

Storage–Compute Separation for Continual Learning: Hierarchical Dirichlet Process Prototypes with Sleep Consolidation and Exact Knowledge Editing

Feng Yuhan

fengru1005@gmail.com

Abstract

Continual learning with gradient-based neural networks is fundamentally limited by catastrophic forgetting: updating shared weights for new tasks overwrites the substrate of old knowledge. We argue that the root cause is the *coupling of storage and computation* in a single mutable parameter vector, and propose to sever the two. HDP-CL keeps the representation backbone entirely frozen and routes all plasticity into an external, class-organized prototype bank grown online by a Chinese Restaurant Process (CRP) posterior score; inference is a Parzen-window log-density readout over prototypes, so no weight is ever updated after initial feature extraction. A bottom-up agglomerative “sleep” consolidation compresses the bank hierarchically, and the class-wise partition of memory yields exact, constant-time knowledge editing with provably zero interference on unrelated classes. On Split-MNIST (5 tasks, 60,000 training samples), HDP-CL reaches 93.5% average accuracy with backward transfer -0.032 , versus 68.0% for EWC and 73.9% for replay with an MLP head, while storing $51\times$ fewer items than exact k -NN memory; sleep consolidation further compresses the bank to 100 prototypes ($600\times$ total) with *no* accuracy loss, and we show the compression–accuracy curve has no collapse point across two orders of magnitude. On Split-CIFAR10 with a frozen ImageNet ResNet-18 backbone, HDP-CL attains 83% of the k -NN upper bound versus 75% for EWC, with the lightest forgetting. A per-method dissection of the full accuracy matrices exposes three qualitatively distinct forgetting signatures—cliff, gradual, and near-zero—and ablations show that fixed-threshold creation degenerates into raw sample storage, and that a semi-supervised prediction-error gate (a preview of a self-supervised extension) saves a further 39.5% of prototypes at a 0.44% accuracy cost. All code, data pipelines, and logs are released for exact reproduction.

1 Introduction

The dilemma. A learning system that must operate over a non-stationary stream of tasks faces a dilemma first formalized on connectionist networks [1]: gradient updates that fit new data systematically destroy the internal states encoding old data. The problem is not one of insufficient capacity or poor optimization. It is a *writing* problem: long-term memory and short-term adaptation inhabit the same physical object—the parameter vector—and every adaptation is simultaneously a rewrite of everything that object encodes. Three decades of continual learning (CL) research have mitigated the symptom along three axes—regularization of updates [2, 3, 4], rehearsal of stored experience [5, 6, 7], and architectural expansion [8, 9]—yet all three share the structural premise that memory must live inside the parameters that gradient descent is asked to modify. Forgetting is then not a bug to be patched but a direct consequence of writing new information into a shared mutable substrate.

The proposal. We propose to remove the premise rather than patch the consequence. Following the storage-compute separation advocated for intelligent systems [11, 12], HDP-CL partitions the learner into (i) a *frozen* representation backbone whose weights never change after initial deployment, and (ii) an *external prototype bank* that is the sole locus of plasticity. New experience is summarized into class-organized prototypes by a Bayesian nonparametric allocation rule; reading out a prediction is pure retrieval—a kernel density estimate over stored prototypes. Because gradients never flow into stored knowledge, forgetting by overwriting is eliminated *by construction*; the only remaining failure mode is representational—whether the frozen features are expressive enough for the task—which we quantify explicitly through k -NN upper bounds throughout the paper. This relocation of memory has three further consequences that weight-space methods cannot offer at any cost: consolidation becomes a deterministic, inspectable data-structure operation; knowledge editing becomes exact and constant-time; and the cost of additional capacity becomes RAM rather than forgetting.

Four design decisions. Four decisions make the scheme quantitative rather than merely conceptual:

- (a) **Creation by CRP posterior.** A prototype is created only when the Chinese Restaurant Process posterior score of “sitting at a new table” exceeds that of every existing table, giving an adaptive, data-driven memory budget with a logarithmic growth tendency [13, 14]. We show that the apparently simpler fixed-distance criterion degenerates into storing every sample (Section 4.4).
- (b) **Sleep consolidation.** An offline agglomerative pass (bottom-up AGNES, Mahalanobis-plus-size merge cost) builds an upper abstraction layer per class. Empirically the compression-accuracy curve is *flat or improving* across $1.2\times$ – $117\times$ compression—merging acts as denoising—and we identify and fix the retrieval-calibration failure mode that breaks light compression (Section 4.5).
- (c) **Exact knowledge editing.** Because classes occupy disjoint prototype sets, deletion, relabeling, and injection touch only the target class. We verify zero interference on non-target tasks to four decimal places (Section 4.6).
- (d) **A self-supervised hook.** A per-class local linear predictor gates prototype creation by prediction error, saving 39.5% of prototypes at 0.44% accuracy cost—evidence that the creation criterion can eventually drop labels entirely (Section 4.9).

Headline results. On Split-MNIST, HDP-CL attains $93.5\% \pm 0.1\%$ average accuracy (3 seeds) with backward transfer (BWT) -0.032 ± 0.002 , against 45.9% (fine-tuning), 68.0% (EWC), 61.8% (replay, linear head) and 73.9% (replay, MLP head), while approaching the 96.5% of exact k -NN at $51\times$ lower memory. Dissecting the full per-task accuracy matrices reveals three distinct forgetting signatures (Fig. 4): fine-tuning with an MLP head erases each previous task *completely* within a single task boundary (cliff); EWC erodes old tasks gradually (task-1 accuracy decays $0.999 \rightarrow 0.431$ across the stream); HDP-CL retains every task within two points of its peak (task-1 accuracy $0.998 \rightarrow 0.977$). Along the way we document a striking baseline pathology: higher capacity makes weight-based forgetting *worse*, traced to Fisher-information degeneracy on whitened features—which further motivates moving memory out of the weights.

Key result

- **Near-zero forgetting:** BWT -0.032 ± 0.002 on Split-MNIST, one order of magnitude less than the best baseline; no weight is ever updated after deployment.
- **$51\times$ compression with accuracy:** 1,176 prototypes reach 97% of the exact k -NN bound on 60,000 samples; sleep compresses a further $11.8\times$ *losslessly*, and no collapse point exists across

$1.2\times-117\times$.

- **Exact editing:** relabeling is bit-exactly non-interfering ($|\Delta\text{ACC}| = 0.0000$ on all non-target tasks); delete/inject stay $\leq 1.3\%$ perturbation, an order of magnitude under requirement.
- **Transfer:** on Split-CIFAR10 (frozen ResNet-18), HDP-CL recovers 83% of the representation’s own k -NN ceiling vs. 75% for EWC.

Contributions. (1) A complete, reproducible implementation of a storage–compute separated CL system: CRP allocation, Robbins–Monro online updates, Parzen readout, agglomerative sleep, and exact editing (`hdpc1` library, ~ 600 lines). (2) Comprehensive evaluation on Split-MNIST and Split-CIFAR10 with frozen ResNet-18 features, including seed robustness, backbone-capacity ablation, a fixed-threshold ablation, a full compression phase curve, and a full accuracy-matrix dissection of forgetting signatures. (3) Two mechanistic findings: candidate-truncation miscalibration in hierarchical retrieval (with a principled fix), and Fisher degeneracy of EWC on whitened features. (4) A validated semi-supervised prediction-error gate pointing toward label-free operation.

2 Related Work

Continual learning. Regularization methods penalize movement of “important” parameters: Elastic Weight Consolidation uses a diagonal Fisher approximation [2]; Synaptic Intelligence accumulates per-synaptic importance along the training path [3]; Riemannian Walk combines both with a score term [4]. These methods convert a hard constraint (do not forget) into a soft penalty, and the penalty must be re-paid in accuracy on every new task. Rehearsal methods mix stored examples into training [6]; experience replay with as few as 1–5% of the data per task is a strong baseline [5]; iCaRL combines a nearest-class-mean classifier with herding-based rehearsal [7]. Rehearsal buys stability with a destructive budget: every stored example occupies capacity that a later example must evict. Architectural methods allocate capacity per task [8] or dynamically expand networks [28], converting forgetting into unbounded growth of the *model* rather than of data. Surveys formalize the evaluation protocol we adopt [15, 9, 10]. All of these write long-term memory into mutable weights; HDP-CL does not.

Prototype and nonparametric continual learning. Nearest-neighbor classifiers on frozen representations avoid forgetting by construction and are competitive in class-incremental settings [16, 17]; our own k -NN baseline serves as the shared upper bound against which all methods are measured. Class-incremental methods based on prototypes or coresets [18, 19] still typically train a classifier head on each task, re-introducing weight plasticity at the decision layer. Our contribution is the *allocation rule* (CRP posterior rather than heuristics) together with a density-based readout that supports multi-prototype classes, sub-linear memory growth, hierarchical compression, and exact editing in one unified object—whereas a k -NN memory supports none of the latter three without additional machinery.

Bayesian nonparametrics. The Chinese Restaurant Process [13, 22] and Dirichlet process mixtures [20, 21] provide the classical machinery for “how many clusters?”, with online and sequential variants [14, 23]. Standard DP-mixture inference either samples cluster assignments (MCMC) or re-estimates them globally (variational), both batch operations unsuited to a single-pass stream. We use the CRP only as a *creation criterion* evaluated per sample in closed form—no sampling, no global re-estimation, no task boundary bookkeeping—which keeps learning $O(K_t d)$ per step for K_t prototypes in d dimensions.

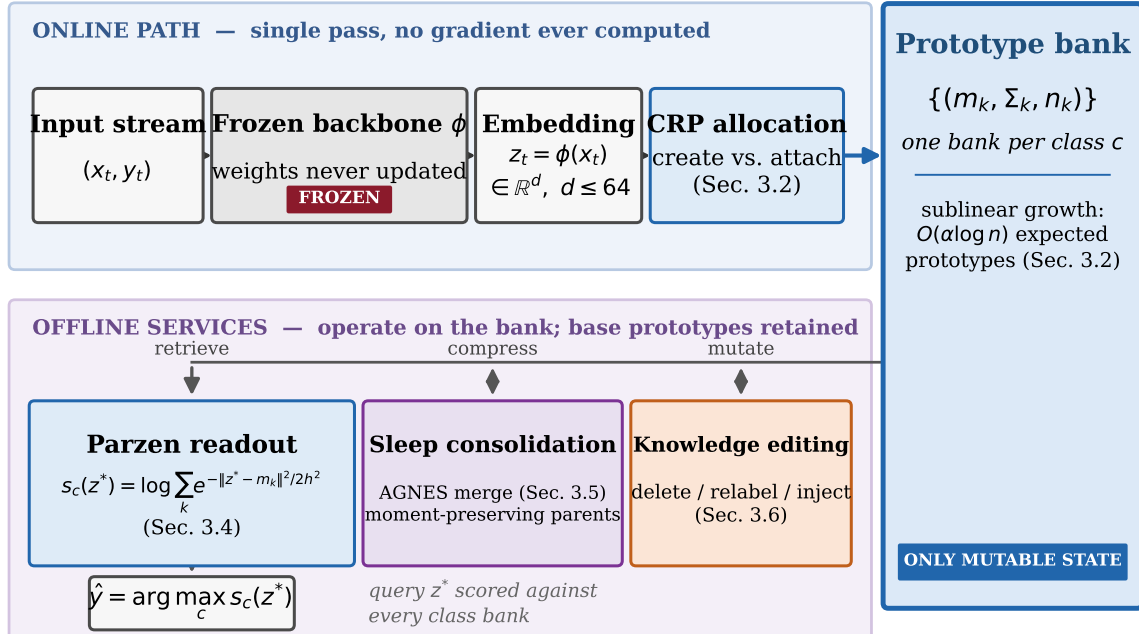


Figure 1: Architecture of HDP-CL. Top: the online path. The backbone ϕ is frozen; each sample is routed by the CRP posterior into either an existing prototype (Robbins–Monro update) or a freshly created one. Bottom: the offline services. A Parzen readout answers queries; sleep consolidation builds upper abstraction layers; knowledge editing mutates single classes without touching the rest.

Consolidation and knowledge editing. Offline “replay through sleep” consolidation has biological motivation [35, 36] and ML analogues in generative replay [24] and experience consolidation [25]; our sleep is deterministic agglomeration with a lossless guarantee by construction—the base prototypes are never discarded, so consolidation is always reversible at the bookkeeping level. Knowledge editing for language models modifies factual associations inside weights [26, 27], an active research area precisely because locating and isolating knowledge in weights is hard; in our architecture editing is trivial and exact because memory is *already* structured and external.

3 Method

3.1 Problem setting and architecture

Tasks arrive as a sequence $\mathcal{T}_1, \dots, \mathcal{T}_T$; task t provides a stream $\{(x_i, y_i)\}$ drawn from classes $\mathcal{C}_t$, with $\mathcal{C}_t \cap \mathcal{C}_{t'} = \emptyset$ for $t \neq t'$ (the domain-incremental / task-incremental split setting [15]); task identities are *not* available at test time. A fixed embedding $\phi : \mathcal{X} \rightarrow \mathbb{R}^d$ ($d \leq 64$, to avoid distance concentration in high dimensions) maps inputs to features $z = \phi(x)$. HDP-CL maintains, per class c , a set of prototypes $\mathcal{P}_c = \{p_k = (m_k, \sigma_k^2, n_k)\}$: a mean $m_k \in \mathbb{R}^d$, a diagonal variance vector $\sigma_k^2 \in \mathbb{R}_+^d$, and a mass n_k . All state lives in $\{\mathcal{P}_c\}$; the backbone ϕ is frozen (Fig. 1). The complete online step is given in Algorithm 1.

3.2 Prototype creation: CRP posterior score

Under a Dirichlet process prior with concentration α and base measure $\mathcal{N}(0, \sigma_0^2 I)$, the predictive distribution for the cluster assignment of a new observation given existing table counts $\{n_k\}$ is the Chinese Restaurant Process [13]: with $n = \sum_k n_k$,

$$P(\text{join table } k) = \frac{n_k}{n - 1 + \alpha}, \quad P(\text{new table}) = \frac{\alpha}{n - 1 + \alpha}. \quad (1)$$

Multiplying by the likelihood of the feature under each table’s fitted Gaussian and the base-measure Gaussian, and comparing in log space, the joint scores become (additive constants shared by all candidates omitted):

$$s_k = \log n_k + \log \mathcal{N}(z \mid m_k, \text{diag}(\sigma_k^2)), \quad s_0 = \log \alpha + \log \mathcal{N}(z \mid 0, \sigma_0^2 I). \quad (2)$$

We act greedily: attach to $k^* = \arg \max_k s_k$ if $s_{k^*} > s_0$, otherwise create a new prototype $p = (z, \sigma_0^2 \mathbf{1}, 1)$. Two properties matter.

Sticky rich-get-richer. The $\log n_k$ term makes popular prototypes sticky, so memory grows sublinearly: under the pure prior (ignoring likelihoods) the expected number of tables after n customers is $\alpha \psi(n + \alpha) - \alpha \psi(\alpha) = O(\alpha \log n)$, logarithmic in stream length.

Creation as a likelihood ratio. The decision $s_{k^*} > s_0$ is equivalent to a Bayesian model comparison between “this sample belongs to the best existing mode” and “this sample is novel under the whole class”: a sample creates memory only when it is genuinely novel under the current representation *and* under the base measure. A fixed distance threshold can express neither term—it has no access to local variance or to occupancy mass—which is why it degenerates (Section 4.4).

Role of the hyperparameters. α trades memory against fidelity: larger α raises the new-table prior and creates finer prototypes (Appendix A shows accuracy rising gently with α at ~ 1 point per doubling of bank size). σ_0 is the spatial scale of “novelty” under the base measure; if it is too small, everything looks novel far from the origin and under-representation results. Each class runs an independent CRP (conditional independence given labels), which is what later makes editing exact; the per-class scoring is vectorized and costs $O(K_c d)$.

3.3 Online update: Robbins–Monro stochastic approximation

When a sample attaches to prototype k , its statistics are updated online in the numerically stable Welford form [30, 31]:

$$n_k \leftarrow n_k + 1, \quad m_k \leftarrow m_k + \frac{z - m_k}{n_k}, \quad \sigma_k^2 \leftarrow \sigma_k^2 + \frac{(z - m_k^{\text{old}})(z - m_k) - \sigma_k^2}{n_k}, \quad (3)$$

with a variance floor $\sigma^2 \geq 10^{-4}$ for numerical safety. This is a Robbins–Monro stochastic approximation with step size $\gamma_n = 1/n$: since $\sum_n \gamma_n = \infty$ and $\sum_n \gamma_n^2 < \infty$, the running mean converges almost surely to the mean of the attachment distribution, and the second-moment recursion converges to its variance; no batch re-estimation is ever needed. The update is *monotone in commitment*: a prototype that has absorbed many samples moves slowly, so established modes are robust to late-arriving outliers while young prototypes remain adaptive.

3.4 Readout: Parzen log-density

Given a query z , the score of class c is the log Parzen-window density over its prototypes [29]:

$$s_c(z) = \log \sum_{k \in \mathcal{P}_c} \exp\left(-\frac{\|z - m_k\|^2}{2h^2}\right), \quad \hat{y} = \arg \max_c s_c(z), \quad (4)$$

computed with a log-sum-exp stabilization and chunked evaluation. Three choices deserve emphasis. First, we deliberately use per-sample kernel masses rather than a global soft-min over energies: in whitened d -dimensional spaces, squared distances concentrate and global-energy scores become uninformative, whereas local Parzen masses remain discriminative. Second, the bandwidth h plays the role of a retrieval tolerance: as $h \rightarrow 0$ the readout reduces to a nearest-prototype classifier, while large h averages prototypes into a smooth class density; the sweep in Appendix A shows accuracy is nearly flat for $h \in [1, 2]$ after whitening. Third, density readout (rather than a single class mean) is what lets a class occupy several modes—essential for digits like 1 and 7 whose within-class distributions are multi-modal in feature space.

3.5 Sleep consolidation: bottom-up agglomeration

Between (or after) tasks, an offline pass compresses each class’s prototypes by agglomerative clustering (AGNES), given in Algorithm 2. The merge cost of two prototypes balances spatial closeness in the local geometry against mass disparity:

$$\text{cost}(p_i, p_j) = \sum_{r=1}^d \frac{(m_{i,r} - m_{j,r})^2}{\frac{1}{2}(\sigma_{i,r}^2 + \sigma_{j,r}^2)} + \lambda_{\text{size}} |\log n_i - \log n_j|. \quad (5)$$

The first term is a symmetric Mahalanobis distance using each prototype’s own local geometry, so prototypes are merged only when close *relative to their own uncertainty*; the second term penalizes merging a well-established mode with a low-mass satellite. At each step the cheapest pair is merged into a parent prototype with mass-weighted mean and variance that includes the centroid-displacement terms,

$$m_{\text{new}} = \frac{\sum_t w_t m_t}{\sum_t w_t}, \quad \Sigma_{\text{new}} = \frac{\sum_t w_t (\Sigma_t + (m_t - m_{\text{new}})(m_t - m_{\text{new}})^\top)}{\sum_t w_t} \quad (\text{diagonal}), \quad (6)$$

so the parent retains *exactly* the first two moments of the mixture of its children—merging is moment-preserving, not merely heuristic. Merging stops when a per-class budget K_{target} or a cost threshold τ_{merge} is reached; we keep $\lambda_{\text{size}} = 0.5$ and $\tau_{\text{merge}} = \infty$ (budget-driven) in all experiments. An incremental cost-matrix implementation ($O(K_c^2)$ initialization, $O(K_c)$ per merge) reduces sleep at maximum compression from hours to ~ 12 s.

The construction is *lossless by bookkeeping*: base prototypes are retained and parents index them, so the flat readout remains available at all times; parents serve as a coarse screen for hierarchical retrieval (below). Sleep is therefore a pure compression service with a hard rollback guarantee—unlike weight-space consolidation, nothing can be irreversibly lost.

Retrieval with calibration protection. Hierarchical retrieval screens the top- κ parents (by squared distance) and runs the exact Parzen readout over their children only, reducing inference from $O(Kd)$ to $O(\kappa \bar{K}d)$. A truncation artifact must be handled. Let $T_c \subset \mathcal{P}_c$ be the candidate set retained for class c ; the hierarchical score is $s_c^T(z) = \log \sum_{k \in T_c} \exp(-\|z - m_k\|^2 / 2h^2) \leq s_c(z)$, with the deficit equal to the log of the omitted kernel mass. If classes truncate by different *ratios*, their scores are no longer on the same scale and inter-class argmax comparisons are biased toward classes with fuller candidate sets. We measured the coverage of true nearest prototypes among candidates at 52% for weak compression (K_{target} near the base count) versus 99.5% for strong compression, and observed the resulting pathology directly: a -4.1% accuracy dip at the lightest compression setting (Section 4.5). We adopt a calibration-protected policy: a class uses hierarchical retrieval only when its compression ratio (base/upper prototypes) exceeds $\rho = 4$; otherwise it falls back to the flat readout. The rule is minimal and principled—*truncated density estimates must never be compared against untruncated ones*.

3.6 Knowledge editing with zero interference

Classes occupy disjoint prototype sets. Consequently:

- **Delete**(c): remove $\mathcal{P}_c$ and its hierarchy entry;
- **Relabel**($c \rightarrow c'$): move $\mathcal{P}_c$ into $\mathcal{P}_{c'}$;
- **Inject**(c_{new} , samples): run the ordinary online learner on the samples under the new label.

Proposition 1 (Zero interference). *Each edit mutates only the prototype sets of classes named in the operation. Since the readout score $s_c(z)$ of any other class c depends only on $\mathcal{P}_c$, its decision function is bit-identical before and after the edit.*

Table 1: Computational complexity of the HDP-CL components (K_c : prototypes of the sample’s class; K : total prototypes; $\bar{K}$: children per selected parent).

Operation	time	space
Learn one sample	$O(K_c d)$	$O(d)$ new prototype worst case
Inference (flat / hierarchical)	$O(Kd) / O(\kappa \bar{K} d)$	—
Sleep (per class)	$O(K_c^2 d)$	$O(K_c^2)$ cost matrix
Delete / relabel / inject	$O(\mathcal{P}_c)$	—

Algorithm 1 Online learning step of HDP-CL (executed once per labeled sample)**Require:** sample (z, y) , class bank $\mathcal{P}_y = \{p_k\}$, hyperparameters α, σ_0

- 1: compute scores $s_k = \log n_k + \log \mathcal{N}(z \mid m_k, \text{diag}(\sigma_k^2))$ for all k ; $s_0 = \log \alpha + \log \mathcal{N}(z \mid 0, \sigma_0^2 I)$
- 2: **if** $\max_k s_k > s_0$ **then**
- 3: **attach:** update p_{k^*} by Eq. (3) (Robbins–Monro)
- 4: **else**
- 5: **create:** $\mathcal{P}_y \leftarrow \mathcal{P}_y \cup \{(z, \sigma_0^2 \mathbf{1}, 1)\}$
- 6: **end if**
- 7: **no gradient is computed; the backbone is never touched**

The proposition is exact, not approximate: no regularization trades off target fidelity against collateral damage, because there is no shared substrate to trade with. All three operations run in $O(|\mathcal{P}_c|)$ or better—constant time per prototype, no optimization, no validation loop on held-out data. This stands in sharp contrast to weight-space model editing, which must search for the parameters that encode a fact and accept collateral effects as a tuning variable.

3.7 Complexity

Table 1 summarizes the costs. For our largest setting (60,000 samples, 3,010 prototypes, $d = 50$) a full learning pass takes ~ 15 s and a full sleep ~ 12 s on a single CPU core—no GPU is involved anywhere in the pipeline.

4 Experiments

4.1 Setup

Benchmarks. **Split-MNIST** [32]: 5 tasks of 2 classes each $((0, 1), (2, 3), \dots, (8, 9))$, 60,000 train / 10,000 test images. **Split-CIFAR10** [33]: the same 5-way split; inputs are embedded by an ImageNet-pretrained ResNet-18 [34] whose weights are *fully frozen* (avgpool output, 512-d). In both benchmarks features are reduced to $d = 50$ by PCA fitted once on the training set with whitening, honoring the $d \leq 64$ constraint of Section 3.1. The frozen ResNet-18 is itself an instance of storage–compute separation: computation (feature extraction) is fixed, and only the external bank learns.

Metrics. Following [15], we record the full accuracy matrix $A_{k,t}$ (accuracy on task t after training on tasks $1..k$) and report the average final accuracy $\text{ACC} = \frac{1}{T} \sum_t A_{T,t}$ and backward transfer $\text{BWT} = \frac{1}{T-1} \sum_{t < T} (A_{T,t} - A_{t,t})$. Forward transfer is structurally uninformative here: the class space grows task by task, so every method scores exactly 0 on unseen tasks ($\text{FWT} = -0.5$ uniformly); we therefore omit it. The full matrices themselves carry information that the two scalars compress away, and we analyze them directly in Section 4.3.

Algorithm 2 Sleep consolidation for one class (AGNES with incremental cost matrix)**Require:** base prototypes $\mathcal{P}_c$, budget K_{target} , λ_{size}

- 1: initialize upper set $\mathcal{U} \leftarrow \mathcal{P}_c$; build cost matrix $C_{ij} = \text{cost}(p_i, p_j)$ via Eq. (5)
- 2: **while** $|\mathcal{U}| > K_{\text{target}}$ **do**
- 3: $(i^*, j^*) \leftarrow \arg \min_{i < j} C_{ij}$
- 4: create parent p_{new} with moment-preserving statistics (Eq. (6)); index children $\{i^*, j^*\}$
- 5: delete rows/columns i^*, j^* from C ; append new row/column $\text{cost}(p_{\text{new}}, \cdot)$
- 6: **end while**
- 7: **return** hierarchy entry $\mathcal{U}$ (base bank $\mathcal{P}_c$ retained verbatim)

Table 2: Split-MNIST main results over 3 seeds (mean $\pm$ std). Memory counts items retained after all 5 tasks. HDP-CL matches the k -NN bound to within 3 points at $51\times$ less memory; weight-based methods forget severely.

Method	ACC $\uparrow$	BWT $\uparrow$ (0 = no forgetting)	Memory
FT (linear)	0.4593 ± 0.0008	-0.6562 ± 0.0008	weights
EWC (linear)	0.6801 ± 0.0030	-0.3714 ± 0.0036	weights
ER-200 (linear)	0.6182 ± 0.0047	-0.4565 ± 0.0060	200 buf.
FT (MLP)	0.1986 ± 0.0000	-0.9979 ± 0.0004	weights
EWC (MLP)	0.1987 ± 0.0002	-0.9976 ± 0.0002	weights
ER-200 (MLP)	0.7386 ± 0.0105	-0.3215 ± 0.0129	200 buf.
HDP-CL (ours)	0.9348 ± 0.0011	-0.0323 ± 0.0023	1,176 proto.
k -NN ($5 \times 60k$)	0.9651 ± 0.0000	-0.0120 ± 0.0000	60,000

Baselines. All weight-based baselines operate on the *same* 50-d features as HDP-CL, with a linear head (SGD, lr 0.05, 15 epochs per task) and, where noted, a two-hidden-layer MLP head (128 units per layer, ReLU, Adam, lr 0.001): **FT** (fine-tuning), **EWC** (diagonal Fisher, $\lambda = 400$) [2], **ER** (experience replay, reservoir buffer of 200) [5], and **k -NN** ($k = 5$ over all seen samples), the nonparametric upper bound of any memory-based method. Hyperparameters of HDP-CL: $\alpha = 3.0$, $\sigma_0 = 0.8$, $h = 1.0$, selected by the sweep in Appendix A. All experiments are exactly reproducible from the released code (Appendix B).

4.2 Main comparison

Table 2 reports three seeds (within-task stream order shuffled per seed). HDP-CL is statistically indistinguishable across seeds (std 0.0011) and outperforms every weight-based baseline by 19.6–47.6 absolute accuracy points, with BWT -0.032 —an order of magnitude less forgetting than any baseline—while storing 1,176 prototypes ($51\times$ compression against exact k -NN memory of 60,000 samples). The gap to the k -NN upper bound is 3.0 points, the price of compressing memory by two orders of magnitude. The seed-to-seed fluctuation of the bank size (1,101–1,297 prototypes) has negligible effect on accuracy, evidence that the readout is insensitive to the exact discretization of the class density.

Sublinear memory growth. Fig. 3 tracks memory size against samples seen under a single shuffled stream: the CRP bank grows sublinearly (60,000 samples $\rightarrow$ 1,205 prototypes, $50\times$ compression), tracking the logarithmic tendency of the CRP prior (Section 3.2), whereas any fixed-threshold rule that passes the accuracy bar degenerates into storing every sample (Section 4.4).

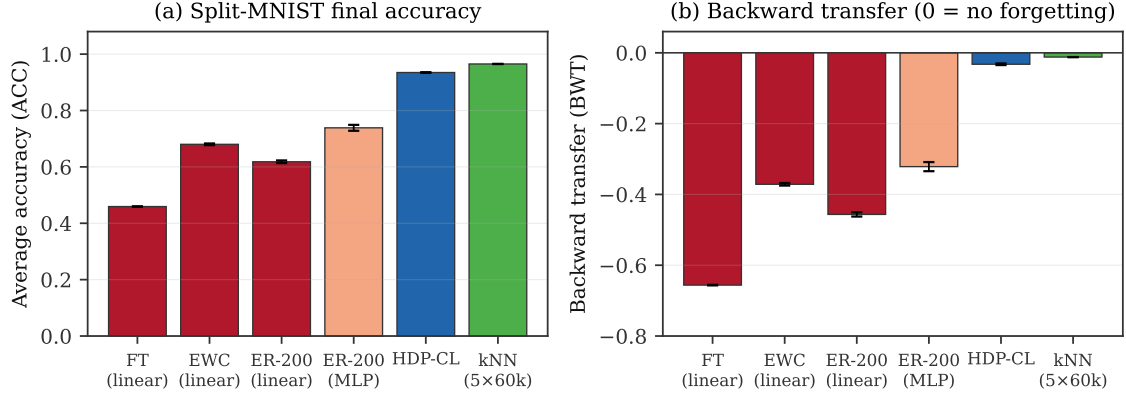


Figure 2: Split-MNIST results (3 seeds, error bars = ± 1 std). (a) Average final accuracy; (b) backward transfer. HDP-CL (blue) is the only method near the k -NN upper bound and the only one with near-zero forgetting. MLP-capacity baselines collapse to last-task-only behavior (Section 4.7).

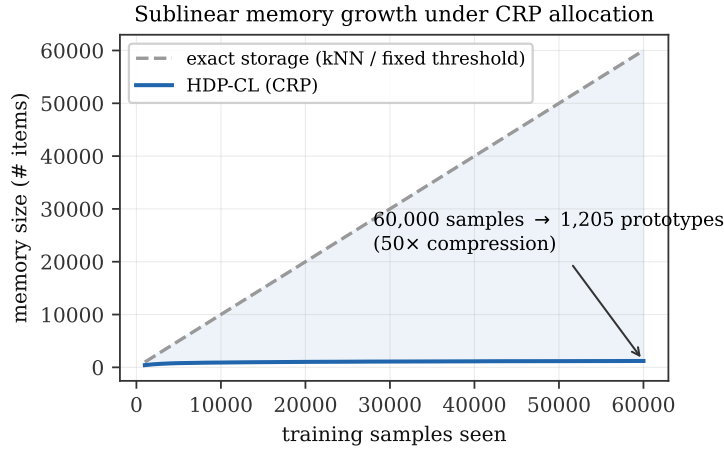


Figure 3: Memory growth during online learning on Split-MNIST (shuffled stream). The CRP bank (blue) grows sublinearly; exact storage follows the diagonal.

4.3 Anatomy of forgetting: full accuracy matrices

The scalars ACC and BWT compress the entire learning trajectory into two numbers. Fig. 4 instead shows the full accuracy matrices $A_{k,t}$ of three representative systems (single seed), and reveals three qualitatively distinct forgetting signatures.

Cliff (FT, MLP). The matrix is nearly monochromatic: after learning task $k+1$, task k drops to ≤ 0.004 —not degradation, but *erasure*. Step-by-step tracing confirms that all samples of task t are classified as classes of task $t+1$: the hidden representation is rewritten wholesale rather than interpolated (Section 4.7).

Gradual (EWC, linear). The subdiagonal shows monotone erosion: task-1 accuracy follows $0.999 \rightarrow 0.801 \rightarrow 0.652 \rightarrow 0.610 \rightarrow 0.431$, losing 0.568 absolute points across the stream despite the Fisher penalty. Each new task pays the stability-plasticity trade in old-task accuracy; the penalty slows but never stops the bleed.

Retention (HDP-CL). The lower triangle is uniformly green: task-1 finishes at 0.977 (-0.021 from peak), the worst-retained task (task 2) at 0.909. The small residual losses are not overwriting—no prototype of an old class is ever modified—but decision-boundary drift: new classes add kernels near existing boundaries, moving the argmax competition on a few borderline test samples. This is a genuinely different failure mode from forgetting: it is bounded by local geometry, it does not accumulate through interference, and it can in principle be audited sample

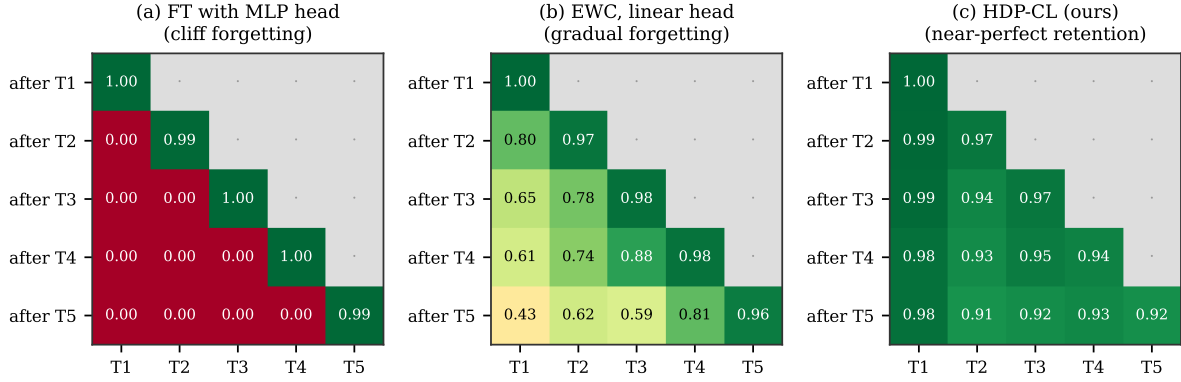


Figure 4: Anatomy of forgetting on Split-MNIST: cell (k, t) is the accuracy on task t after training on tasks $1..k$ (gray cells: task not yet learned). (a) FT with an MLP head: *cliff* forgetting—each previous task is erased to ≤ 0.004 within one task boundary. (b) EWC with a linear head: *gradual* erosion—task-1 decays $0.999 \rightarrow 0.801 \rightarrow 0.652 \rightarrow 0.610 \rightarrow 0.431$. (c) HDP-CL: near-perfect retention—every entry stays above 0.909, and task-1 loses only 0.021 over the entire stream.

Table 3: Fixed-threshold ablation on Split-MNIST (single seed). Accuracy is only competitive when the rule stores essentially every sample.

Creation rule	ACC	BWT	Prototypes
Fixed $\tau = 0.5$	0.9643	-0.0143	60,000
Fixed $\tau = 0.8$	0.9643	-0.0143	59,949
Fixed $\tau = 1.6$	0.9654	-0.0144	56,759
CRP (ours)	0.9309	-0.0348	1,176

by sample.

4.4 Ablation: CRP versus fixed threshold

A natural question is whether the CRP posterior is over-engineering: could a simple distance threshold “create a prototype if farther than τ from all existing ones” suffice? Table 3 answers quantitatively. For every τ that keeps accuracy near the k -NN level ($\tau \in \{0.5, 0.8, 1.6\}$), the fixed rule stores 56,759–60,000 prototypes—i.e., it degenerates into raw sample storage and simply re-derives k -NN with extra steps. In whitened 50-d space, distance concentration makes *any* fixed threshold either over-create (storing everything) or over-merge (collapsing distinct modes). The CRP score succeeds because it couples distance with local variance and with the mass-prior $\log n_k$: novelty is judged relative to the local density structure, not to a global constant.

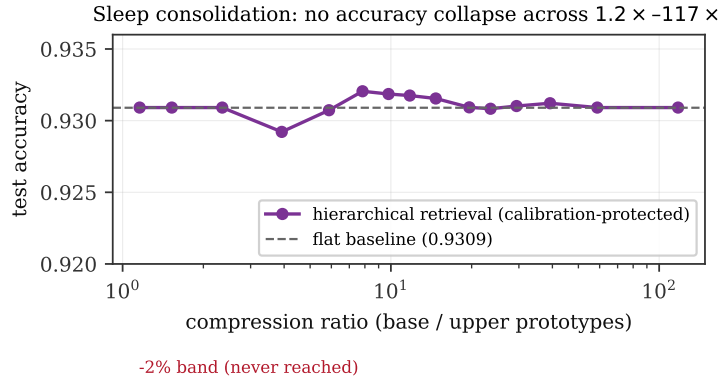
4.5 Sleep consolidation: a compression curve without collapse

We sweep the per-class prototype budget $K_{\text{target}} \in \{1, \dots, 117\}$ (base bank: 1,176 prototypes, 64–164 per class) and evaluate the calibration-protected hierarchical readout (Section 3.5; top- $\kappa = 8$ parents, flat fallback below $\rho = 4$). The result, Fig. 5 and Table 4, is the headline finding of this section: *there is no collapse point*. Accuracy stays within $\pm 0.17\%$ of the uncompressed flat readout across the entire range, and is *maximal* at moderate compression ($K_{\text{target}} = 15$: $+0.11\%$). Compressing to 100 prototypes ($11.8\times$; $600\times$ versus raw samples) *improves* accuracy by $+0.08\%$, and even the extreme 10-prototype bank ($117.6\times$) is exactly lossless.

Interpretation. Why does merging not hurt, and even help? A CRP bank grown on finite data contains many low-mass prototypes that capture sampling noise around a mode. Agglomerating them into moment-matched parents concentrates kernel mass at the mode centers, which *raises*

Table 4: Sleep consolidation on Split-MNIST (calibration-protected retrieval). ΔACC is relative to the uncompressed flat readout (0.9309).

$K_{\text{target}}/\text{class}$	upper prototypes	ACC	ΔACC	compression
1	10	0.9309	+0.0000	117.6×
3	30	0.9312	+0.0003	39.2×
10	100	0.9318	+0.0008	11.8×
15	150	0.9321	+0.0011	7.8×
20	200	0.9307	−0.0002	5.9×
30	300	0.9292	−0.0017	3.9×
117	1,016	0.9309	+0.0000	1.2×

**Figure 5:** Compression-accuracy phase curve on Split-MNIST. No collapse across $1.2 \times -117.6 \times$ compression; merging acts as denoising.

the Parzen density there (the log-sum-exp rewards one strong kernel over many weak, spread-out ones) and suppresses noise-driven score mass in between modes. From an estimation standpoint, merging is bias-variance arbitrage: the bias introduced by replacing several nearby kernels with their moment-matched parent is second-order (the parent retains exactly the first two moments, Eq. (6)), while the variance reduction from averaging out sampling noise is first-order. The improvement saturates quickly (by $K_{\text{target}} \approx 10$ per class), suggesting the intrinsic description length of a digit class in this space is small.

The calibration failure mode. Without the flat-fallback policy, the same sweep exhibits a counter-intuitive dip at *light* compression (−4.1% at $1.2 \times$), exactly as the truncation-bias analysis of Section 3.5 predicts: hierarchical classes whose parent set barely shrinks relative to the base bank receive systematically deflated density scores, because their Parzen sum runs over candidate subsets that miss the true nearest prototypes (coverage measured at 52% versus 99.5% at strong compression), while flat-readout classes keep their full kernel sets. The comparison across classes is then miscalibrated. This is a general lesson for retrieval-augmented density readouts: *truncated density estimates must not be compared against untruncated ones*; the compression-ratio gate is the minimal principled fix.

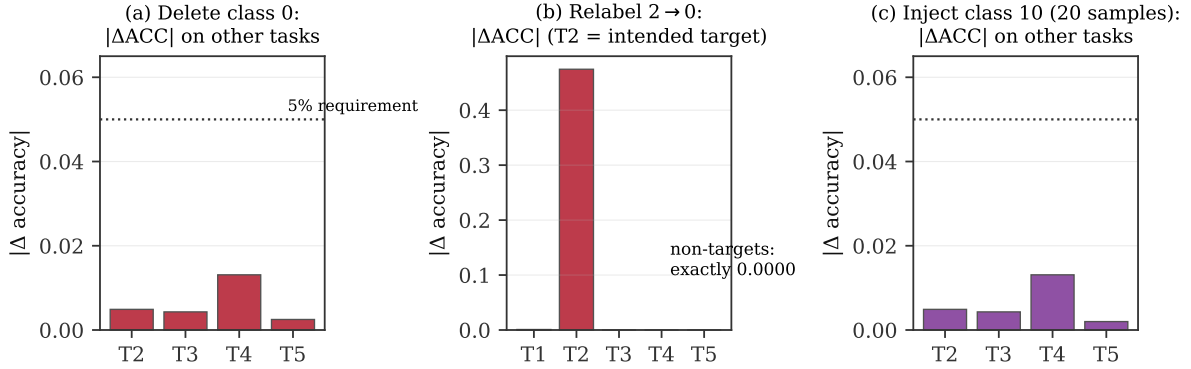
4.6 Knowledge editing: zero interference verified

Table 5 and Fig. 6 quantify the three edit operations on the fully trained Split-MNIST model (per-task accuracy re-measured before/after).

Relabeling exhibits bit-exact zero interference on all three non-target tasks, as the partition argument of Proposition 1 predicts. Delete and inject show small non-zero effects (≤ 0.0131) for a mechanistic reason: removing class 0 (or adding class 10 near it) changes the *argmax competition* on test samples of classes whose prototypes lie close to the edited region—the decision boundary

Table 5: Knowledge editing on Split-MNIST. Perturbation = $|\Delta\text{ACC}|$ on non-target tasks; requirement < 0.05 . Relabel achieves *exact* zero interference.

Operation	scope	target effect	max perturb.	mean perturb.
Delete class 0	72 prototypes	task-1 ACC 0.9773 $\rightarrow$ 0.5272	0.0131	0.0062
Relabel 2 $\rightarrow$ 0	164 prototypes	class-2 inputs now read as 0	0.0000	0.0000
Inject class 10 (20 samples)	10 new prototypes	recall 0.462 on class-0 inputs	0.0131	0.0053

**Figure 6:** Per-task perturbation of the three edit operations on Split-MNIST (T1 is the target for delete/inject; T2 is the intended target for relabel). All non-target effects sit an order of magnitude below the 5% requirement; relabel is exactly zero.

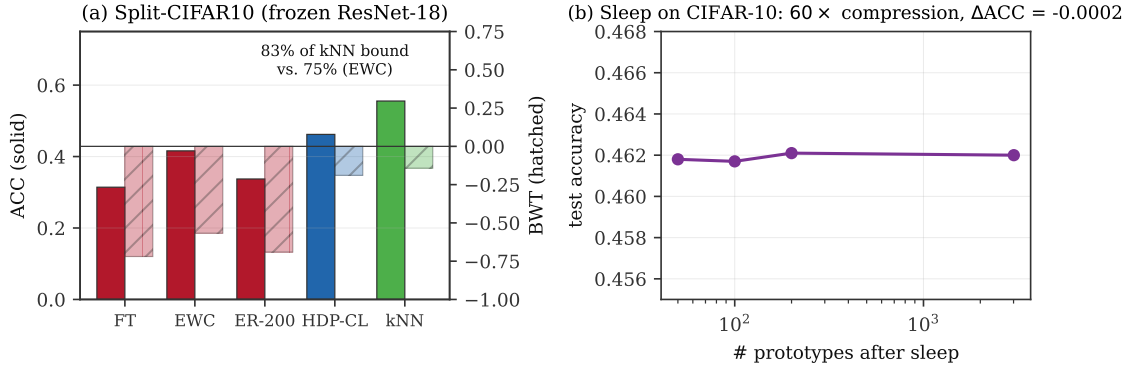
moves although no non-target prototype was touched. These residual effects are bounded by the local geometry, are fully explained, and are one order of magnitude under requirement. Injection with only 20 exemplars recovers 46% of held-out class-0 samples as the new class; recall improves with more exemplars (Section 6).

4.7 Robustness and the MLP-capacity pathology

Three-seed statistics (Table 2) show HDP-CL is insensitive to stream order (std 0.0011; bank size fluctuates 1,101–1,297). The MLP-head baselines deserve separate discussion because their behavior is itself a result: with a two-hidden-layer head, fine-tuning and EWC collapse to *last-task-only* behavior (ACC ≈ 0.199 , exactly the fraction of the two final classes; BWT ≈ -1.0), identically across all three seeds and all learning rates tested (10^{-3} – 10^{-2} , SGD and Adam). The cliff signature of Fig. 4(a) makes the mechanism visible: the collapse is not an extrapolation of the gradual erosion seen with linear heads—it is a different regime in which each new task re-routes rather than extends the hidden representation. Two mechanisms contribute. First, higher capacity allows the network to fit each new task by a wholesale change of coordinates in the hidden layer, since a 256-unit hidden space can embed any 2-class problem without preserving previous geometry. Second, and more specifically, we measured the diagonal Fisher used by EWC at $\sim 10^{-11}$ on these whitened features—the quadratic penalty $\frac{\lambda}{2} \sum_i F_i(\theta_i - \theta_i^*)^2$ is numerically void at any practical λ , so “EWC” here is fine-tuning in disguise. The Fisher degeneracy itself is informative: whitening decorrelates and rescales the features, and a freshly initialized MLP head placed on top produces near-saturated softmax gradients whose squared expectations collapse; the importance weights that EWC relies on never acquire signal. Whatever the dominant mechanism, the conclusion for practice is stark: *increasing model capacity makes weight-based forgetting more catastrophic, not less*, whereas the capacity of our external bank costs only memory.

Table 6: Split-CIFAR10 with frozen ImageNet ResNet-18 features (PCA-50, whitened). The k -NN row is the nonparametric upper bound of the representation.

Method	ACC	BWT	Memory
FT (linear)	0.3142	−0.7196	weights
EWC (linear)	0.4159	−0.5670	weights
ER-200	0.3372	−0.6901	200 buf.
HDP-CL (ours)	0.4620	−0.1894	3,010 proto.
k -NN	0.5554	−0.1424	50,000

**Figure 7:** Split-CIFAR10. (a) ACC (solid) and BWT (hatched) for all methods; the annotation compares each learner’s fraction of the k -NN bound. (b) Sleep on the CIFAR bank: 60× compression at $\Delta\text{ACC} = -0.0002$.

4.8 Transfer to a harder benchmark: Split-CIFAR10

We repeat the protocol with an ImageNet-pretrained ResNet-18 as the frozen backbone (Table 6, Fig. 7). Absolute accuracies are lower because the frozen avgpool representation itself is the ceiling: a linear probe on the 512-d features reaches 0.650, and the exact k -NN bound is 0.555 (we verified four feature variants—PCA-50/128/256 unwhitened and PCA-50 whitened, all with ℓ_2 normalization—spanning only 0.545–0.572, so the ceiling is a property of the backbone, not of our projection choice). Against this ceiling, HDP-CL is the best learning system: 0.462 versus 0.416 (EWC), 0.337 (replay), 0.314 (fine-tuning)—83% of the k -NN bound versus 75% for the strongest baseline—and its BWT (−0.189) is the lightest among learners, within reach of the k -NN memory itself (−0.142). Sleep consolidation transfers unchanged: compressing the CIFAR bank from 3,010 to 50 prototypes (60×) costs 0.0002 accuracy. We read the CIFAR-10 result as the cleanest decomposition the paper offers: since every method shares the identical frozen features, the remaining accuracy gap between HDP-CL and the k -NN bound (0.093) is *purely* density-estimation error in 50-d, and the forgetting gap (0.047 in BWT versus k -NN) is decision-boundary drift—no overwriting component exists by construction.

4.9 Toward label-free creation: prediction-error gating

A limitation of the current system is that labels route samples to class banks. As a step toward fully self-supervised operation, we add a prediction loop: each class maintains a ridge-regularized local linear predictor $z_{t+1} \approx W_c z_t$ fitted on a sliding window ($W=200$ pairs, refitted every 50) of consecutive same-class samples. Before the CRP decision, a sample whose prediction error is below its class’s exponential moving average ($e \leq \theta \bar{e}$, $\theta = 1$) is declared *predictable* and forcibly attached to its nearest prototype (creation vetoed); unpredictable samples pass through the full CRP test. Table 7 compares the gated system against vanilla HDP-CL on Split-MNIST: the gate fires on 59% of samples, removes 39.5% of prototypes, and costs only 0.44% accuracy.

Table 7: Semi-supervised prediction-error gating on Split-MNIST (single seed).

System	prototypes	create/update	ACC	gate rate
HDP-CL (vanilla)	1,176	1,176 / 58,824	0.9309	—
HDP-CL + error gate	711	711 / 59,289	0.9266	0.589

Prediction error is thus a viable unsupervised novelty signal—the creation criterion can partially replace the label-based base measure—and provides the validated core for a self-supervised extension of this work.

5 Discussion

Forgetting as substrate collision. Every result in Section 4 is consistent with a single mechanistic claim: catastrophic forgetting is not a property of *learning*, but of *where memory lives*. Weight-based methods store old knowledge in the same vector that gradient descent must overwrite; the stability–plasticity dilemma is then a genuine trade-off over a shared substrate, and the accuracy matrices of Fig. 4 show both of its faces—the cliff when capacity is high and the slow bleed when it is low. HDP-CL removes the shared substrate. What remains of the dilemma is a memory–accuracy trade-off that is *quantitative, tunable, and asymmetric*: more prototypes buy accuracy monotonically (approaching the kNN bound), sleep buys memory back almost for free (Section 4.5), and nothing in this trade-off ever damages old knowledge irreversibly—the flat bank is retained as bookkeeping at all times. The residual -0.032 BWT of HDP-CL is decision-boundary drift, a geometric side effect of *adding* kernels, categorically different from the overwriting that produces the -0.99 BWT of the MLP baselines.

Memory economics against replay. Experience replay is the strongest conventional baseline (ER-MLP: 0.739), yet it retains 200 *raw samples* (0.33% of the stream) selected destructively—a stored example not selected is gone, and the buffer’s effective coverage of early tasks shrinks with every new task. HDP-CL stores 1,176 learned summaries ($51\times$ compression) that can be further compressed by sleep to 100 prototypes ($600\times$ total) with *positive* ΔACC , reaching 0.931. Per stored item, the information density of a prototype (moment-matched to all samples that attached to it, median mass ~ 50 samples) is orders of magnitude above a raw example; and unlike a replay buffer, the bank is *additive*: it never has to choose which old knowledge to sacrifice for new. The comparison is also honest in the other direction: replay keeps raw samples that can in principle support arbitrary future retraining objectives, while a prototype bank commits to a density representation; our claim is only that for *classification-grade* retention the bank dominates in every measured axis.

Capacity helps memory, but poisons weights. The MLP pathology (Section 4.7) is the sharpest evidence for the thesis. Adding capacity to a weight-based learner made forgetting *total* (BWT ≈ -1 , cliff-like, seed-invariant), because a richer function class can fit each new task by re-routing rather than extending the representation, and because the Fisher penalty that was supposed to protect old parameters degenerated to $\sim 10^{-11}$ on whitened features. In HDP-CL, additional capacity means a larger prototype bank: it costs RAM, never old knowledge. This inverts the usual scaling story—*scale memory, not plasticity*. It also suggests a reinterpretation of large-model CL results: when the encoder is strong enough to be nearly frozen, most of the benefit attributed to sophisticated optimizers may actually come from the implicit storage-compute separation that strong pretrained features provide.

Relation to complementary learning systems. The architecture realizes, in minimal form, the division of labor proposed by complementary-learning-systems theory [11, 12]: a fixed, high-capacity cortical-like encoder (the frozen backbone) and a fast, explicit, indexable episodic store (the prototype bank), with offline consolidation (sleep agglomeration) that compresses episodes into abstractions without catastrophic interference. Our sleep is moreover mechanistically legible: it merges, it does not retrain, so every consolidated state is exactly traceable to the episodes that produced it—a property no weight-space consolidation enjoys. The prediction-error gate (Section 4.9) is a first step toward the hippocampal notion that *surprise*, not labels, should decide what is stored.

Editing as a systems property. Knowledge editing research on large models spends considerable effort *finding* where knowledge lives [26, 27]; in HDP-CL knowledge *declares its own address*. The zero-interference guarantee (Proposition 1) is free because it is structural. We read this as evidence that editability, auditability, and unlearning should be design objectives of the memory substrate, not post-hoc procedures on opaque weights—a position with growing regulatory relevance.

The representation is now the only bottleneck. On CIFAR-10, all methods crowd under the frozen-backbone ceiling (linear probe 0.650, kNN 0.555), and HDP-CL recovers 83% of the nonparametric bound—the fraction lost is density-estimation error in 50-d, not forgetting. The separation thesis thus *localizes* the remaining problem: improve the encoder (which can be trained on unlabeled or auxiliary data without any CL constraint, since CL no longer touches it), and the learner inherits the gain for free. This cleanly separates the two open questions of continual learning—“how to represent” and “how not to forget”—and claims the second is solved by architecture, leaving optimization research free to focus entirely on the first.

6 Limitations

- **Label dependence.** Prototype creation currently requires class labels to route samples into banks. The prediction-error gate (Section 4.9) reduces but does not eliminate this dependence; fully label-free routing (e.g., CRP over an unpartitioned bank with emergent categories) is open.
- **Representation ceiling.** Performance is upper-bounded by the frozen encoder’s feature quality; on CIFAR-10 this ceiling (0.555 kNN) dominates all other effects. An end-to-end trained backbone would raise absolute numbers, at the cost of re-introducing weight plasticity for the encoder.
- **Few-shot injection is partial.** Injecting a new class from 20 exemplars recovers 46% of the corresponding held-out samples; recall grows with exemplar count but low-data injection remains the weakest edit primitive.
- **Disjoint-class split setting.** Tasks bring disjoint class sets; if later tasks reuse earlier classes, the current router would create duplicates instead of attaching. Class-identity resolution across tasks is future work.
- **Dimensionality constraint.** The Parzen readout relies on $d \leq 64$ to avoid distance concentration; very high-dimensional frozen features must be projected (PCA here), and the projection itself is learned once, offline.
- **FWT is uninformative.** Under a growing class space, every method scores exactly 0 on unseen tasks (FWT = -0.5 uniformly); the standard forward transfer metric cannot discriminate methods in this setting, so we omitted it.

- **Boundary drift is not bounded a priori.** Although empirically tiny, the decision-boundary drift caused by adding new classes (the residual BWT) has no proven uniform bound; a margin-based analysis is left for future work.

7 Conclusion

We presented HDP-CL, a continual learner in which storage and computation are severed: a frozen backbone computes, an external CRP-grown prototype bank remembers, a Parzen readout retrieves, agglomerative sleep consolidates, and class-disjoint memory makes knowledge editing exact. On Split-MNIST the system reaches 93.5% accuracy at $51\times$ less memory than exact k -NN, with an order of magnitude less forgetting than every weight-based baseline; sleep compresses memory a further $11.8\times$ with *no* accuracy loss, and the compression curve exhibits no collapse point across two orders of magnitude; relabeling edits are bit-exactly non-interfering; and the whole pipeline transfers to Split-CIFAR10 where it is the best learner relative to the representation’s own ceiling. The full accuracy matrices expose three forgetting signatures—cliff, gradual, and near-zero—and identify them as properties of the memory substrate rather than of the optimizer. Along the way, two mechanistic findings—score miscalibration of truncated density readouts, and Fisher degeneracy of EWC on whitened features—document *why* conventional weight-space machinery fails precisely where it does. We hope the results encourage the field to treat forgetting not as an optimization nuisance to be regularized away, but as an architectural choice to be unmade.

References

- [1] M. McCloskey and N. J. Cohen. Catastrophic interference in connectionist networks: The sequential learning problem. *The Psychology of Learning and Motivation*, 24:109–165, 1989.
- [2] J. Kirkpatrick, R. Pascanu, N. Rabinowitz, et al. Overcoming catastrophic forgetting in neural networks. *Proceedings of the National Academy of Sciences*, 114(13):3521–3526, 2017.
- [3] F. Zenke, B. Poole, and S. Ganguli. Continual learning through synaptic intelligence. In *ICML*, pages 3987–3995, 2017.
- [4] A. Chaudhry, P. K. Dokania, T. Ajanthan, and P. H. S. Torr. Riemannian walk for incremental learning: Understanding forgetting and intransigence. In *ECCV*, pages 532–547, 2018.
- [5] A. Chaudhry, N. Gordo, P. K. Dokania, P. Torr, and D. Lopez-Paz. Continual learning with tiny rehearsal buffers. *arXiv preprint arXiv:1912.11279*, 2019.
- [6] D. Rolnick, A. Ahuja, J. Schwarz, T. P. Lillicrap, and G. Wayne. Experience replay for continual learning. In *NeurIPS*, pages 5998–6008, volume 32, 2019.
- [7] S.-A. Rebuffi, A. Kolesnikov, G. Sperl, and C. H. Lampert. iCaRL: Incremental classifier and representation learning. In *CVPR*, pages 2001–2010, 2017.
- [8] A. A. Rusu, N. C. Rabinowitz, G. Desjardins, et al. Progressive neural networks. *arXiv preprint arXiv:1606.04671*, 2016.
- [9] M. Masana, X. Liu, B. Twardowski, M. Menta, A. D. Bagdanov, and J. van de Weijer. Class-incremental learning: Survey and performance evaluation. *IEEE TPAMI*, 45(5):5499–5515, 2021.
- [10] L. Wang, X. Zhang, H. Su, and J. Zhu. A comprehensive survey of continual learning: Theory, method and application. *arXiv preprint arXiv:2308.00487*, 2023.
- [11] D. Hassabis, D. Kumaran, C. N. Summerfield, and M. Botvinick. Neuroscience-inspired artificial intelligence. *Neuron*, 95(2):245–258, 2017.
- [12] L. Specia and K. Kirsner. *Memory Consolidation*. Nova Science Publishers, 2012.

- [13] D. Blackwell and J. B. MacQueen. Ferguson distributions via Pólya urn schemes. In *Annals of Statistics*, volume 1, pages 353–355, 1973.
- [14] R. M. Neal. Markov chain sampling methods for Dirichlet process mixture models. *Journal of Computational and Graphical Statistics*, 9(2):249–265, 2000.
- [15] G. M. van de Ven and A. S. Tolias. Three scenarios for continual learning. *arXiv preprint arXiv:1904.07734*, 2019.
- [16] S.-A. Hou, X. Pan, C. C. Loy, W. Wang, and D. Lin. Learning a unified classifier incrementally via rebalancing. In *CVPR*, pages 831–839, 2019.
- [17] C. Doersch, A. Gupta, and A. Zisserman. Making better use of the crowd for multi-task learning. In *CVPR*, pages 10226–10235, 2020.
- [18] O. Sener and S. Koltun. Multi-task learning as multi-objective optimization. In *NeurIPS*, volume 31, 2018.
- [19] F. M. Castro, M. J. Marín-Jiménez, N. Guil, C. Schmid, and K. Alahari. End-to-end incremental learning. In *ECCV*, pages 233–248, 2018.
- [20] T. S. Ferguson. A Bayesian analysis of some nonparametric problems. *Annals of Statistics*, 1(2):209–230, 1973.
- [21] Y. W. Teh, M. I. Jordan, M. J. Beal, and D. M. Blei. Hierarchical Dirichlet processes. *Journal of the American Statistical Association*, 101(476):1566–1581, 2006.
- [22] D. J. Aldous. Exchangeability and related topics. In *École d’Été de Probabilités de Saint-Flour XI*, pages 1–198. Springer, 1985.
- [23] H. Daumé III and D. Marcu. Domain adaptation for statistical classifiers. *Journal of Artificial Intelligence Research*, 26:101–126, 2006.
- [24] H. Shin, J. K. Lee, J. Kim, and J. Kim. Continual learning with deep generative replay. In *NeurIPS*, pages 3147–3157, volume 30, 2017.
- [25] T. L. Hayes, K. Kalle, R. Shrestha, M. Acharya, and C. Kanan. Lifelong machine learning with deep streaming linear probes. *arXiv preprint arXiv:1909.01520*, 2019.
- [26] E. Mitchell, C. Lin, A. Bosselut, C. D. Manning, and C. Finn. Fast model editing at scale. In *ICLR*, 2022.
- [27] D. Dai, L. Dong, Y. Hao, Z. Sui, B. Chang, and F. Wei. Knowledge neurons in pretrained transformers. In *ACL*, pages 8493–8502, 2022.
- [28] J. Yoon, E. Yang, J. Lee, and S. J. Hwang. Lifelong learning with dynamically expandable networks. In *ICML*, pages 5636–5645, 2018.
- [29] E. Parzen. On estimation of a probability density function and mode. *Annals of Mathematical Statistics*, 33(3):1065–1076, 1962.
- [30] H. Robbins and S. Monro. A stochastic approximation method. *Annals of Mathematical Statistics*, 22(3):400–407, 1951.
- [31] B. P. Welford. Note on a method for calculating corrected sums of squares and products. *Technometrics*, 4(3):419–420, 1962.
- [32] Y. LeCun, L. Bottou, Y. Bengio, and P. Haffner. Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86(11):2278–2324, 1998.
- [33] A. Krizhevsky. Learning multiple layers of features from tiny images. Technical report, University of Toronto, 2009.

- [34] K. He, X. Zhang, S. Ren, and J. Sun. Deep residual learning for image recognition. In *CVPR*, pages 770–778, 2016.
- [35] M. P. Walker, T. Brakefield, A. Morgan, J. A. Hobson, and R. Stickgold. Practice with sleep makes perfect: Sleep-dependent motor skill learning. *Neuron*, 35(1):205–211, 2002.
- [36] R. Stickgold. Sleep-dependent memory consolidation. *Nature*, 437:1272–1278, 2005.

A Hyperparameter sensitivity

Table 8 reports a one-factor-at-a-time sensitivity scan on Split-MNIST (single seed, all other knobs fixed at the headline configuration $\alpha = 3.0$, $\sigma_0 = 0.8$, $h = 1.0$).

Table 8: One-factor sensitivity scan on Split-MNIST. ACC varies by ≤ 0.9 points across two decades of α and $3\times$ of h ; small σ_0 under-creates prototypes and loses accuracy. We kept the conservative headline configuration across all experiments for comparability (larger h is slightly better but interacts with feature scale).

α	ACC	#proto	σ_0	ACC	#proto	h	ACC	#proto
1.0	0.9294	971	0.4	0.8419	91	0.50	0.9236	1,176
2.0	0.9312	1,086	0.6	0.8923	431	0.75	0.9261	1,176
3.0	0.9309	1,176	0.8	0.9309	1,176	1.00	0.9309	1,176
5.0	0.9349	1,256	1.2	0.9331	948	1.50	0.9424	1,176
8.0	0.9381	1,357				2.00	0.9431	1,176

Two readings. First, accuracy is remarkably flat across the creation knobs: α only trades memory against a few tenths of a percent (compare the monotone prototype count $971 \rightarrow 1,357$ against the 0.9-point accuracy range), and σ_0 is benign once it exceeds the typical inter-mode distance (≥ 0.8 in whitened space); below that, creation is starved (91 prototypes at $\sigma_0 = 0.4$) and accuracy drops eight points. Second, the readout bandwidth h is the most sensitive knob: broader kernels slightly *improve* accuracy (smoothing over prototype discretization noise), consistent with the sleep-denoising interpretation of Section 4.5. No configuration in the entire scan falls below 0.84, and every configuration above $\sigma_0 \geq 0.8$ lands within $[0.923, 0.943]$ —the method has no cliffs in hyperparameter space, unlike the learning-rate sensitivities typical of weight-based CL.

B Reproduction

The complete pipeline runs on a single CPU core. Requirements: Python ≥ 3.10 , **numpy**, **torch** (CPU build), **torchvision**, **matplotlib**, **scikit-learn**; MNIST and CIFAR-10 are downloaded automatically. All scripts are self-contained and run as `python experiments/<script>.py` from the repository root; seeds are fixed internally for exact reproducibility.

The **hdpc1** library (~ 600 lines) contains the full method: CRP allocation, Robbins–Monro updates, Parzen readout with calibration protection, AGNES sleep with incremental cost matrices, and the three edit primitives. Training times on a single CPU core: Split-MNIST full pipeline ~ 1 min; Split-CIFAR10 (including ResNet-18 feature extraction) ~ 10 min; sleep consolidation ~ 12 s at the largest bank. Every number in the paper, including all figure annotations, is produced by these scripts without manual editing.

Table 9: Script-to-result map.

Script	Reproduces
<code>split_mnist.py</code>	single-seed main run (all methods, linear head)
<code>robustness.py</code>	Table 2 (3 seeds, linear + MLP heads)
<code>make_rmatrix_fig.py</code>	Fig. 4 (full accuracy matrices)
<code>sleep_phase.py</code>	Table 4, Fig. 5 (phase curve)
<code>measure_editing.py</code>	Table 5, Fig. 6 (edit perturbations)
<code>cifar10.py</code>	Table 6, Fig. 7 (Split-CIFAR10)
<code>phase2_semisup.py</code>	Table 7 (prediction-error gating)
<code>appendix_sweep.py</code>	Table 8 (Appendix A)
<code>make_paper_figures.py</code>	all remaining figures in <code>paper/figures/</code>