

Motivation

- **Clustering**-based features widely used.
- Most methods choose ONE resolution γ — that is global, fixed and ad hoc.
- Real data exhibits multi-scale structure.
- Sweeping $\gamma \in [0, \infty)$ is intractable.
- Recent work introduces *configurations* — a finite set of stable partitions (see below)

Q: adaptively weight all stable partitions (configurations) for downstream prediction?

Preliminaries: Configurations

Def. 3.1: A partition of n items is represented by an assignment vector $\omega \in \{1, \dots, |\omega|\}^n$.

Hamiltonian energy (LeCun et al., 2006):

$$H(\omega) = h_a + \gamma h_r \quad (1)$$

$$h_a = - \sum_{k=1}^{|\omega|} \sum_{i < j} w_{ij} \mathbf{1}_{\omega_i = \omega_j = k} \quad (\text{attraction})$$

$$h_r = \sum_{k=1}^{|\omega|} \left(\sum_{i < j} w_{ij} \mathbf{1}_{\omega_i = k} \right)^2 \quad (\text{repulsion})$$

where $w_{ij} \geq 0$ are pairwise geodesic distance.

Def. 3.2: Let $\omega^*(\gamma) := \arg \min H(\omega)$. A *configuration* $\omega = \omega^*(\gamma)$ for all $\gamma \in [\gamma^-, \gamma^+]$. Sweeping $\gamma \in [0, \infty)$ yields a finite set of distinct configurations $\Omega = \{\omega_1, \dots, \omega_m\}$.

Parallel-DT (Pitsianis et al., 2023) extracts all such coarse→fine configurations in $O(n \log n)$.

Configuration-Mixed Prediction

Jointly learn predictor f_ψ and selector $\mathbf{w}_\theta(\mathbf{x})$:

$$\mathcal{L}_{\text{CMP}} := \frac{1}{n} \sum_{t=1}^n \ell(f_\psi([x_t; z_\theta(x_t)]), y_t), \quad (2)$$

$$z_\theta(\mathbf{x}) := \sum_{i=1}^m w_{\theta,i}(\mathbf{x}) \phi(\omega_i, \mathbf{x}),$$

$$(\theta^*, \psi^*) \in \arg \min_{\theta, \psi} \mathcal{L}_{\text{CMP}}(\theta, \psi).$$

where $\mathbf{w}_\theta(\mathbf{x}) = (w_{\theta,1}(\mathbf{x}), \dots, w_{\theta,m}(\mathbf{x}))$; ϕ encodes $\mathbf{x}$ under ω_i ; and ℓ is a prediction loss.

We group five predictor classes by how the mixing is done:

1. $\mathcal{F}_{\text{raw}}$: $f_\psi(\mathbf{X})$ — uses raw features only. = Base
2. $\mathcal{F}_{\text{single}}$: $f_\psi(\mathbf{X}, \omega_i)$ — augments with one configuration. = +Single
3. $\mathcal{F}_{\text{static}}$: $f_\psi(\mathbf{X}) + \bar{\mathbf{w}}^T \Omega$ — with a global weight. = +Static (global $\bar{\mathbf{w}}$)
4. $\mathcal{F}_{\text{standard}}$: $f_\psi(\mathbf{X}, \Omega)$ — concatenates all as features. = +Config
5. $\mathcal{F}_{\text{adaptive}}$: $f_\psi(\mathbf{X}, \mathbf{w}_\theta(\mathbf{X})^T \Omega)$ — learns $\mathbf{w}_\theta(\mathbf{X})$. = +MixConfig

By construction, $\mathcal{F}_{\text{raw}} \subseteq \mathcal{F}_{\text{single/static}} \subseteq \mathcal{F}_{\text{standard}} \subseteq \mathcal{F}_{\text{adaptive}}$.

Prop. 3.4: The inclusion $\mathcal{F}_{\text{static}} \subset \mathcal{F}_{\text{adaptive}}$ is strict. A deep $\mathcal{F}_{\text{standard}}$ may in principle match it. Empirically, it does not (see Ablations).

MixConfig

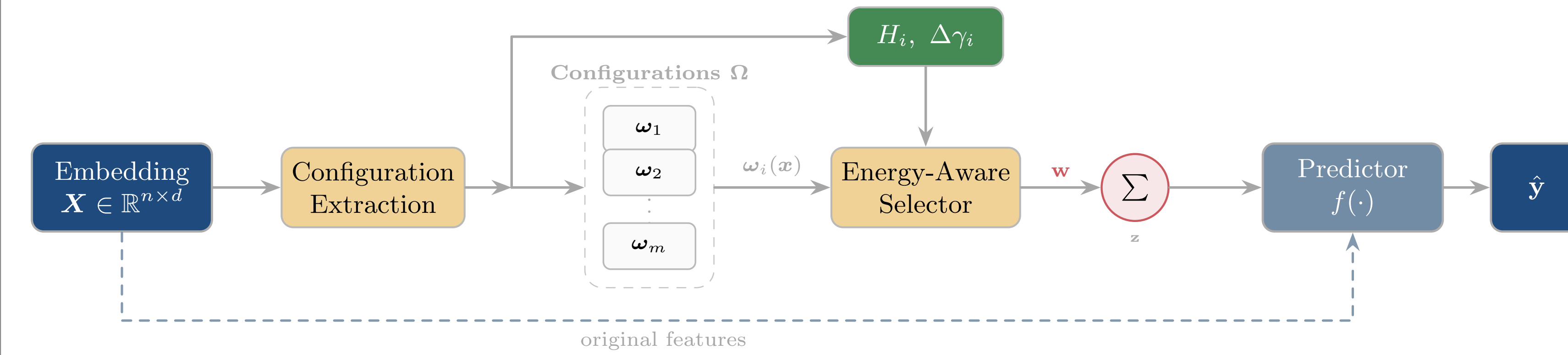


Figure 1. Frozen embedding → MixConfig (extract configs + energy-aware mix) → predictor.

The energy-aware selector learns *which resolutions matter for which samples*, leveraging stability stats as inductive bias:

$$\mathbf{h} = \text{MLP}_{\text{enc}}(\mathbf{x}) \quad (\text{sample context}) \quad (3)$$

$$\mathbf{c}_i = \text{Embed}_i(\omega_i(\mathbf{x})) \quad (\text{cluster assignment}) \quad (4)$$

$$\mathbf{e}_i = [H_i, h_a^{(i)}, h_r^{(i)}, \Delta\gamma_i] \quad (\text{energy statistics}) \quad (5)$$

$$s_i = \text{MLP}_{\text{score}}([\mathbf{h}; \mathbf{c}_i; \mathbf{e}_i]) \quad (\text{compatibility score}) \quad (6)$$

$$\mathbf{w}_i(\mathbf{x}) = \text{softmax}(s_i(\mathbf{x})) \quad \mathbf{w}(\mathbf{x}) \in \Delta^{m-1} \quad (7)$$

We then form a mixed configuration representation $\mathbf{z}(\mathbf{x}) = \sum_{i=1}^m w_i(\mathbf{x}) \mathbf{c}_i$ and pass the augmented features $[\mathbf{x}; \mathbf{z}(\mathbf{x})]$ to the downstream predictor.

Consistent Gains Across Four Domains

We evaluate MixConfig across four benchmark families. Our configuration-based methods (+Config, +MixConfig) consistently win.

Benchmark	Embedding	Base	+DeepCluster	+HDBSCAN	+Config	+MixConfig
<i>Tabular (OpenML-CC18)</i>						
Binary (AUC $\uparrow$)	TabPFN	.895 $\pm$.004	.900** $\pm$.005	.898 $\pm$.004	.903** $\pm$.004	.907** $\pm$.003
Multi-class (Acc. $\uparrow$)	FT-Trans.	.782 $\pm$.006	.779 $\pm$.007	.785 $\pm$.006	.789** $\pm$.005	.793** $\pm$.004
<i>Vision</i>						
CIFAR-100 (Top-1 $\uparrow$)	CLIP	.742 $\pm$.003	.740 $\pm$.005	.746 $\pm$.003	.751** $\pm$.003	.758** $\pm$.002
ImageNet-1K (Top-1 $\uparrow$)	CLIP	.763 $\pm$.002	.764 $\pm$.003	.768 $\pm$.003	.773** $\pm$.002	.782** $\pm$.002
<i>Molecular</i>						
MoIHIV (AUC $\uparrow$)	GIN	.804 $\pm$.008	.801 $\pm$.010	.809 $\pm$.007	.814** $\pm$.007	.819** $\pm$.005
QM9 (MAE $\downarrow$)	DimeNet++	.0112 $\pm$.0002	.0114 $\pm$.0004	.0109 $\pm$.0002	.0106** $\pm$.0002	.0106** $\pm$.0001
<i>Text</i>						
SST-2 (Acc. $\uparrow$)	RoBERTa	.945 $\pm$.003	.944 $\pm$.004	.948 $\pm$.003	.951** $\pm$.003	.955** $\pm$.002
AG News (Acc. $\uparrow$)	BERT	.942 $\pm$.004	.946** $\pm$.004	.944 $\pm$.004	.948** $\pm$.003	.951** $\pm$.003

Table 1. Bold = best; ** $p < 0.05$ vs Base (paired t -test, 5 seeds); all Cohen’s $d > 0.8$.

Ablations

We isolate MixConfig’s components and compare mixing strategies.

Variant	BBBP	CIFAR-100
MixConfig (full)	.903	.758
w/o energy features	.882	.738
w/o sample context $\mathbf{h}$	.890	.747
w/o cluster embeddings $\mathbf{c}_i$	.886	.743
Using only m' of m configs:		
$m' = 2$	.864	.716
$m' = 4$	.882	.737
$m' = 6$	.895	.751
$m' = m$ (all)	.903	.758

Table 2. Each component helps.

Gains Grow as Labels Shrink

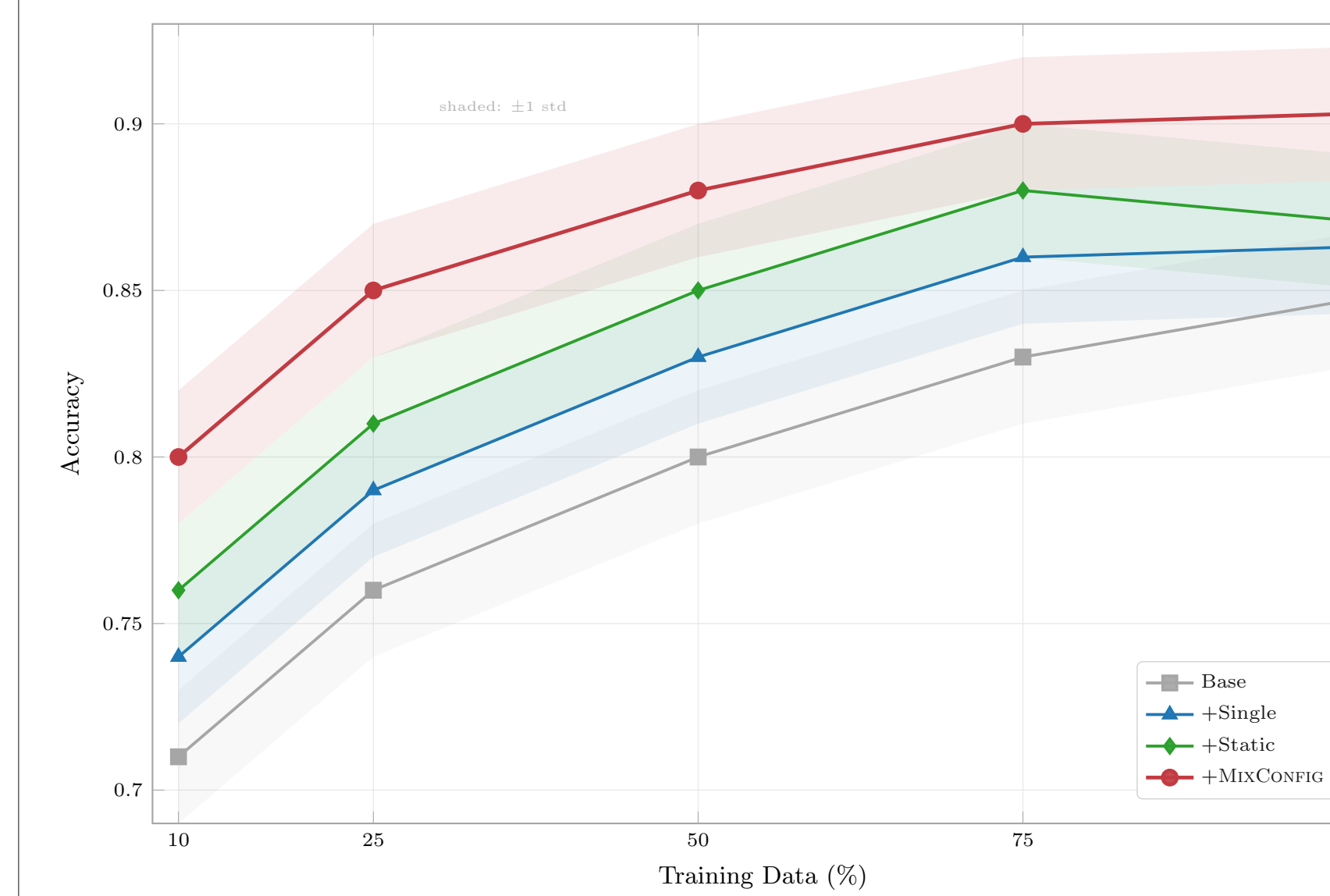


Figure 2. Performance vs. training set size.

Deployment

- Predictor-agnostic on frozen embeddings; $d_c=32$ selector MLP.
- Batch-transductive ($n \gtrsim 50$); inductive k NN fallback — no test-time extraction, retains $\geq 75\%$ of gains.
- One-time extraction $\sim 150\times$ cheaper than a resolution sweep.

Takeaways

- **Energy-aware mixing:** uniform / global $\bar{\mathbf{w}}$ / stronger-MLP controls all $<$ MixConfig; adds 32 dims, beats +Config’s 78.
- **Plug-and-play** on any frozen embedding & predictor.
- **Biggest gains when labels are scarce** ($\sim 5\times$ efficiency).
- **Best on hierarchical, heterogeneous, low-data problems; not** single-scale data or tiny batches.

Links & References

LeCun, Y., Chopra, S., Hadsell, R., Ranzato, M., and Huang, F. J. A Tutorial on Energy-Based Learning. 2006.

Pitsianis, N., Floros, D., Liu, T., and Sun, X. Parallel Clustering with Resolution Variation. In *2023 IEEE High Performance Extreme Computing Conference (HPEC)*, 2023.

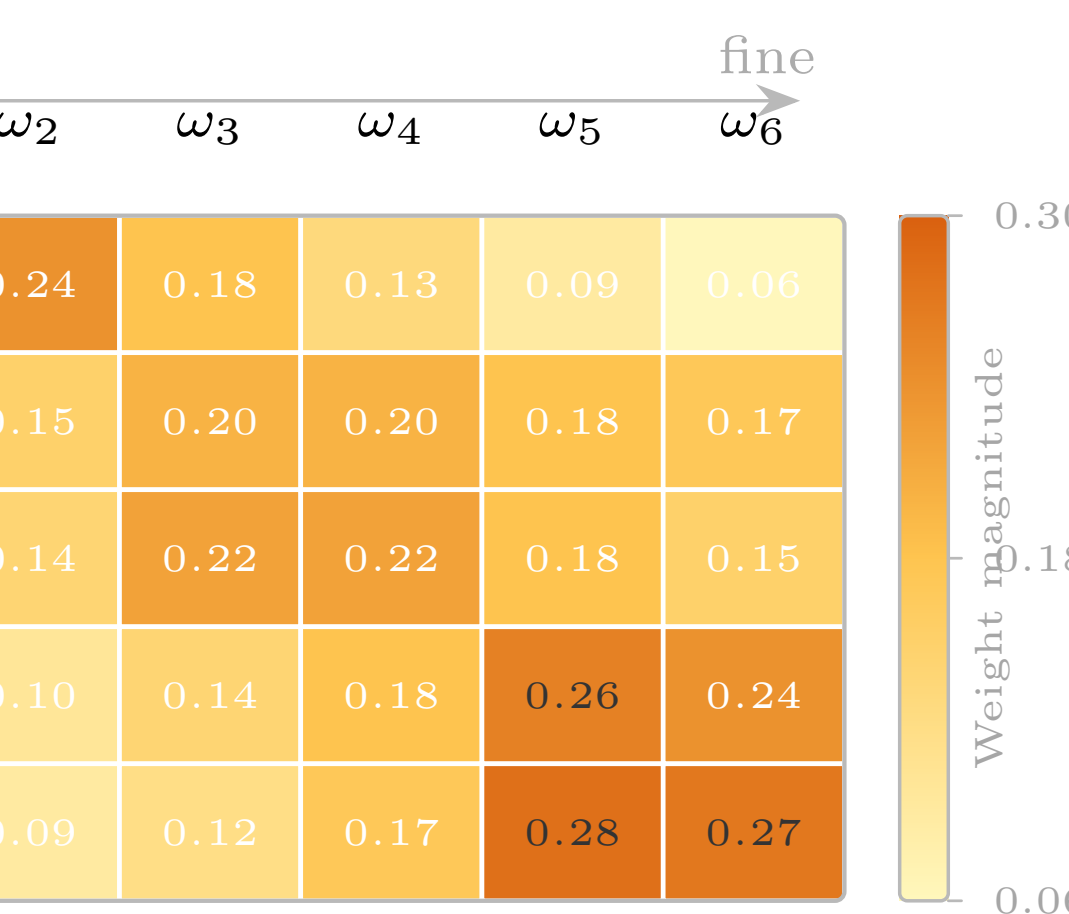


Figure 3. Learned selector weights, averaged over samples from each superclass.



Code



Paper