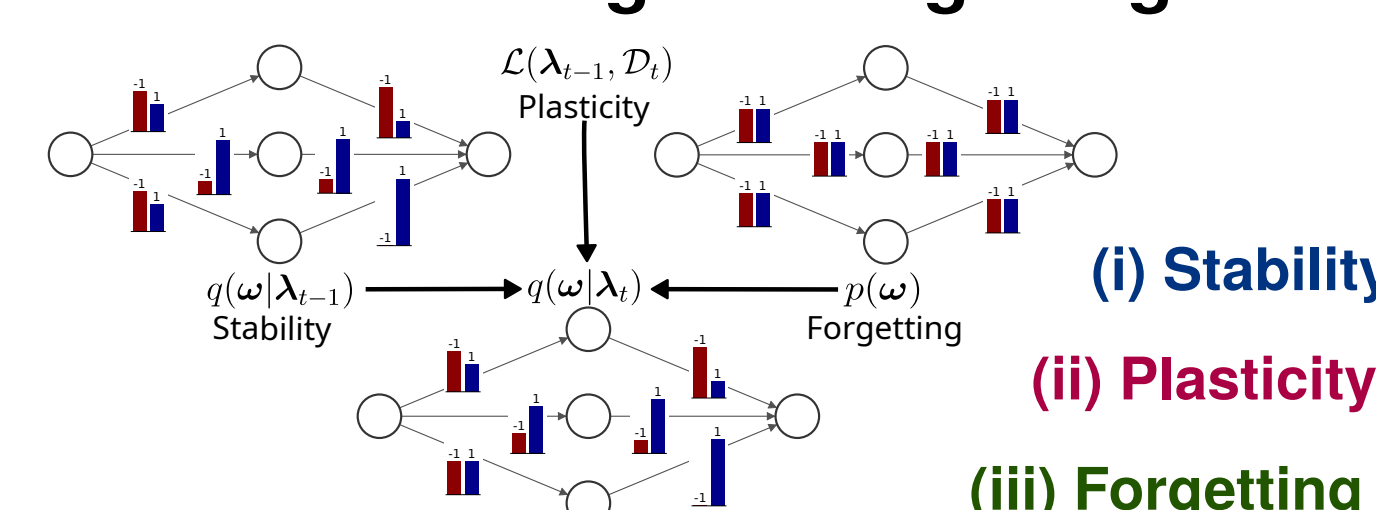
Bernoulli weights

$$\omega \in \{-1, +1\}^s \mid q(\omega|\lambda) = \prod_{i=1}^s \frac{\exp(\lambda^{(i)} \omega^{(i)})}{2 \cosh(\lambda^{(i)})}$$

Natural Parameter

$$p^{(i)} = \text{sigmoid}(2\lambda^{(i)})$$

Mean-field Variational Inference Continual learning and forgetting



Divergence leads to **deterministic behavior** (catastrophic forgetting)

3. BiMU: LEARNING RULE FOR BINARY ONLINE LIFELONG LEARNING

Binary Metaplasticity from Uncertainty

$$\lambda_t^{(i)} = \lambda_{t-1}^{(i)} - \eta(\lambda_{t-1}^{(i)}) \left[\frac{\partial \mathcal{L}}{\partial \lambda^{(i)}} \Big|_{\lambda_{t-1}} + \frac{\lambda_{t-1}^{(i)} - \lambda_{\text{prior}}^{(i)}}{N \cosh^2(\lambda_{t-1}^{(i)})} \right]$$

BiMU updates the natural parameter of each weight in **online settings**

(i) Metaplastic learning rate (Stability)

(ii) Gradient term (Plasticity)

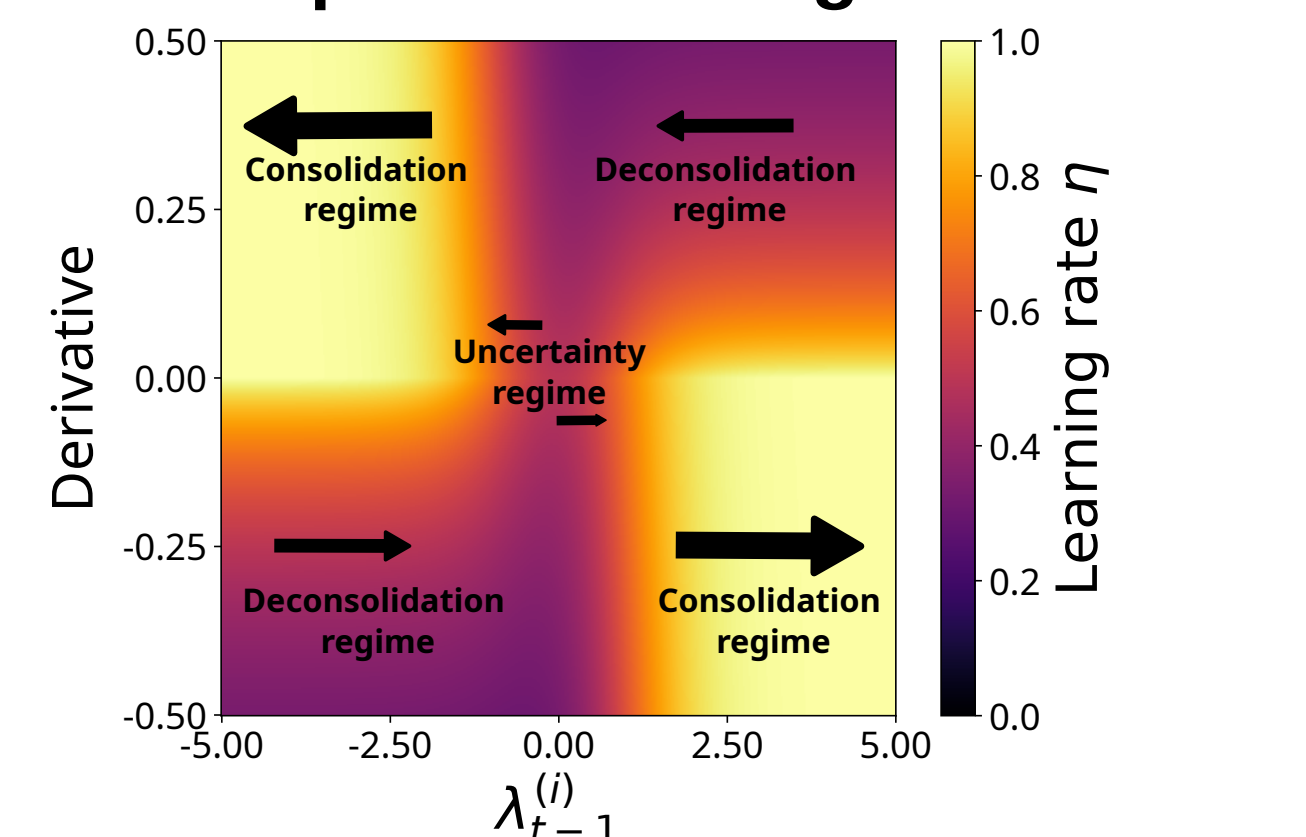
(iii) Uncertainty-gated pull towards prior (Forgetting)

- Slows strong opposing gradient** to the most probable sign
- Preserves non-consolidated** synapses

- Drives the natural parameter learning** (probability of sign)

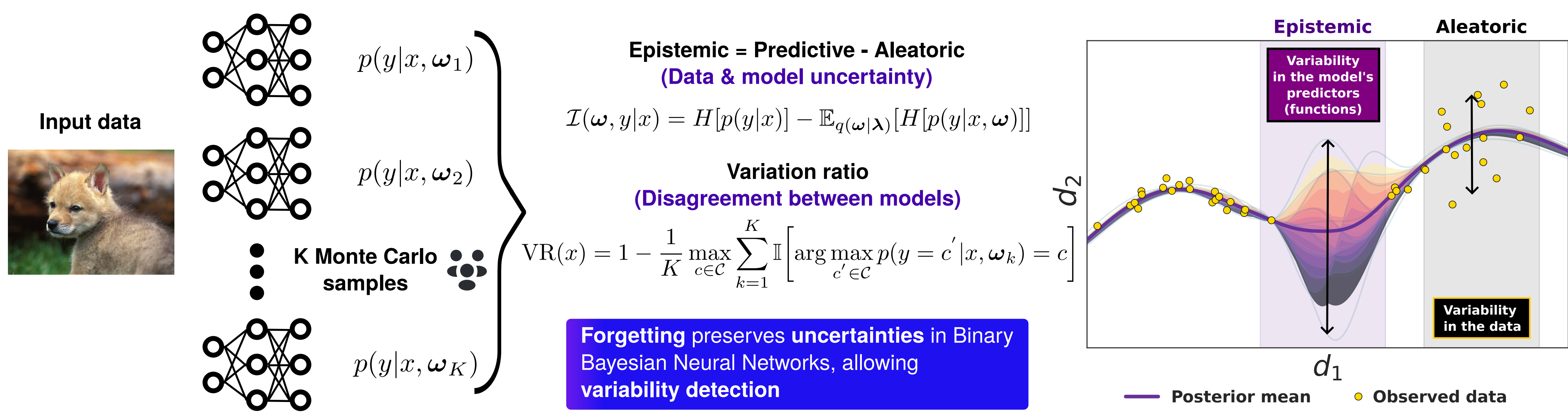
- Sliding window N** controls the **forgetting**
- Synapses with **high variance** pull to the **prior**

Metaplastic learning rate



$$\frac{1}{\eta(\lambda_{t-1}^{(i)})} = \frac{1}{\cosh^2(\lambda_{t-1}^{(i)})} + 2 \tanh(\lambda_{t-1}^{(i)}) \frac{\partial \mathcal{L}}{\partial \lambda^{(i)}} \Big|_{\lambda_{t-1}} + \frac{1}{\alpha_{\text{max}}} + 2 \left| \frac{\partial \mathcal{L}}{\partial \lambda^{(i)}} \Big|_{\lambda_{t-1}} \right|$$

4. METHODOLOGY FOR UNCERTAINTY QUANTIFICATION

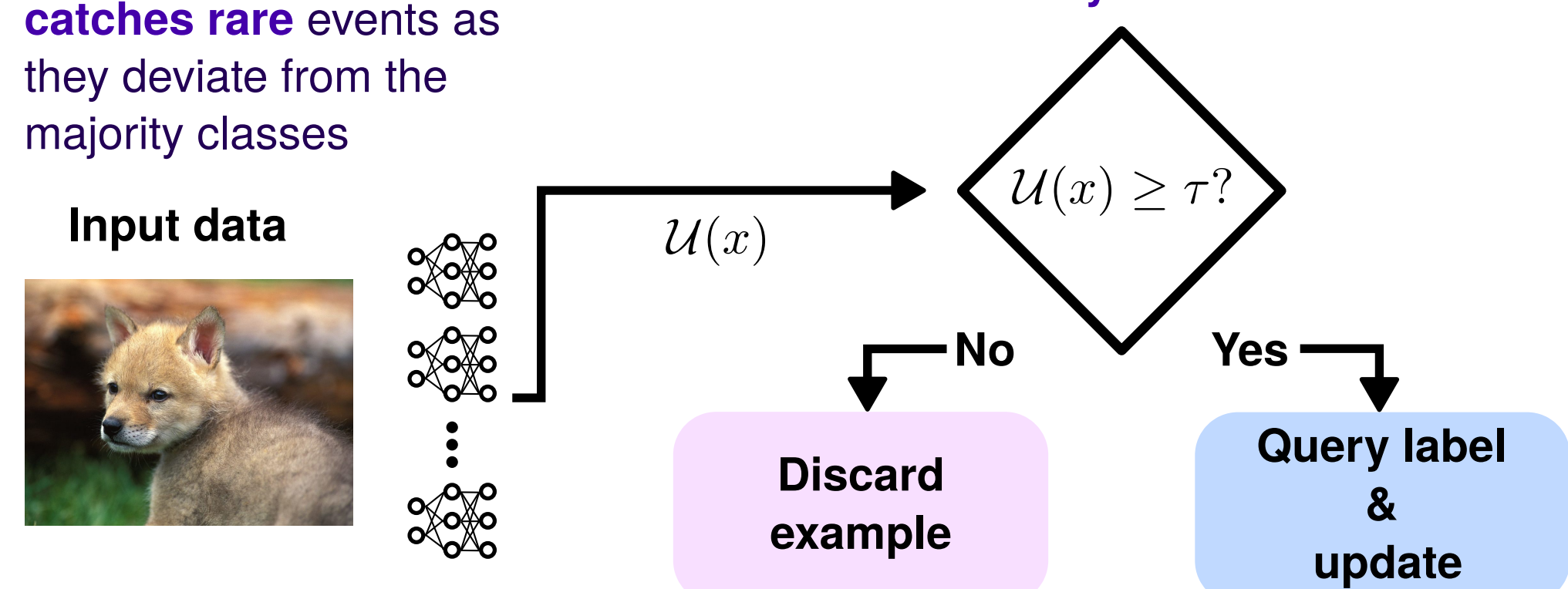


5. METHODOLOGY FOR ACTIVE LEARNING

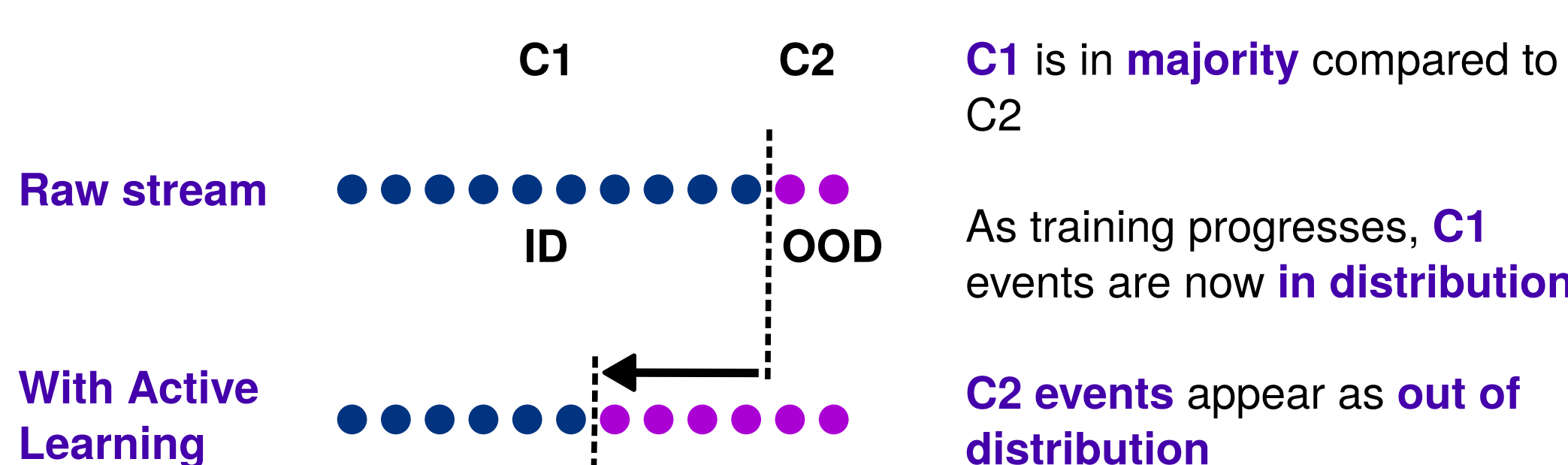
Active learning at the edge

Intuition: uncertainty catches rare events as they deviate from the majority classes

Query labels from examples that exceed the **uncertainty threshold**



Why active learning helps with data imbalance?



Active learning queries more rare/uncertain examples, recovering informative events

6. EXPERIMENTS & RESULTS: 32x UPDATES REDUCTION ON LIFELONG LEARNING

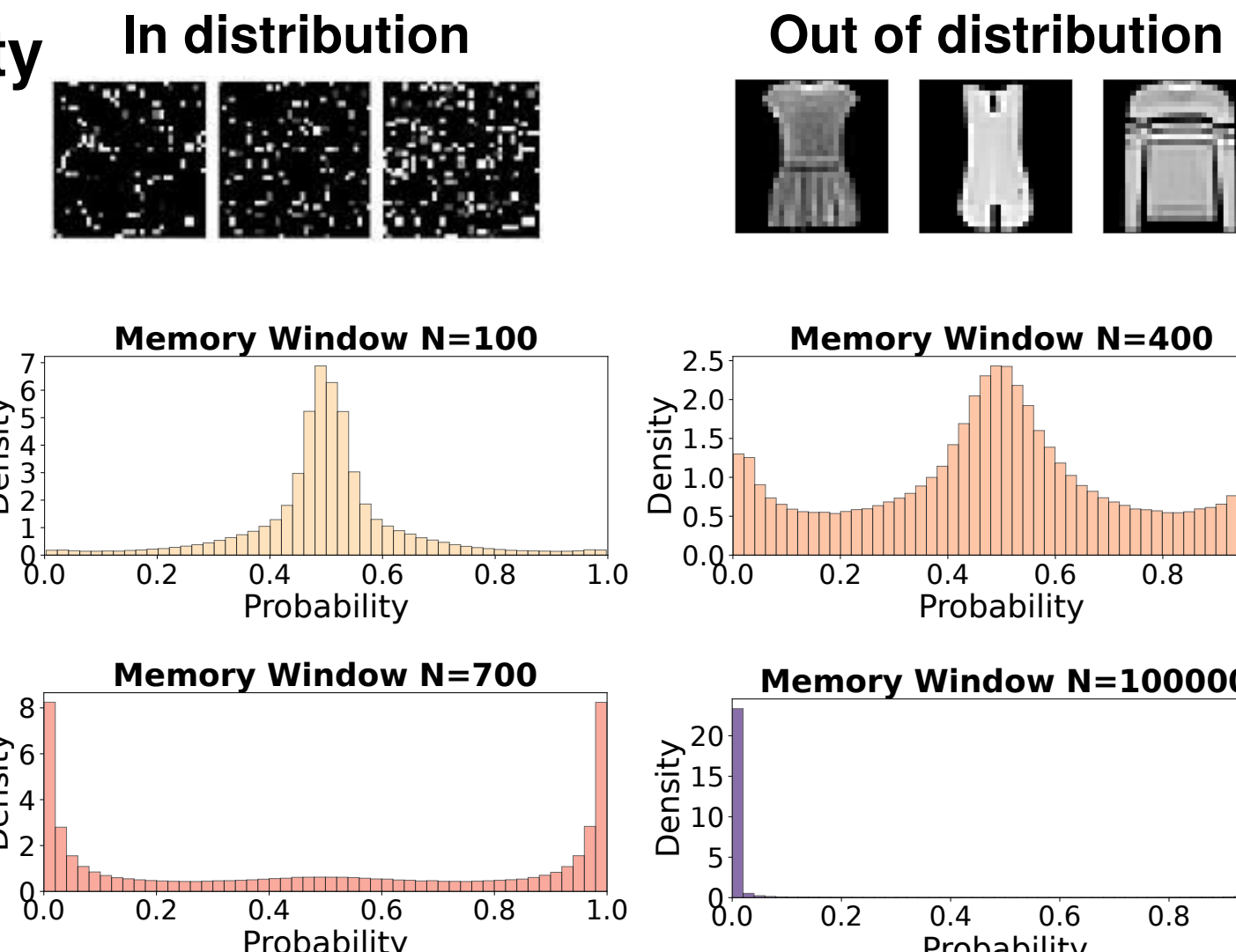
A. Long-Horizon continual learning: mitigating forgetting and loss of plasticity (1000-tasks Permuted MNIST)

Method	Task bounds	MEAN ACC. 5 tasks (%)	OOD det. (AUC)	MMRR	Acc. 1 task (%)
<i>Binary neural networks</i>					
BiMU	NO	90.30 $\pm$ 0.38	0.99 $\pm$ 0.00	139.47	94.67 $\pm$ 0.11
BAYESBiNN	YES	41.12 $\pm$ 1.62	0.57 $\pm$ 0.12	2.04	93.22 $\pm$ 0.09
SYN. META.	YES	10.27 $\pm$ 0.01	-	1.64	71.40 $\pm$ 1.48
STE	NO	29.35 $\pm$ 0.96	0.69 $\pm$ 0.04	9.32	77.56 $\pm$ 1.35
<i>Real-valued neural networks</i>					
MESU	NO	91.69 $\pm$ 0.58	0.95 $\pm$ 0.03	261.10	96.10 $\pm$ 0.18
EWC O.	YES	81.78 $\pm$ 0.82	0.66 $\pm$ 0.11	6.63	96.06 $\pm$ 0.11
SI	YES	74.41 $\pm$ 1.19	0.66 $\pm$ 0.17	5.11	95.55 $\pm$ 0.28
SGD	NO	66.64 $\pm$ 2.70	0.87 $\pm$ 0.05	43.50	96.03 $\pm$ 0.34

Protocol: 100 hidden units, fully online

BiMU is stable (high mean accuracy), **plastic** (high OOD detection), **forgetful** (high MMRR)

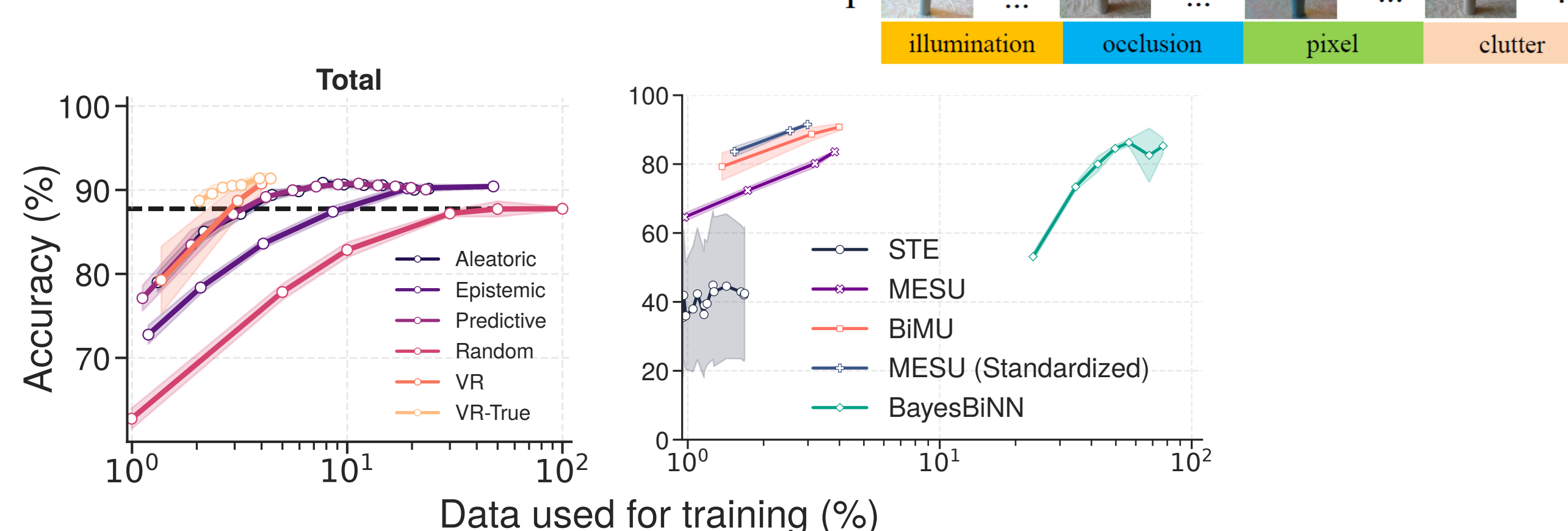
- Uncertainties** are maintained over **1000 tasks**
- Memory window N** enforces **forgetting or stability**



B. Imbalanced active continual learning: balancing classes, reducing annotations (OpenLORIS-Object)

Protocol: Frozen features, linear classifier, fully online
Classes with less than 700 training examples have their training images removed up to 50-80%

- Variation Ratio (VR)** active learning enables a **32x reduction in updates**
- Outperforms full training** with only **4.0%** of the budget
- BiMU (Binary)** performs **similarly** to **MESU (Real-valued)**



Low frequency accuracy drops significantly when using active learning with random examples

VR maintains the accuracy of **both high frequency classes** and **low frequency classes**

At the 32x reduction, BiMU with VR recovers **28% accuracy** on the **low frequency classes**

Takeaway: BiMU allows efficient continual learning at the edge by preserving Bernoulli distribution uncertainties, using fast inference to compute disagreement, leading to selective updates