
Seemingly Redundant Modules Enhance Robust Odor Learning in Fruit Flies

Haiyang Li^{1*} Liao Yu^{2*} Qiang Yu^{3†} Yunliang Zang^{1,4†}

¹Academy of Medical Engineering and Translational Medicine, Tianjin University, Tianjin, China

²School of Mathematical Sciences, Beihang University, Beijing, China

³School of Artificial Intelligence, Tianjin University, Tianjin, China

⁴Xiamen Intretech Inc, Xiamen, Fujian, China

haiyangli@tju.edu.cn

yuliao_16@buaa.edu.cn

yuqiang@tju.edu.cn

yunliangzang@tju.edu.cn

Abstract

Biological circuits have evolved to incorporate multiple modules that perform similar functions. In the fly olfactory circuit, both lateral inhibition (LI) and neuronal spike frequency adaptation (SFA) are thought to enhance pattern separation for odor learning. However, it remains unclear whether these mechanisms play redundant or distinct roles in this process. In this study, we present a computational model of the fly olfactory circuit to investigate odor discrimination under varying noise conditions that simulate complex environments. Our results show that LI primarily enhances odor discrimination in low- and medium-noise scenarios, but this benefit diminishes and may reverse under higher-noise conditions. In contrast, SFA consistently improves discrimination across all noise levels. LI is preferentially engaged in low- and medium-noise environments, whereas SFA dominates in high-noise settings. When combined, these two sparsification mechanisms enable optimal discrimination performance. This work demonstrates that seemingly redundant modules in biological circuits can, in fact, be essential for achieving optimal learning in complex contexts. The code is available at: <https://github.com/L-Ocean/Fly-SNN>.

1 Introduction

Biological redundancy is commonly observed in the brain, where different regions, pathways, or mechanisms can perform similar functions [1–4]. Different motifs of biological redundancy may exist; for instance, different modules may have evolved to fulfill similar roles, ensuring robust neural functions under pathological conditions [5–8]. Alternatively, these mechanisms may be required to achieve near-optimal learning in more complex environments. Comparative analyses of the roles of putatively redundant modules in learning can clarify how the brain adapts and, in turn, inform theories that guide neuromorphic design [9, 10]. For instance, core computational features of the fly olfactory circuit have motivated the FlyLoRA architecture [11], which enhances task decoupling and parameter efficiency.

The fly olfactory system, a canonical cerebellum-like circuit, is a tractable model for dissecting neural computation owing to its relative simplicity and well-characterized anatomy and function

* Equal contribution.

† Corresponding author.

[12–18]. Two motifs—lateral inhibition (LI) and spike-frequency adaptation (SFA)—are prominent in the mushroom body and related circuits and both are known to shape neuronal responses [19–21]. LI is mediated by inhibitory interneurons that constrain the spatial spread of excitation and sharpen population representations [22, 23]. SFA reflects spike-triggered neuronal adaptation currents that accumulate during sustained stimulation, producing a progressive reduction in firing rate and emphasizing stimulus onsets and changes [24, 25].

Prior work has extensively characterized the roles of LI and SFA in neural information processing. LI drives competition and winner-take-all dynamics [26], increases population sparseness [27], enhances input contrast [28], and decorrelates activity patterns [29], it also facilitates subtractive or divisive gain modulation and strengthens regularization [30]. SFA functions as a nonlinear high-pass filter [31], promoting intensity-invariant coding [32], encoding changes in input statistics [33], sparsifying temporal responses, and forming short-term memory with non-stored retrieval (distinct from synaptic memory) [25].

Although the roles of LI and SFA in sparsifying neuronal responses are well documented, their relative contributions to learning—particularly under naturalistic conditions—remain less well characterized. Both mechanisms are expected to transform odor inputs into spatially and temporally sparse codes, attenuate noise, and enhance pattern separation [34–38](Figure 1). However, how they shape noisy odor representations and support discrimination learning across different noise regimes is not well understood. From an AI perspective, designing noise-robust classifiers is a classical challenge; elucidating how a fly learns to classify noisy odors may inform the development of robust, biologically inspired algorithms [39].

In this work, we developed a fly olfactory circuit model to investigate the roles of LI and SFA in odor discrimination tasks under varying noise conditions that simulate complex environments. Our main finding is that LI enhances learning performance in low- and medium-noise conditions, but this benefit gradually diminishes and may reverse when odors become noisier. In contrast, SFA consistently improves odor learning regardless of noise levels. LI tends to be more effective in low- and medium-noise environments, while SFA shows superior performance under high-noise conditions. When combined, the enhancement effects of these two mechanisms can be added up to achieve the optimal performance. These results suggest that seemingly redundant modules may be selectively recruited to optimize learning under different noisy conditions.

2 Method Overview

We developed a spiking neural network model of the fly olfactory circuit. The input layer consists of 50 olfactory receptor neurons (ORNs) that transduce odor stimuli into spikes. Each ORN projects to a corresponding projection neuron (PN), yielding 50 PNs. We also included local interneurons (LNs), which receive excitatory input from ORNs and provide lateral inhibition onto PNs. PN activity is relayed to 2,000 Kenyon cells (KCs) in the mushroom body; each KC samples inputs from approximately six PNs on average. Finally, KCs converge onto mushroom body output neurons (MBONs), which serve as readout units; in our task configuration, each MBON corresponds to an odor class being learned [40, 41].

The fly olfactory pathway exhibits three canonical features that are critical for discrimination [42]: large expansion (PN $\rightarrow$ KC), sparse connectivity, and sparse coding. Here, we focus on the role of spatiotemporally sparse spike coding, which conventional ANN models cannot capture; other architectural and connectivity parameters were constrained to experimentally observed values.

A schematic of the network and its putative role in odor discrimination is shown in Figure 1. After the KC stage, the sparsification mechanisms can transform odor responses in ways that facilitate odor discrimination. For example, they may preserve the original inter-class separability while increasing intra-class compactness (top), increase inter-class separation while keeping intra-class compactness unchanged (middle), or simultaneously achieve high intra-class compactness and improved inter-class separability (bottom). Regardless of the specific transformation, these changes are supposed to enhance the decision boundaries for odor discrimination.

Synaptic plasticity was restricted to the KC $\rightarrow$ MBON connections; all other synapses were fixed. Details of the training procedure and plasticity rule are provided in the Learning Algorithm section.

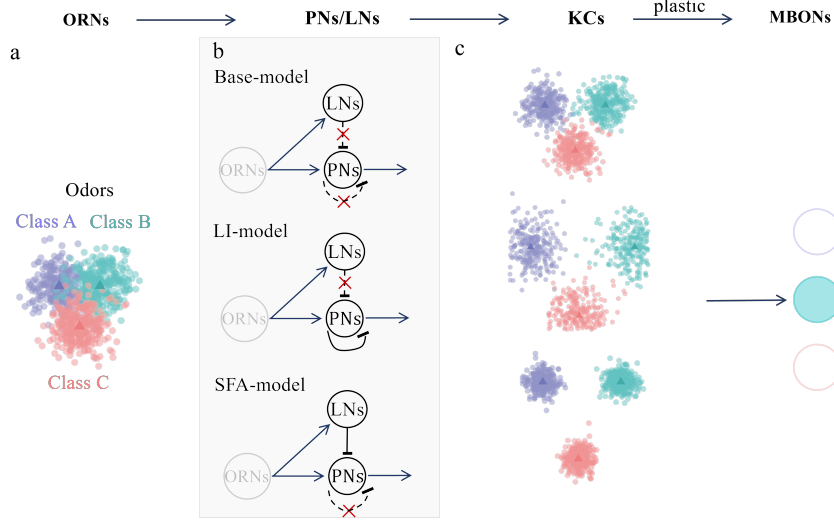


Figure 1: **Schematic of the fly olfactory circuit model.** Odor inputs are sensed and encoded by ORNs. After passing through the ORNs, odor-triggered spikes in PNs can be shaped by two factors: LI caused by LNs and SFA by inherent adaptation currents. Neuronal spikes in KCs may subsequently show different variation patterns to favor odor discrimination in MBONs.

Input Data. We adapted the odor-discrimination dataset from [43]. For each odor class i , we define a prototype ORN response vector $\mathbf{I}_i \in \mathbb{R}^{50}$. Its components are independently sampled from a uniform distribution: $I_{i,j} \sim \mathcal{U}(0, 1)$ for $j = 1, \dots, 50$. An individual sample from class i , denoted $\mathbf{I}_i^{(k)}$, is generated by adding an additive noise vector to the prototype: $\mathbf{I}_i^{(k)} = \mathbf{I}_i + \mathbf{Y}^{(k)}$, where every component of the noise vector $\mathbf{Y}^{(k)} \in \mathbb{R}^{50}$ is independently sampled from a zero-mean Gaussian distribution, $Y_j^{(k)} \sim \mathcal{N}(0, \sigma_{\text{noise}}^2)$. In accordance with the non-negativity of ORN firing rates, we applied element-wise clipping so that $\mathbf{I}_i^{(k)} \geq 0$.

Neuron Model. All neurons were modeled as leaky integrate-and-fire (LIF) units with a soft reset. The membrane potential of neuron j in the layer X evolves according to:

$$\tau_m^X \frac{dV_j^X(t)}{dt} = -V_j^X(t) + I_{\text{input},j}^X(t) + I_{\text{bias},j}^X + I_{\text{SFA},j}^X(t) + I_{\text{LI},j}^X(t) \quad (1)$$

Where τ_m^X is the membrane time constant. $I_{\text{input},j}^X(t)$ denotes the input current to neuron j in layer X , arising from odor-evoked stimulation or synaptic drive from presynaptic neurons; $I_{\text{bias},j}^X$ is a constant bias that sets a baseline drive (used primarily in PNs and LNs to provide a small excitatory input); $I_{\text{SFA},j}^X(t)$ is a spike-triggered adaptation current that depends on the neuron's recent spiking history and typically exerts a net inhibitory effect; and $I_{\text{LI},j}^X(t)$ is the lateral inhibitory current from LNs onto PNs, scaled by LN spiking activity.

In the model, a spike is emitted when $V_j^X(t)$ reaches the threshold V_{th}^X . Upon spiking, the membrane potential undergoes a soft reset: $V_j^X(t^+) = V_j^X(t) - V_{\text{th}}^X$. For simplicity, we used the same membrane time constant and threshold across layers ($\tau_m^X = \tau_m$, $V_{\text{th}}^X = V_{\text{th}}$), except in the readout (MBON) layer, where the threshold was set to 1.5 times higher. Given the dense $\text{KC} \rightarrow \text{MBON}$ connectivity, this higher threshold mitigates excessive spiking driven by the convergence of inputs from thousands of KCs.

LI Mechanism. LI onto PNs is mediated by LNs and modeled as an inhibitory current driven by an exponentially decaying trace of LN spiking. The inhibitory current onto PN j is:

$$I_{\text{LI},j}^{\text{PN}}(t) = \sum_k w_{\text{LN} \rightarrow \text{PN},jk} T_k^{\text{LN}}(t) \quad (2)$$

Where $w_{\text{LN} \rightarrow \text{PN},jk} \leq 0$ are inhibitory synaptic weights, and $T_k^{\text{LN}}(t)$ is the spike-triggered trace of LN k that integrates recent activity and decays over time:

$$\tau_{\text{trace_LN}} \frac{dT_k^{\text{LN}}(t)}{dt} = -T_k^{\text{LN}}(t) + S_k^{\text{LN}}(t) \quad (3)$$

$$S_k^{\text{LN}}(t) = \sum_l \delta(t - t_{kl}^{\text{LN}}) \quad (4)$$

with $\tau_{\text{trace_LN}}$ the decay time constant, $S_k^{\text{LN}}(t)$ the spike train of LN k , and t_{kl}^{LN} the l -th spike time. To compensate for the added inhibition and maintain comparable PN firing rates across conditions, we apply a modest increase in PN input drive when LI is enabled.

SFA Mechanism. SFA was implemented in PNs, LNs, and KCs as an inhibitory, spike-triggered current driven by an exponentially decaying state variable. For neuron j in layer X ,

$$I_{\text{SFA},j}^X(t) = w_{\text{SFA}}^X A_j^X(t) \quad (5)$$

$$\tau_{\text{SFA}}^X \frac{dA_j^X(t)}{dt} = -A_j^X(t) + S_j^X(t) \quad (6)$$

$$S_j^X(t) = \sum_l \delta(t - t_{j,l}^X) \quad (7)$$

where $A_j^X(t)$ integrates the recent spiking history and decays with time constant τ_{SFA}^X ; $w_{\text{SFA}}^X \leq 0$ sets the strength (sign) of the inhibitory adaptation current; $S_j^X(t)$ is the spike train of neuron j (Dirac delta representation), and $t_{j,l}^X$ denotes the l -th spike time.

To maintain comparable firing rates across conditions when SFA is enabled, we applied a small positive bias current to PNs and LNs.

Learning Algorithm. In our model, classification is based on the average membrane potential of the MBONs over a fixed evaluation window rather than on spike counts. Membrane potentials vary more smoothly than discrete spike trains and thus provide a more stable signal for gradient-based optimization [44, 45]. The complete training procedure is outlined in Algorithm 1.

Algorithm 1 Simplified SNN Training with Adaptive Mechanisms for Olfactory Tasks

```

1: Input: Training dataset  $D_{\text{train}}$ , Learnable weights  $W_{\text{KC} \rightarrow \text{MBON}}$ , Non-learnable parameters  $P$ .
2: for epoch = 1, 2, ..., num_epochs do
3:   for each batch  $(x_{\text{batch}}, y_{\text{batch}})$  in  $D_{\text{train}}$  do
4:     Phase 1: Temporal Forward Simulation
5:     Initialize all SNN states  $H_0$ . Set  $H_{\text{hist}} = []$ .
6:     for step  $t = 0$  to num_steps - 1 do
7:        $I_t = \text{Input}(x_{\text{batch}}, \text{step})$  ▷ Get input current
8:        $H_{t+1} = F(I_t, H_t, P)$  ▷ SNN forward pass and state update
9:       Record MBON membrane potentials in  $H_{\text{hist}}$ .
10:    end for
11:    Phase 2: Backpropagation and Weight Update
12:     $\hat{Y}_{\text{batch}} = G(H_{\text{hist}})$  ▷ Process MBON output
13:    loss =  $L(\hat{Y}_{\text{batch}}, y_{\text{batch}})$  ▷ Cross-Entropy Loss
14:    loss.backward() ▷ Perform Backpropagation Through Time
15:    Optimizer.step() ▷ Update  $W_{\text{KC} \rightarrow \text{MBON}}$ 
16:  end for
17: end for

```

For input sample i , we recorded each MBON's membrane potential over the evaluation window $[t_s, t_e]$ and computed its time average. Let N_{MBON} be the number of MBONs. The mean potential for MBON j and the corresponding vector across MBONs are

$$\bar{V}_{m,i}^j = \frac{1}{t_e - t_s} \int_{t_s}^{t_e} V_j(t) dt, \quad \bar{\mathbf{V}}_{m,i} \in \mathbb{R}^{N_{\text{MBON}}} \quad (8)$$

We convert $\bar{\mathbf{V}}_{m,i}$ to class probabilities using a softmax function and train the network with a cross-entropy loss:

$$p_i^j = \text{softmax}(\bar{\mathbf{V}}_{m,i})_j \quad (9)$$

$$L_i = - \sum_j^{N_{\text{MBON}}} y_i^j \log(p_i^j) \quad (10)$$

Where y_i^j is the one-hot target for sample i .

We optimized the learnable $\text{KC} \rightarrow \text{MBON}$ synaptic weights $W_{\text{KC} \rightarrow \text{MBON}}$ using backpropagation through time (BPTT), unrolling the network over the evaluation window. Because LIF spikes arise from a step nonlinearity with zero derivative almost everywhere, we used a surrogate gradient during the backward pass. Specifically, we replaced the derivative of the Heaviside with a smooth arctan-based surrogate:

$$\sigma'(u) \approx \frac{k_1}{1 + (k_2 u)^2} \quad (11)$$

where $k_1, k_2 > 0$ are scaling constants.

Training uses Adam with L2 weight decay (weight regularization) to improve generalization. We also employed a ReduceLROnPlateau scheduler that monitors classification accuracy on a held-out set and reduced the learning rate by a factor $\alpha = 0.2$ if accuracy did not improve for $n = 10$ epochs.

Each model was trained for 100 epochs using mini-batches of size 256. After each batch, the loss was computed, gradients were backpropagated through time, and parameters were updated using the optimizer.

3 Experiments

Dataset and models: We generated an odor dataset comprising 30,000 training samples and 10,000 test samples, each drawn at random as described above. For most simulations, we evaluated three network configurations: (i) the Baseline model, which includes neither LI nor SFA; (ii) the LI model, which adds LI onto PNs; and (iii) the SFA model, which implements SFA in PNs, LNs, and KCs. In Figure 4, we also simulated the full model with both LI and SFA.

Experimental settings: Simulations were implemented in Python (snnTorch) and run on an NVIDIA A800 GPU. Network dynamics were simulated with a 1-ms time step; each trial comprised a 10-ms pre-stimulus baseline followed by a 30-ms odor presentation. All neurons shared a membrane time constant of 10 ms. Firing thresholds were 0.8 for PNs, LNs, and KCs, and 1.2 for MBONs. Connectivity featured sparse $\text{PN} \rightarrow \text{KC}$ projections (each KC received inputs from 6 PNs; fixed weight 0.3) and fully connected $\text{KC} \rightarrow \text{MBON}$ synapses initialized uniformly in $[0, 0.08]$. LI used a 5-ms LN trace time constant, and SFA used a 50-ms adaptation time constant. Models were trained for 100 epochs (batch size 256) with Adam (initial learning rate 1.0×10^{-4}). Performance was evaluated from the time-averaged MBON membrane potentials over the stimulus window.

4 Results

We evaluated our fly olfactory network model on a noisy odor discrimination task by incorporating LI and SFA mechanisms. The results show that these two seemingly redundant mechanisms play complementary roles in learning across varying noise conditions, thereby optimizing odor discrimination.

4.1 Odor discrimination in noise-free conditions

Our results show that the fly olfactory circuit model learns odor discrimination effectively (Figure 2). We first examined discrimination in noise-free conditions. In the baseline model—without LI or SFA—the discrimination accuracy reaches 91.7% for 1,000 odor classes. Although accuracy decreases as the number of classes increases, it remains 74.72% for 10,000 odor classes. These

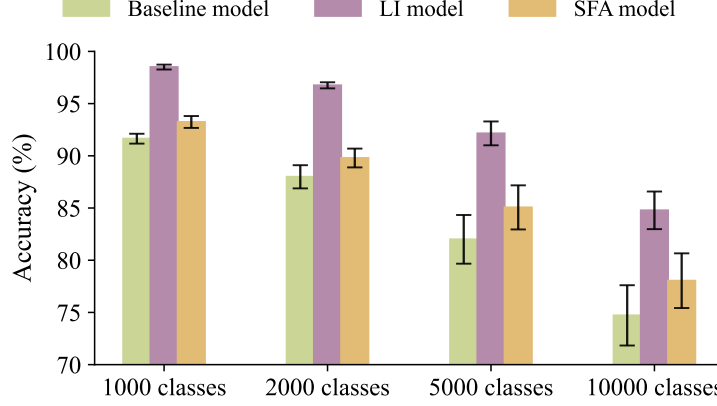


Figure 2: **Odor discrimination accuracy for fly olfactory circuit variants in a noise-free setting.** Performance is shown for three configurations—Baseline, LI, and SFA models—as a function of the number of odor classes (1,000–10,000).

findings indicate that the circuit’s other features confer strong pattern-classification capacity even without explicit sparsifying mechanisms in the circuit [42, 46].

Both LI and SFA improve odor discrimination relative to the Baseline model across all tested numbers of classes. By comparison, the LI model achieves significantly higher accuracy than the SFA model under the same conditions.

4.2 Odor discrimination in noisy conditions

Table 1 presents a comparative analysis of discrimination performance under varying noise intensities for the Baseline, LI, and SFA models, with odor classes ranging from 1,000 to 5,000. At low- and medium-noise levels—defined as noise intensity (N.I.) < 0.20 for 1,000- and 2,000-class odor discrimination, and N.I. < 0.15 for 5,000-class odor discrimination—the results are consistent with the noise-free context (LI model $>$ SFA model $>$ Baseline model), although overall test accuracy decreases. These findings indicate that sparsification of neuronal spikes in KCs via LI is more effective than SFA in enhancing odor discrimination under no- and low-noise conditions. However, when odor inputs become noisier (N.I. ≥ 0.20 for 1,000- and 2,000-class discrimination, and N.I. ≥ 0.15 for 5,000-class discrimination), the SFA model consistently achieves the highest discrimination accuracy.

Table 1: **Comparative performance of the Baseline, SFA, and LI models for 1,000-, 2,000-, and 5,000-class odor discrimination under different noise intensities.** Magenta values indicate the highest performance under each condition, violet values represent the second-best performance, and blue values denote the lowest performance.

	N.I.	1000 classes			2000 classes			5000 classes		
		Baseline	LI	SFA	Baseline	LI	SFA	Baseline	LI	SFA
Acc. (%)	0	91.70	98.30	93.50	88.85	96.85	90.90	82.00	92.15	85.06
	0.05	78.46	96.35	82.78	68.62	92.47	74.81	52.22	77.47	58.93
	0.1	74.61	91.85	78.26	61.70	83.07	67.87	38.65	57.83	45.18
	0.15	74.04	83.63	79.94	61.14	71.52	68.54	34.47	42.15	43.89
	0.2	72.64	74.78	78.77	58.80	58.87	67.34	32.70	31.91	43.04
	0.25	67.32	62.63	74.34	52.26	45.84	61.55	27.74	22.26	36.94
	0.3	59.03	53.82	69.34	43.34	37.50	55.78	20.26	16.55	30.38

We systematically analyzed the impact of the strength of each sparsification mechanism—LI and SFA—on discrimination performance (Figure 3). The simulation results indicate that both facilitation effects depend on noise levels, but in distinct ways. For LI, discrimination accuracy initially improves

as N.I. increases, but then gradually diminishes and can even reverse at higher noise levels. Across most of the tested N.I. range, stronger inhibition promotes discrimination accuracy. Compared to the Baseline model, strong inhibition improves accuracy by 6.80% when N.I. = 0.0 and 18.52% when N.I. = 0.1. Notably, although LI impairs discrimination at high N.I., (−4.5% when N.I. = 0.30), the reduction is still less pronounced for stronger inhibition.

In contrast to LI, SFA consistently facilitates odor discrimination across all tested N.I. ranges. As with LI, greater SFA strength generally yields higher accuracy, with improvements of 1.80% when N.I. = 0.0, 3.65% when N.I. = 0.1, and 11.56% when N.I. = 0.30.

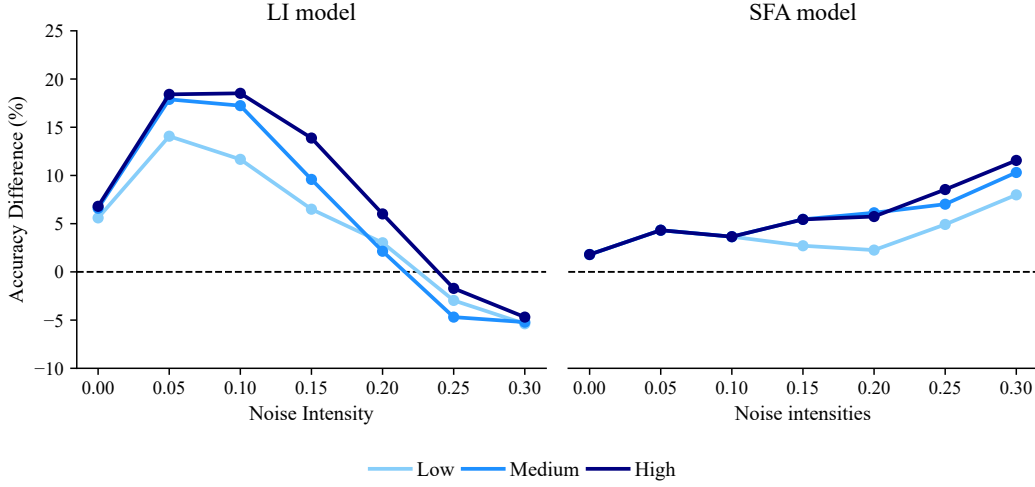


Figure 3: Changes in discrimination performance of the SFA and LI models relative to the Baseline model across varying degrees of inhibition and adaptation under different noise intensities. Results are shown for 1,000-class odor discrimination only. The “Low,” “Medium,” and “High” conditions (distinguished by color) represent increasing strengths of the respective mechanisms, achieved by systematically adjusting the relevant synaptic weights: $w_{LN \rightarrow PN}$ in Eq. 2 for LI, and $w_{SFA,X}$ in Eq. 5 for SFA. For both mechanisms, “Low,” “Medium,” and “High” correspond to weight increases in an approximate 1:2:3 ratio.

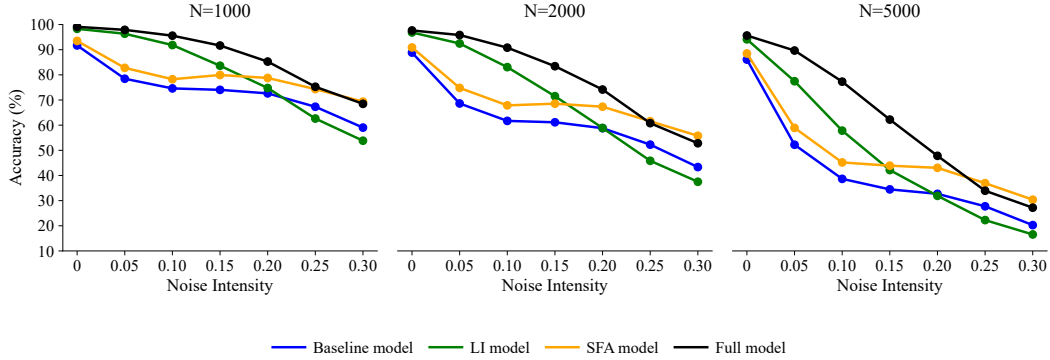


Figure 4: Discrimination performance of the SFA model, LI model, Full (SFA + LI) model, and Baseline model under different noise intensities. The models were tested on odor discrimination tasks with 1,000, 2,000, and 5,000 odor classes.

In addition, we investigated whether LI and SFA could be dynamically combined to achieve optimal odor discrimination learning (Figure 4). Our results show that orchestrating these two mechanisms produces higher discrimination accuracy than using either mechanism alone under low- and medium-noise levels. Only under high-noise conditions does the SFA model outperform all other models. The benefits of these two mechanisms are additive. These findings suggest that, although each mechanism

is most effective within specific noise regimes, LI and SFA can work together in a complementary manner, combining their individual effects to optimize the learning process.

4.3 Learning speed

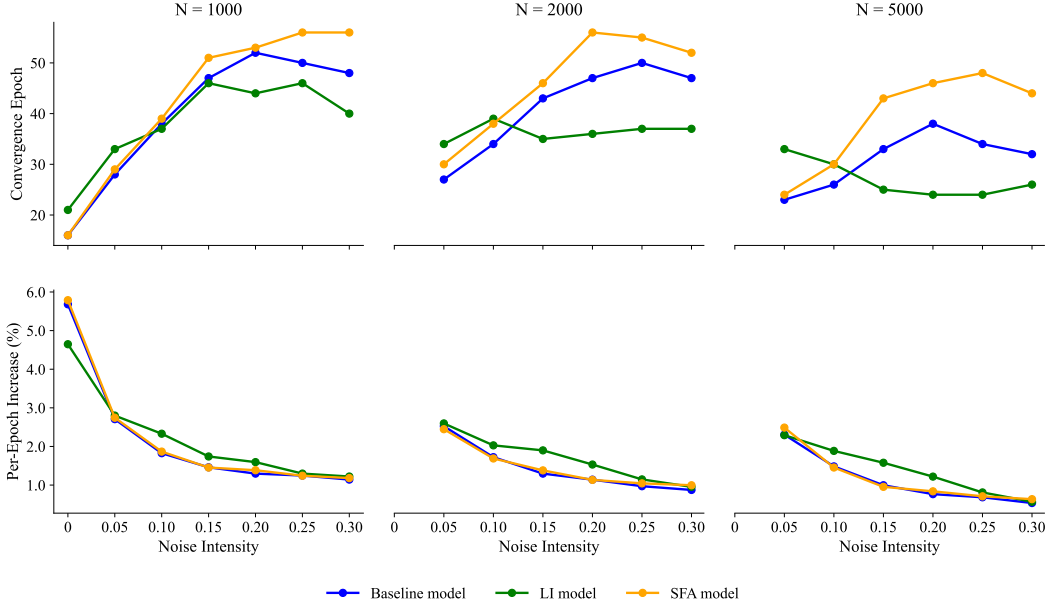


Figure 5: **Impacts of LI and SFA on model convergence speed.** The top (bottom) panel shows the number of training epochs required for convergence (the accuracy gain per epoch) for the Baseline, LI, and SFA models, plotted against noise intensity for different odor category sizes: 1,000, 2,000, and 5,000 classes. To avoid potential misinterpretations from relying solely on maximum accuracy and to objectively assess learning progress, convergence is defined as the point at which model accuracy growth plateaus. Specifically, a model is considered converged when the average improvement in accuracy over $n = 10$ consecutive epochs falls below a predefined threshold (threshold = 0.003).

To further assess the impacts of LI and SFA on learning efficiency, we provide a quantitative analysis of how these two sparsification mechanisms shape the time course of odor discrimination learning. As shown in Figure 5 (top), similar to the final discrimination accuracy values (Table 1), the number of training epochs required for convergence depends on noise level. In the low-noise range, the LI model requires more epochs to reach the defined converged state, regardless of odor category size. SFA follows LI, while the Baseline model reaches convergence fastest. The convergence speed can be explained by combining the final discrimination accuracy (Table 1) with the accuracy gain per epoch (Figure 5, bottom). In the low-noise range, both LI and SFA require higher final accuracy thresholds to converge, but their per-epoch gains show no advantage over the Baseline model, resulting in slower convergence.

In higher-noise ranges, however, the LI model exhibits the fastest convergence, followed by the Baseline model, with the SFA model consistently the slowest. In this regime, the LI model generally maintains an advantage over the Baseline model in per-epoch accuracy gain, explaining its faster convergence (Figure 5, top). When N.I. = 0.3, the per-epoch accuracy gain is similar across models; however, SFA requires a higher final accuracy threshold to reach convergence compared to both the Baseline and LI models (Table 1). Consequently, LI reaches the converged state first, the Baseline model second, and the SFA model last.

These results suggest that, depending on the noise level, learning speed—alongside discrimination accuracy—may be an important factor in determining which sparsification mechanism is recruited during noisy odor discrimination learning.

4.4 Sensitivity analysis

4.4.1 Sensitivity to noise types

The previous results were obtained by simulating the odor discrimination task with Gaussian noise. To assess the generalizability of our findings, we evaluated the effects of the sparsification mechanisms under different noise environments. In addition to Gaussian noise, we simulated noise generated by an Ornstein–Uhlenbeck (OU) process, which captures the temporal correlations often observed in natural odors. Odor samples with OU noise were generated using a procedure analogous to that used for Gaussian noise.

The results, summarized in Table 2, show that the effects of LI and SFA on discrimination performance with OU noise are generally consistent with those observed using Gaussian noise. LI achieves the highest discrimination accuracy at low- and medium-noise levels, whereas SFA is most effective at higher noise intensities. The main difference is that, within the tested noise range, LI—like SFA—consistently improves the discrimination of noisy odors compared to the Baseline model, rather than impairing performance at high noise levels, although its benefit gradually diminishes. Overall, these findings confirm that the complementary effects of LI and SFA are robust across the tested noise types.

Table 2: **Comparative performance of the Baseline, LI, and SFA models for 1,000-class odor discrimination under different OU noise intensities.**

	N.I.	Baseline	LI	SFA
Acc. (%)	0.1	86.02	97.98	88.66
	0.3	78.66	96.43	82.55
	0.5	76.01	93.10	78.85
	0.9	74.54	87.84	78.50
	1.5	67.55	69.72	77.49

4.4.2 Parameter sensitivity analysis

To further verify the reliability of our simulation results, we conducted a comprehensive parameter sensitivity analysis. Three key hyperparameters—random seed, learning rate, and batch size—were varied while keeping all other experimental conditions fixed to assess their impacts on odor discrimination.

Our tests reveal that:

- Varying the random seed typically changed accuracy by less than 1.3%.
- Adjusting the learning rate to 0.5–2.0 of its default value resulted in accuracy changes generally under 1.6%.
- Scaling the batch size to 0.5–2.0 of its standard value produced accuracy differences typically below 3.0%.

These results demonstrate that, within reasonable variation ranges, model performance is highly robust to hyperparameter choices. Importantly, across all parameter variations, our core conclusions remain unchanged: the LI model achieves optimal performance in low- and medium-noise environments, whereas the SFA model performs best under high-noise conditions. This analysis confirms that our main findings are reliable and independent of specific hyperparameter settings, providing a strong foundation for the interpretations presented in this study.

5 Conclusion

In this paper, we present a computational demonstration of how two distinct neural sparsification mechanisms—circuit-level LI and neuronal-level SFA—provide complementary advantages for noisy odor discrimination. Both mechanisms, well-documented in biological neural systems, actively shape and refine spatiotemporal neuronal responses. Our simulation results reveal a noise-dependent

recruitment pattern: LI is preferentially engaged under low-noise conditions, whereas SFA dominates in high-noise environments. The fly can orchestrate these two sparsification mechanisms to achieve optimal discrimination performance. These findings illustrate how seemingly redundant circuit modules in biological systems may, in fact, represent an optimized strategy for maintaining robust learning performance across diverse environmental conditions. This study sits at the interface of computational neuroscience and spiking neural networks, leveraging AI methods to investigate learning processes in the fly olfactory circuit. While this interdisciplinary approach is enriching, it also carries several limitations, as discussed below.

6 Limitations

We present a computational model of the fly olfactory circuit and investigate the roles of LI and SFA in odor discrimination learning. While our findings offer valuable insights, several limitations point to promising directions for future research.

First, regarding the training and test datasets, our simulations used artificially generated odors [43]. Although the odor generation method was experimentally inspired and provided a controlled environment, future work should validate odor discrimination performance using more naturalistic and physiologically realistic odor stimuli.

Second, our current model assumes that plasticity is confined to synaptic connections between KCs and MBONs. The synaptic update rule follows backpropagation through time, which is generally considered biologically implausible. While we believe this does not alter our conclusions, future studies should explore biologically plausible learning algorithms. Moreover, as in most cerebellum-like circuit models, the synaptic weights between the input and hidden layers (ORN–KC synapses) are fixed [18, 47]. These connections may also undergo other forms of plasticity, such as unsupervised learning via Oja’s rule [48]. Exploring these possibilities would provide a more comprehensive understanding of how learning unfolds across the entire circuit.

Finally, although our parameter robustness analysis and exploration of different noise types support the generalizability of our findings, the parameter ranges and noise characteristics examined remain limited. Broader investigations into diverse environmental conditions and increased model complexity could further strengthen the conclusions.

7 Acknowledgement

This work was supported by the National Key Research and Development Program of China (2023YFF1204200), the National Natural Science Foundation of China (62476197, 12372060, 92370103, 62176179), and the Xiaomi Foundation. The funders had no role in study design, data collection and analysis, decision to publish, or preparation of the manuscript.

References

- [1] Shreesh P Mysore and Eric I Knudsen. Reciprocal inhibition of inhibition: a circuit motif for flexible categorization in stimulus selection. *Neuron*, 73(1):193–205, 2012.
- [2] Laura N Driscoll, Krishna Shenoy, and David Sussillo. Flexible multitask computation in recurrent networks utilizes shared dynamical motifs. *Nature Neuroscience*, 27(7):1349–1363, 2024.
- [3] Brad K Hulse, Hannah Haberkern, Romain Franconville, Daniel B Turner-Evans, Shin-ya Takemura, Tanya Wolff, Marcella Noorman, Marisa Dreher, Chuntao Dan, Ruchi Parekh, et al. A connectome of the drosophila central complex reveals network motifs suitable for flexible navigation and context-dependent action selection. *Elife*, 10, 2021.
- [4] Yunliang Zang, Eve Marder, and Shimon Marom. Sodium channel slow inactivation normalizes firing in axons with uneven conductance distributions. *Current Biology*, 33(9):1818–1824.e3, 2023. ISSN 0960-9822. doi: <https://doi.org/10.1016/j.cub.2023.03.043>. URL <https://www.sciencedirect.com/science/article/pii/S0960982223003317>.
- [5] Jessica A Cardin. Functional flexibility in cortical circuits. *Current opinion in neurobiology*, 58:175–180, 2019.

- [6] Marijn CW Kroes and Guillén Fernández. Dynamic neural systems enable adaptive, flexible memories. *Neuroscience & Biobehavioral Reviews*, 36(7):1646–1666, 2012.
- [7] Theresa A Jones. Motor compensation and its effects on neural reorganization after stroke. *Nature Reviews Neuroscience*, 18(5):267–280, 2017.
- [8] Yukio Nishimura and Tadashi Isa. Cortical and subcortical compensatory mechanisms after spinal cord injury in monkeys. *Experimental neurology*, 235(1):152–161, 2012.
- [9] Chittotosh Ganguly, Sai Sukruth Bezugam, Elisabeth Abs, Melika Payvand, Sounak Dey, and Manan Suri. Spike frequency adaptation: bridging neural models and neuromorphic applications. *Communications Engineering*, 3(1):22, 2024.
- [10] Jason Yik, Korneel Van den Berghe, Douwe den Blanken, Younes Bouhadjar, Maxime Fabre, Paul Hueber, Weijie Ke, Mina A Khoei, Denis Kleyko, Noah Pacik-Nelson, et al. The neurobench framework for benchmarking neuromorphic computing algorithms and systems. *Nature Communications*, 16(1):1545, 2025.
- [11] Heming Zou, Yunliang Zang, Wutong Xu, Yao Zhu, and Xiangyang Ji. Flylora: Boosting task decoupling and parameter efficiency via implicit rank-wise mixture-of-experts, 2025. URL <https://arxiv.org/abs/2510.08396>.
- [12] Tiffany Kee, Pavel Sanda, Nitin Gupta, Mark Stopfer, and Maxim Bazhenov. Feed-forward versus feedback inhibition in a basic olfactory circuit. *PLoS computational biology*, 11(10):e1004531, 2015.
- [13] Alex C Keene and Scott Waddell. Drosophila olfactory memory: single genes to complex neural circuits. *Nature Reviews Neuroscience*, 8(5):341–354, 2007.
- [14] Africa Couto, Mattias Alenius, and Barry J Dickson. Molecular, anatomical, and functional organization of the drosophila olfactory system. *Current Biology*, 15(17):1535–1547, 2005.
- [15] Zhihao Zheng, Feng Li, Corey Fisher, Iqbal J Ali, Nadiya Sharifi, Steven Calle-Schuler, Joseph Hsu, Najla Masoodpanah, Lucia Kmecova, Tom Kazimiers, et al. Structured sampling of olfactory input by the fly mushroom body. *Current Biology*, 32(15):3334–3349, 2022.
- [16] Julia E Manoim-Wolkovitz, Tal Camchy, Eyal Rozenfeld, Hao-Hsin Chang, Hadas Lerner, Ya-Hui Chou, Ran Darshan, and Moshe Parnas. Nonlinear high-activity neuronal excitation enhances odor discrimination. *Current Biology*, 35(7):1521–1538, 2025.
- [17] Collins Assisi, Mark Stopfer, and Maxim Bazhenov. Optimality of sparse olfactory representations is not affected by network plasticity. *PLoS computational biology*, 16(2):e1007461, 2020.
- [18] Yunliang Zang and Erik De Schutter. Recent data on the cerebellum require new models and theories. *Current Opinion in Neurobiology*, 82:102765, 2023. ISSN 0959-4388. doi: <https://doi.org/10.1016/j.conb.2023.102765>. URL <https://www.sciencedirect.com/science/article/pii/S0959438823000909>.
- [19] Rachel I Wilson. Early olfactory processing in drosophila: mechanisms and principles. *Annual review of neuroscience*, 36(1):217–241, 2013.
- [20] Germain U Busto, Isaac Cervantes-Sandoval, and Ronald L Davis. Olfactory learning in drosophila. *Physiology*, 25(6):338–346, 2010.
- [21] Katherine I Nagel and Rachel I Wilson. Biophysical mechanisms underlying olfactory receptor neuron dynamics. *Nature neuroscience*, 14(2):208–216, 2011.
- [22] Luis M Franco, Zeynep Okray, Gerit A Linneweber, Bassem A Hassan, and Emre Yaksi. Reduced lateral inhibition impairs olfactory computations and behaviors in a drosophila model of fragile x syndrome. *Current Biology*, 27(8):1111–1123, 2017.
- [23] Adam M Large, Nathan W Vogler, Samantha Mielo, and Anne-Marie M Oswald. Balanced feedforward inhibition and dominant recurrent inhibition in olfactory cortex. *Proceedings of the National Academy of Sciences*, 113(8):2276–2281, 2016.
- [24] Farzad Farkhooi, Anja Froese, Eilif Muller, Randolph Menzel, and Martin P Nawrot. Cellular adaptation facilitates sparse and reliable coding in sensory pathways. *PLoS computational biology*, 9(10):e1003251, 2013.
- [25] Nils A Koch, Benjamin W Corrigan, Michael Feyerabend, Roberto A Gulli, Michelle S Jimenez-Sosa, Mohamad Abbass, Julia K Sunstrum, Sara Matovic, Megan Roussy, Rogelio Luna, et al. Spike frequency adaptation in primate lateral prefrontal cortex neurons results from interplay between intrinsic properties and circuit dynamics. *Cell Reports*, 44(1), 2025.

- [26] Mark Stopfer. Olfactory coding: inhibition reshapes odor responses. *Current biology*, 15(24):R996–R998, 2005.
- [27] Rachel I Wilson and Gilles Laurent. Role of gabaergic inhibition in shaping odor-evoked spatiotemporal patterns in the drosophila antennal lobe. *Journal of Neuroscience*, 25(40):9069–9079, 2005.
- [28] Joseph Del Rosario, Stefano Coletta, Soon Ho Kim, Zach Mobbille, Kayla Peelman, Brice Williams, Alan J Otsuki, Alejandra Del Castillo Valerio, Kendell Worden, Lou T Blanpain, et al. Lateral inhibition in v1 controls neural and perceptual contrast sensitivity. *Nature Neuroscience*, pages 1–12, 2025.
- [29] Sonya Giridhar, Brent Doiron, and Nathaniel N Urban. Timescale-dependent shaping of correlation by olfactory bulb lateral inhibition. *Proceedings of the National Academy of Sciences*, 108(14):5843–5848, 2011.
- [30] Shawn R Olsen and Rachel I Wilson. Lateral presynaptic inhibition mediates gain control in an olfactory circuit. *Nature*, 452(7190):956–960, 2008.
- [31] Jan Benda. Neural adaptation. *Current Biology*, 31(3):R110–R116, 2021.
- [32] Steven A Prescott and Terrence J Sejnowski. Spike-rate coding and spike-time coding are affected oppositely by different adaptation mechanisms. *Journal of Neuroscience*, 28(50):13649–13661, 2008.
- [33] Jan Clemens, Nofar Ozeri-Engelhard, and Mala Murthy. Fast intensity adaptation enhances the encoding of sound in drosophila. *Nature communications*, 9(1):134, 2018.
- [34] George Barnum and Elizabeth J Hong. Olfactory coding. *Current Biology*, 32(23):R1296–R1301, 2022.
- [35] Karol Gregor, Arthur Szlam, and Yann LeCun. Structured sparse coding via lateral inhibition. *Advances in Neural Information Processing Systems*, 24, 2011.
- [36] Shaul Druckmann, Tao Hu, and Dmitri Chklovskii. A mechanistic model of early sensory processing based on subtracting sparse representations. *Advances in Neural Information Processing Systems*, 25, 2012.
- [37] Sean X Luo, Richard Axel, and LF Abbott. Generating sparse and selective third-order responses in the olfactory system of the fly. *Proceedings of the National Academy of Sciences*, 107(23):10713–10718, 2010.
- [38] Matthew Chalk, Olivier Marre, and Gašper Tkačik. Toward a unified theory of efficient, predictive, and sparse coding. *Proceedings of the National Academy of Sciences*, 115(1):186–191, 2018.
- [39] Alex Lamy, Ziyuan Zhong, Aditya K Menon, and Nakul Verma. Noise-tolerant fair classification. *Advances in neural information processing systems*, 32, 2019.
- [40] Minna Ng, Robert D Roorda, Susana Q Lima, Boris V Zemelman, Patrick Morcillo, and Gero Miesenböck. Transmission of olfactory information between three populations of neurons in the antennal lobe of the fly. *Neuron*, 36(3):463–474, 2002.
- [41] Nicolas Y Masse, Glenn C Turner, and Gregory SXE Jefferis. Olfactory information processing in drosophila. *Current Biology*, 19(16):R700–R713, 2009.
- [42] Heming Zou, Yunliang Zang, and Xiangyang Ji. Structural features of the fly olfactory circuit mitigate the stability-plasticity dilemma in continual learning. *arXiv preprint arXiv:2502.01427*, 2025.
- [43] Peter Y Wang, Yi Sun, Richard Axel, LF Abbott, and Guangyu Robert Yang. Evolving the olfactory system with machine learning. *Neuron*, 109(23):3879–3892, 2021.
- [44] Yufei Guo, Yuanpei Chen, Liwen Zhang, Xiaode Liu, Yinglei Wang, Xuhui Huang, and Zhe Ma. Im-loss: information maximization loss for spiking neural networks. *Advances in Neural Information Processing Systems*, 35:156–166, 2022.
- [45] Hangchi Shen, Qian Zheng, Huamin Wang, and Gang Pan. Rethinking the membrane dynamics and optimization objectives of spiking neural networks. *Advances in Neural Information Processing Systems*, 37: 92697–92720, 2024.
- [46] Yunliang Zang, Eve Marder, and Shimon Marom. Sodium channel slow inactivation normalizes firing in axons with uneven conductance distributions. *Current biology*, 33(9):1818–1824, 2023.
- [47] Yunliang Zang and Erik De Schutter. Climbing fibers provide graded error signals in cerebellar learning. *Frontiers in Systems Neuroscience*, 13:56, 2019.
- [48] Erkki Oja. A simplified neuron model as a principal component analyzer. *Journal of Mathematical Biology*, 15(3):267–273, 1982.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: We have summarized the main contributions and scope in the abstract and introduction sections.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: We mentioned some limitations related to model performance in Appendix B.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[NA\]](#)

Justification: This thesis does not include formal theoretical results and mathematical proofs. It is mainly supported by experimental and simulation data of the computational model.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [\[Yes\]](#)

Justification: We introduced all the parameters related to the experiments in the Method overview and Experiments. And the code is available.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: We provide code in the github.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: We introduced the functions of different parameters in the Method overview.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: We present the results with error bars in Figure 2 and conduct a parameter sensitivity analysis.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).

- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: We summarize our computational resources in Experiment.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: We have read the Code of Ethics and make sure to preserve anonymity.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: We have discuss potential societal impacts in Introduction.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.

- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: Our work poses no such risks.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [NA]

Justification: We develop the model from scratch.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.

- If this information is not available online, the authors are encouraged to reach out to the asset’s creators.

13. **New assets**

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [\[Yes\]](#)

Justification: Our paper does not release new assets.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and research with human subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [\[NA\]](#)

Justification: Our paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. **Institutional review board (IRB) approvals or equivalent for research with human subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [\[NA\]](#)

Justification: Our paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: The core method development in this research does not involve LLMs as any important, original, or non-standard components.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.

A Discrimination Accuracy over Training Epochs under Different Noise Intensities

The main text primarily presents the final discrimination accuracy values under specific experimental conditions. To better illustrate the learning dynamics, this section includes results showing the time course of recognition accuracy for the three models (Baseline, LI, and SFA models) as they learn to discriminate 1,000 odor classes under varying noise intensities (N.I. = 0.1, 0.2, and 0.3), as shown in Figure 6.

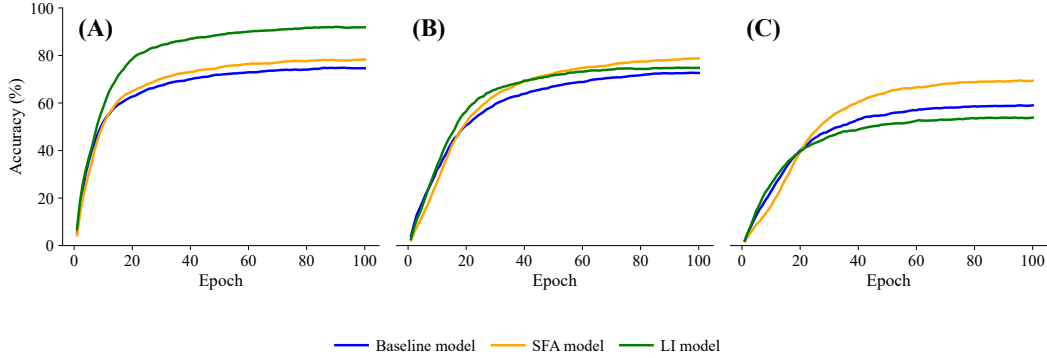


Figure 6: **Evolution of discrimination accuracy over training epochs under different noise intensities.** The subfigures show results for: (A) N.I. = 0.1, (B) N.I. = 0.2, and (C) N.I. = 0.3.

B Model Performance with Enhanced Strength of Input Signal

The experiments described in the main text imposed certain limits on the strength of the input odor signals to better isolate performance differences attributable to the LI and SFA mechanisms. Under conditions with a large number of odor categories and substantial noise interference, this constraint led to relatively low absolute accuracy values across all models. In this section, we demonstrate that a moderate increase in input signal strength can significantly enhance overall discrimination performance.

Table 3: **Comparative performance of the Baseline, SFA, and LI models for 10,000-class odor discrimination under different noise intensities with enhanced strength of input signal.**

	N.I.	Baseline	LI	SFA
Acc. (%)	0	99.82	99.99	99.97
	0.05	98.46	99.90	99.45
	0.1	96.87	98.59	98.14
	0.15	90.13	86.58	93.27
	0.2	69.80	60.79	81.31
	0.25	45.21	34.17	60.07
	0.3	25.94	18.40	37.45

Table 3 summarizes the final discrimination accuracy of the Baseline, LI, and SFA models after increasing the input signal strength. Enhancing the input signal strength led to substantial accuracy improvements for all models across the tested conditions. For example, with 10,000 odor classes and noise = 0.1, the accuracy of the Baseline model rose from 21.80% to 96.87%.