

# V1-LSM: Biologically Structured Liquid State Machines for Event-Based Vision

Devashri Deulkar<sup>1</sup> and Dr. Saumik Bhattacharya<sup>1</sup>

Indian Institute of Technology Kharagpur, West Bengal, India - 721302  
deulkardr@kgpian.iitkgp.ac.in, saumik@ece.iitkgp.ac.in

**Abstract.** Liquid state machines (LSMs) are well suited to event-based vision because they process time-varying spike streams through recurrent spiking dynamics. However, standard LSM reservoirs are usually random and homogeneous, unlike the structured organization of the primary visual cortex (V1). We propose V1-LSM, a modular reservoir for N-MNIST classification that introduces coarse cortical structure through retinotopic hypercolumns, layered intra-column processing, and sparse lateral coupling between neighboring columns. Each column implements a layered processing pathway: Layer 4 for input gating, Layer 2/3 as the recurrent liquid, Layer 5 for stable state extraction, and Layer 6 for gain modulation. A linear readout is trained on the resulting column-level features, while the reservoir itself remains fixed. In matched-neuron experiments on the N-MNIST dataset, the V1-inspired model improves test accuracy from 86.45% to 91.67% at 512 neurons and from 88.23% to 92.37% at 1024 neurons. These results suggest that coarse cortical organization can provide a practical inductive bias for event-based reservoir computing.

**Keywords:** Liquid state machine (LSM) · Reservoir computing · Spiking Neural Network (SNN) · Primary visual cortex (V1) · Event-based vision.

## 1 Introduction

Inspired by the biological retina, event-based vision sensors produce sparse, asynchronous events rather than dense, static image frames. This makes them naturally compatible with spiking neural networks, where information is represented through spike timing and dynamic state. Liquid state machines (LSMs) provide a simple way to use this kind of data: a recurrent spiking reservoir transforms the input stream into a high-dimensional state, and a trained readout maps that state to the task output [10]. In standard LSMs, the internal weights of the reservoir are kept static, which ensures training remain simple and makes the model attractive for neuromorphic systems.

While standard LSMs have random and homogeneous reservoirs, the primary visual cortex is not a single homogeneous pool. It has retinotopic organization, recurrent local circuits, excitatory-inhibitory balance, layer-specific pathways, and lateral interactions. In this work, we investigate whether incorporating such

coarse cortical structure into the reservoir improves performance for event-based vision tasks. We design a modular LSM organized into retinotopic hypercolumns, each with a layered pathway: an input stage analogous to Layer 4, a recurrent liquid inspired by Layers 2/3, and an output stage that aggregates column-level activity. Neighboring hypercolumns are connected through sparse lateral projections, introducing spatial structure into the reservoir dynamics.

We evaluate this V1-inspired architecture on N-MNIST classification using a fixed reservoir and a trained linear readout. The goal is not to reproduce the full biological complexity of the cortex, but to test whether simple, biologically motivated structural constraints can provide a useful inductive bias for reservoir computing in event-based vision, and whether a more cortex-aware reservoir is a better starting point for future neuromorphic vision systems.

### 1.1 Contributions

The main contributions are:

1. We propose V1-LSM, a biologically structured liquid state machine with retinotopic hypercolumns, layered intra-column processing (L2-6), and sparse lateral connections between neighboring columns.
2. We show that the V1-inspired structure improves classification accuracy over the base LSM on the N-MNIST dataset, raising test accuracy from 86.45% to 91.67% at 512 neurons and from 88.23% to 92.37% at 1024 neurons.
3. We identify coarse cortical organization as a useful inductive bias for event-based reservoir computing, suggesting a design direction for future neuromorphic vision systems.

## 2 Background

### 2.1 Spiking Neural Networks

Spiking Neural Networks (SNNs) are often described as the third generation of neural network models, where neurons communicate through discrete spikes over time rather than continuous activations [9]. This temporal perspective is especially relevant to event-driven inputs, such as event-based vision, as it naturally captures both the temporal and dynamic range of input.

**Leaky Integrate-and-Fire Neurons** This work uses the leaky integrate-and-fire (LIF) neuron model, a simplified biophysical model of a neuron. A LIF neuron is described by

$$\tau_m \frac{dV(t)}{dt} = -(V(t) - V_{\text{rest}}) + RI(t), \quad (1)$$

where  $V(t)$  is the membrane potential,  $\tau_m$  is the membrane time constant,  $V_{\text{rest}}$  is the resting potential, and  $I(t)$  is the input current. In this setting, the internal

state evolves dynamically: a neuron integrates incoming current, leaks toward a resting value, and emits a spike once its membrane potential crosses a threshold. This enables temporal coding and online processing, both fundamental to biological neural computation.

## 2.2 Event-Based Vision and N-MNIST

Unlike traditional cameras that capture dense frames at a fixed rate, event sensors asynchronously measure per-pixel brightness changes, producing sparse spatiotemporal event streams [5]. We evaluate our architecture on N-MNIST, a standard neuromorphic benchmark generated by recording static MNIST digits with a dynamic vision sensor undergoing saccade-like motions [11].

For preprocessing, the  $34 \times 34$  N-MNIST sensor field is processed using  $9 \times 9$  overlapping patches, retaining only positive-polarity events. This extraction yields 16 spatial channels. The data is then temporally binned over a 100 ms window, producing a final input tensor with 16 spatial channels and 1000 temporal bins.

## 2.3 Liquid State Machines

The Liquid State Machine (LSM) [10] is a computational framework that models the real-time processing of cortical microcircuits by processing time-varying inputs continuously. An LSM separates temporal integration from task-specific learning. A continuous input stream  $u(t)$  perturbs a fixed, recurrent spiking network (the reservoir), generating a high-dimensional transient liquid state:

$$x^M(t) = (L^M u)(t). \quad (2)$$

A memoryless readout then maps this state to the desired output:

$$y(t) = f^M(x^M(t)). \quad (3)$$

Because only the readout weights are trained, learning remains highly efficient. The liquid provides two critical computational properties: a fading memory that retains recent input history, and a separation property that maps distinct temporal inputs to linearly independent spatial states. This translates complex temporal dynamics into a static classification task, making LSMs ideal for neuromorphic vision.

However, practical LSM implementations typically rely on random, unstructured reservoirs, ignoring the inherent spatial organization of visual data. Conversely, attempting to build a perfectly detailed biophysical model of the visual cortex is computationally expensive and risks losing the flexibility that makes the LSM so useful. To bridge this gap, we treat macroscopic cortical topology as a structural abstraction. Prior work indicates that incorporating V1 anatomical features improves low-level feature extraction [15]. Building on this, we hypothesize that introducing coarse, biologically motivated structure specifically, retinotopic columns and functional layers can serve as a powerful inductive bias for the reservoir, improving classification accuracy while strictly preserving the efficiency of a fixed liquid.

### 3 Related Work

Standard LSMs [10] rely on unstructured reservoirs, which limits their capacity to fully leverage inherent spatial input features. To address this, structural constraints have been introduced via modular sub-networks [1] and deep hierarchical LSMs (D-LSMs) [19; 18], which capture richer multi-scale temporal representations. Recent reinforced variants [6; 7] further integrate hierarchical processing with winner-take-all (WTA) competition [12]. While these architectures introduce critical computational depth, they abstract away the topological organization of biological vision, leaving a gap for explicitly incorporating cortical structure.

Conversely, a separate line of research explores computational models inspired by the primary visual cortex (V1). Incorporating macroscopic biological mechanisms such as isolated layer-specific processing [14], lateral interactions [15], and sparse decoding pathways [17; 8; 16] has been shown to naturally enhance visual feature extraction. Furthermore, models incorporating stochastic synaptic dynamics [2; 3] and three-factor learning rules [4] demonstrate the functional richness available in biological microcircuits.

However, these bio-inspired models often require computationally expensive biophysical simulations or evaluate circuit components in isolation. Our work bridges these two domains. By adopting the multi-layered, retinotopic organization of V1 as a structural abstraction, we introduce a practical inductive bias to the standard LSM framework, pairing biological spatial priors with the training efficiency of a fixed dynamical reservoir.

---

**Algorithm 1** Retinotopic Spatiotemporal Preprocessing and Routing

---

**Input:**  $X^t$ ,  $K$ ,  $\{G_k\}_{k=1}^K$

**Output:** Column-specific routed spike streams  $\{X_k^t\}_{k=1}^K$  for all timesteps  $t$

**Retinotopic Routing**

- 1: **for** timestep  $t = 1$  to  $T$  **do**
  - 2:     **for** each column  $k = 1$  to  $K$  **do**
  - 3:          $X_k^t = G_k \odot X^t$   $\triangleright$  Apply fixed spatial mask via element-wise multiplication
  - 4:     **end for**
  - 5: **end for**
  - 6: **return**  $\{X_k^t\}_{k=1}^K$
- 

### 4 The V1-LSM Architecture

The proposed V1-LSM replaces the standard unstructured reservoir with  $K$  retinotopically arranged hypercolumns. As illustrated in Figure 1, this structured liquid is defined by three core mechanisms: retinotopic input routing, sparse

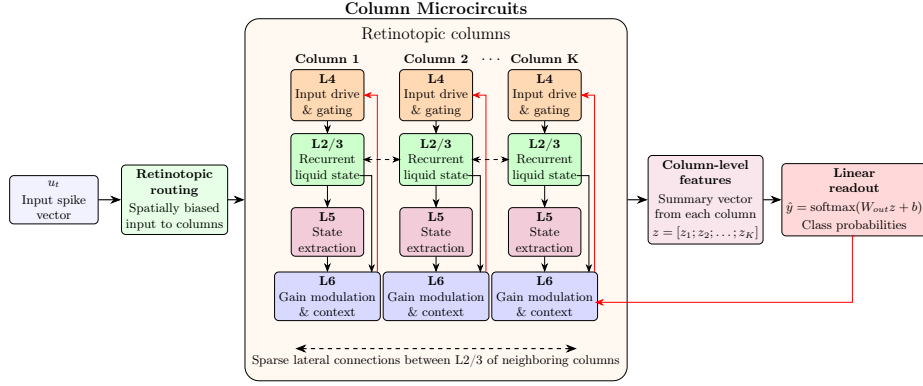


Fig. 1. Flow of the V1-inspired LSM.

lateral connectivity, and multi-layered intra-column microcircuits configured for distinct functional roles (input gating, recurrent processing, state extraction, and gain modulation).

#### 4.1 Retinotopic Input Routing

Standard LSMs broadcast input uniformly across the reservoir. In contrast, the V1-LSM routes specific spatial patches of the visual field to specific columns. The input received by column  $k$  at time  $t$  is defined by a retinotopic gain matrix:

$$u_{k,t} = g_k \odot u_t \quad (4)$$

where  $g_k$  is the column-specific gain vector and  $\odot$  denotes element-wise multiplication. This spatial routing ensures that each column receives a localized, biased view of the visual field, keeping neighboring sensory regions physically close within the reservoir. This routing makes the architecture different from an ensemble of independent reservoirs. The columns process the same visual field, but each column is biased toward a different spatial region.

#### 4.2 Intra-Column Microcircuit

Each of the  $K$  columns is divided into excitatory and inhibitory populations and distributed across four functional layers. The internal processing path for a single column  $k$  flows sequentially through these stages along with feedback from layer 6 :

$$u_{k,t} \rightarrow l_{k,t}^4 \rightarrow l_{k,t}^{23} \rightarrow l_{k,t}^5 \rightarrow z_k \quad (5)$$

- **Layer 4 (Input Drive):** This layer receives the routed input  $u_{k,t}$  and applies an initial stage of inhibitory regulation (feedback from layer 6) to bound the signal before it enters the main recurrent network.

- **Layer 2/3 (Recurrent Liquid):** This layer serves as the primary liquid state. It converts incoming L4 spikes  $l_{k,t}^4$  into high-dimensional liquid state  $l_{k,t}^{23}$ , translating complex temporal dynamics into instantaneous high-dimensional spatial patterns.
- **Layer 5 (State Extraction):** Instead of extracting features directly from the highly dynamic L2/3 layer, L5 provides state stabilization. It combines a running mean of L2/3 excitatory activity with instantaneous components to produce a stable state output  $h_{k,t}^5$ .
- **Layer 6 (Gain Modulation):** This layer integrates local activity from L2/3 and L5, alongside global feedback from the readout layer, to generate adaptive gain signals that dynamically modulate intra-column processing.

### 4.3 Inter-Column Lateral Coupling

The columns are linked by sparse nearest-neighbor lateral connections, allowing local spatial context to propagate across the reservoir. This preserves spatial coherence in the visual field without collapsing the architecture into a single unstructured pool.

The lateral input to column  $k$  is given by

$$\ell_{k,t} = \sum_{j \in \mathcal{N}(k)} A_{kj} W_{kj}^{\text{lat}} x_{j,t}^{23}, \quad (6)$$

where  $\mathcal{N}(k)$  is the set of neighboring columns,  $A_{kj}$  indicates whether a lateral connection exists, and  $W_{kj}^{\text{lat}}$  is the lateral weight matrix.

This weak coupling preserves local column specialization while still allowing information to spread across nearby visual regions.

### 4.4 Feature Extraction and Readout

After the input sequence propagates through the network, a feature summary  $z_k$  is extracted specifically from the L5 states of each hypercolumn. The global reservoir feature vector is formed by concatenating these column-level summaries:

$$z = [z_1; z_2; \dots; z_K]. \quad (7)$$

By passing these normalized statistics to the classifier, rather than the raw, high-dimensional liquid spike trains, the architecture provides a structured and stable representation that perfectly preserves the temporal dynamics of the reservoir. Finally, a trained linear readout maps this concatenated feature vector to the output class probabilities:

$$\hat{y} = \text{softmax}(W_{\text{out}} z + b). \quad (8)$$

Only the readout parameters are trained, liquid parameters remain fixed for experiments.

---

**Algorithm 2** V1-LSM Architecture
 

---

**Input:** Event stream  $X^t$  for  $t = 1 \dots T$ ,  $K$  retinotopic hypercolumns.

**Output:** Predicted class probabilities  $\hat{y}$ .

All layers utilize standard Leaky Integrate-and-Fire (LIF) membrane dynamics.

$S_l^t \in \{0, 1\}$  denotes the binary spike output of layer  $l$  at time  $t$ .

$\Theta(\cdot)$  denotes the Heaviside step function for spike generation.

**For each Hypercolumn**

```

1: for  $t = 1$  to  $T$  do
2:   for each column  $k = 1$  to  $K$  do
3:     Layer 4: Retinotopic Drive & Gating
4:      $I_{4,k}^t = (W_{in,k} X^t) \odot g_{6,k}^{t-1}$   $\triangleright$  Spatially routed, gain-modulated input
5:      $S_{4,k}^t = \Theta(V_{4,k}^t - \theta_4)$ 
6:     Layer 2/3: Recurrent Liquid
7:      $I_{lat,k}^t = \sum_{j \in \mathcal{N}(k)} W_{lat,j \rightarrow k} S_{23,j}^{t-1}$   $\triangleright$  Intercolumn Sparse lateral connection
8:      $I_{23,k}^t = W_{4 \rightarrow 23,k} S_{4,k}^t + W_{23 \rightarrow 23,k} S_{23,k}^{t-1} + I_{lat,k}^t$ 
9:      $S_{23,k}^t = \Theta(V_{23,k}^t - \theta_{23})$   $\triangleright$  Liquid State
10:    Layer 5: State Extraction
11:     $\tilde{S}_{23,k}^t = \alpha \tilde{S}_{23,k}^{t-1} + (1 - \alpha) S_{23,k}^t$   $\triangleright$  Running mean of liquid activity
12:     $H_{5,k}^t = \phi(W_{23 \rightarrow 5,k} [S_{23,k}^t, \tilde{S}_{23,k}^t])$   $\triangleright$  Hypercolumn output
13:    Layer 6: Internal Gain Modulation
14:     $H_{6,k}^t = \phi(W_{23 \rightarrow 6,k} S_{23,k}^t + W_{5 \rightarrow 6,k} H_{5,k}^t + W_{ctx \rightarrow 6,k} C^t)$   $\triangleright$  Global top-down
15:    context  $C^t$ 
16:     $g_{6,k}^t = \text{clip}(1 + W_{6 \rightarrow 4,k} H_{6,k}^t, g_{\min}, g_{\max})$   $\triangleright$  Update L4 gating for timestep
17:     $t + 1$ 
18:   end for
19: end for

20: Readout Interface
21: for each column  $k = 1$  to  $K$  do
22:    $z_k = \frac{1}{T} \sum_{t=1}^T H_{5,k}^t$ 
23: end for
24:  $z = [z_1; z_2; \dots; z_K]$ 
25:  $\hat{y} = \text{softmax}(W_{out} z + b_{out})$ 
26: return  $\hat{y}$ 

```

---

## 5 Experimental Setup

### 5.1 Task and Input Representation

We evaluate our architectures on the N-MNIST digit classification task, which represents traditional MNIST digits as spatiotemporal event streams generated through saccade-like sensor motion [11]. To ensure a fair comparison, both the base LSM and the proposed V1-LSM operate on the exact same preprocessed input representation. The preprocessing pipeline bins events over a 100 ms temporal window with a 0.1 ms simulation time step, aggregating them across 16 spatial patches arranged in a  $4 \times 4$  grid. By filtering exclusively for positive-polarity events, this yields the input tensor

$$U \in R^{C \times T}, \quad (9)$$

where  $C = 16$  is the number of input channels and  $T = 1000$  is the number of time bins. Each sample is therefore represented as a  $16 \times 1000$  input tensor.

### 5.2 Base LSM Baseline

The baseline model is a standard Liquid State Machine (LSM) consisting of a random recurrent spiking reservoir and a trained linear readout. The reservoir utilizes distance-dependent recurrent connectivity and maintains a strict excitatory-inhibitory balance, but lacks any explicit laminar or retinotopic organization. Following the generation of the high-dimensional liquid state, features are extracted and standardized using z-score normalization. A readout network is then trained on these features using a stratified 90/10 train-validation split and early stopping (summarized in Table 2).

### 5.3 Hypercolumn structure

The architecture consists of  $K$  retinotopic hypercolumns, each allocated a budget of 256 spiking neurons. To maximize the computational capacity of the reservoir, the majority of neurons are concentrated in the recurrent Layer 2/3, with the remainder distributed across the functional input and output layers (Table 1).

**Table 1.** Per-column neuron budget in the V1-inspired model.

| Layer | Neurons | Role                                 |
|-------|---------|--------------------------------------|
| L4    | 32      | Receives and gates routed input      |
| L2/3  | 128     | Main recurrent liquid                |
| L5    | 48      | Stable state extraction for readout  |
| L6    | 48      | Gain modulation and context feedback |

While biologically motivated by cortical microcircuit organization [13], this laminar separation is implemented as a structural abstraction rather than a strict biophysical reconstruction.



#### 5.4 Training Configuration and Evaluation

Both architectures use z-score feature normalization and a stratified 90/10 train-validation split (Table 2). To strictly evaluate the structural inductive bias of the V1-LSM, we utilize a fixed-reservoir paradigm: the internal reservoir dynamics act as a static feature extractor, and only the final linear readout is trained. While the implementation supports online plasticity, all intra-reservoir synaptic weights remain fixed for these experiments. Performance is measured via classification test accuracy.

To isolate architectural effects from parameter scaling, we conduct matched-neuron evaluations at two capacities:

1. **512-neuron budget:** A two-column V1-LSM compared against a 512-neuron unstructured baseline.
2. **1024-neuron budget:** A four-column V1-LSM compared against a 1024-neuron unstructured baseline.

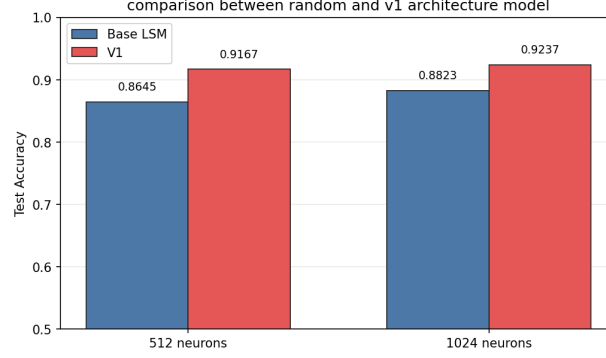
**Table 2.** Key experimental parameters for the evaluated architectures.

| Parameter             | Base LSM                  | V1-inspired LSM           |
|-----------------------|---------------------------|---------------------------|
| Dataset               | N-MNIST                   | N-MNIST                   |
| Input channels        | 16                        | 16                        |
| Validation ratio      | 0.1                       | 0.1                       |
| Feature normalization | z-score                   | z-score                   |
| Readout               | Trained linear classifier | Trained linear classifier |
| V1 columns            | –                         | 2 or 4                    |
| Neurons per column    | –                         | 256                       |
| Lateral density       | –                         | 0.05                      |

## 6 Results

As summarized in Table 3 and Figure 2, the V1-LSM consistently outperforms the unstructured baseline under matched-neuron conditions across both tested capacities. At a 512-neuron budget (two hypercolumns), the structured reservoir achieves a test accuracy of 91.67%, a substantial 5.22 percentage point (pp) improvement over the standard LSM (86.45%). This structural advantage scales effectively: at 1024 neurons (four hypercolumns), the V1-LSM reaches 92.37%, exceeding the matched baseline by 4.14 pp.

These results indicate that the performance gains are intrinsically linked to the structural abstraction of cortical organization. To further contextualize this parameter efficiency, we evaluated the scaling behavior of the unstructured baseline (Table 4). Doubling the baseline reservoir capacity from 1024 to 2048



Base LSM 512 and V1 2 columns both use 512 spiking neurons; Base LSM 1024 and V1 4 columns both use 1024 spiking neurons.

**Fig. 2.** Matched-neuron classification accuracy comparing the unstructured baseline LSM against the V1-inspired structured LSM.

**Table 3.** Classification accuracy under matched-neuron constraints.

| Model                       | Total Neurons | Test Accuracy | Gain over Baseline |
|-----------------------------|---------------|---------------|--------------------|
| Base LSM                    | 512           | 86.45%        | –                  |
| V1-inspired LSM (2 columns) | 512           | 91.67%        | +5.22 pp           |
| Base LSM                    | 1024          | 88.23%        | –                  |
| V1-inspired LSM (4 columns) | 1024          | 92.37%        | +4.14 pp           |

**Table 4.** Scaling behavior of the unstructured baseline LSM, illustrating the diminishing returns of simply increasing random reservoir capacity.

| Base LSM run  | Test Accuracy |
|---------------|---------------|
| Main base run | 75.98%        |
| 512 neurons   | 86.45%        |
| 1024 neurons  | 88.23%        |
| 2048 neurons  | 89.16%        |

neurons yields only marginal gains, peaking at 89.16%. Notably, the 1024-neuron V1-LSM (92.37%) significantly outperforms a random reservoir twice its size.

Finally, these quantitative findings demonstrate that introducing a coarse, biologically motivated topological structure provides a highly effective and parameter-efficient inductive bias for processing event-based visual streams.

## 7 Discussion

By strictly controlling the preprocessing pipeline, task, and neuron budget, our matched experiments support the view that the observed performance gains stem from the structural inductive biases introduced by the architecture.

### 7.1 Mechanisms of Structural Advantage

The improvement over the random reservoir likely reflects how the V1-LSM organizes the liquid, driven by three central design choices:

- **Retinotopic Routing:** Preserves spatial locality in event streams, ensuring that nearby events are processed within the same functional region.
- **Layered Microcircuits:** Separates input drive, recurrent mixing, state extraction, and gain modulation. This provides the readout with a more structured internal representation than a homogeneous reservoir.
- **Sparse Lateral Coupling:** Allows nearby columns to share local context without the dense global recurrence that can destabilize unstructured networks or reduce interpretability.

### 7.2 Scope and Limitations

These results should be read as evidence for a useful structural prior, not as a claim of state-of-the-art neuromorphic performance. The current study is defined by several boundaries:

- **Evaluation Scope:** The current evaluation is limited to the N-MNIST benchmark, and the input pipeline is heavily compressed (16 channels, positive-polarity events only).
- **Component Ablation:** The study does not yet isolate the quantitative contribution of each individual architectural component.
- **Frozen Plasticity:** While reward-modulated plasticity is implemented within the architecture, it is explicitly disabled in the present experiments to evaluate the reservoir’s structural prior in strict isolation.
- **Data Bottlenecks:** A broader challenge in event-based vision is the scarcity of standardized preprocessed datasets. Because extensive effort is currently required for event conversion and input preparation, the development of robust, field-wide data pipelines remains a critical need for accelerating neuromorphic model design.

## 8 Conclusion

This work demonstrates that the structural organization of the primary visual cortex (V1) serves as a highly effective inductive bias for event-based reservoir computing. We restructured a standard liquid state machine into retinotopically organized hypercolumns with layered intra-column processing and sparse lateral coupling, while keeping the reservoir fixed and training only the readout. Matched-neuron evaluations on the N-MNIST dataset confirm that the proposed V1-LSM consistently outperforms unstructured random reservoirs, proving that these gains stem directly from structural organization rather than parameter scaling. Rather than functioning as a strict biophysical replica, this approach establishes that treating macroscopic cortical structure as a modeling abstraction improves reservoir computation without sacrificing the training efficiency of a fixed LSM framework.

Future work will focus on layer-wise and component-wise ablations, evaluations on more complex event-based datasets, richer event encodings, and the use of existing plasticity mechanisms. A natural next step is to extend the same modular design toward higher visual areas such as V2, with the longer-term goal of building more biologically grounded neuromorphic vision systems.

**Acknowledgments.** The authors thank their peers and colleagues for their valuable feedback and insightful comments on earlier versions of this manuscript.

**Disclosure of Interests.** The authors have no competing interests to declare that are relevant to the content of this article.

## References

- [1] Dai, Y., Yamamoto, H., Sakuraba, M., Sato, S.: Computational efficiency of a modular reservoir network for image recognition. *Frontiers in Computational Neuroscience* **15**, 594337 (2021). <https://doi.org/10.3389/fncom.2021.594337>
- [2] El-Laithy, K., Bogdan, M.: Predicting spike-timing of a thalamic neuron using a stochastic synaptic model. In: *Proceedings of the 18th European Symposium on Artificial Neural Networks (ESANN 2010)*. pp. 357–362 (2010)
- [3] El-Laithy, K., Bogdan, M.: Enhancements on the modified stochastic synaptic model: The functional heterogeneity. pp. 389–396 (10 2017). [https://doi.org/10.1007/978-3-319-68600-4\\_45](https://doi.org/10.1007/978-3-319-68600-4_45)
- [4] Frémaux, N., Gerstner, W.: Neuromodulated spike-timing-dependent plasticity, and theory of three-factor learning rules. *Frontiers in Neural Circuits* **9**, 85 (2016). <https://doi.org/10.3389/fncir.2015.00085>
- [5] Gallego, G., Delbruck, T., Orchard, G., Bartolozzi, C., Taba, B., Censi, A., Leutenegger, S., Davison, A.J., Conradt, J., Daniilidis, K., Scaramuzza, D.: Event-based vision: A survey. *CoRR abs/1904.08405* (2019), <http://arxiv.org/abs/1904.08405>
- [6] Krenzer, D., Bogdan, M.: The reinforced liquid state machine: A new training architecture for spiking neural networks. In: *ESANN 2025 Proceedings*. pp. 699–704 (2025). <https://doi.org/10.14428/esann/2025.ES2025-1>
- [7] Krenzer, D., Bogdan, M.: Reinforced liquid state machines—new training strategies for spiking neural networks based on reinforcements. *Frontiers in Computational Neuroscience* **19**, 1569374 (2025). <https://doi.org/10.3389/fncom.2025.1569374>
- [8] Li, X., Yi, H., Luo, S.: Pattern recognition of spiking neural networks based on visual mechanism and supervised synaptic learning. *Neural Plasticity* **2020**(1), 8851351 (2020). <https://doi.org/https://doi.org/10.1155/2020/8851351>, <https://onlinelibrary.wiley.com/doi/abs/10.1155/2020/8851351>
- [9] Maass, W.: Networks of spiking neurons: The third generation of neural network models. *Neural Networks* **10**(9), 1659–1671 (1997). [https://doi.org/10.1016/S0893-6080\(97\)00011-7](https://doi.org/10.1016/S0893-6080(97)00011-7)
- [10] Maass, W., Natschlaeger, T., Markram, H.: Real-time computing without stable states: A new framework for neural computation based on perturbations. *Neural Computation* **14**(11), 2531–2560 (2002). <https://doi.org/10.1162/089976602760407955>
- [11] Orchard, G., Jayawant, A., Cohen, G.K., Thakor, N.: Converting static image datasets to spiking neuromorphic datasets using saccades. *Frontiers in Neuroscience* **9**, 437 (2015). <https://doi.org/10.3389/fnins.2015.00437>

- [12] Oster, M., Douglas, R., Liu, S.C.: Computation with spikes in a winner-take-all network. *Neural Computation* **21**(9), 2437–2465 (2009). <https://doi.org/10.1162/neco.2009.07-08-829>
- [13] Potjans, T.C., Diesmann, M.: The cell-type specific cortical microcircuit: Relating structure and activity in a full-scale spiking network model. *Cerebral Cortex* **24**(3), 785–806 (2014). <https://doi.org/10.1093/cercor/bhs358>
- [14] Probst, D., Maass, W., Markram, H., Gewaltig, M.O.: Liquid computing in a simplified model of cortical layer IV: Learning to balance a ball. In: *Artificial Neural Networks and Machine Learning – ICANN 2012. Lecture Notes in Computer Science*, vol. 7552, pp. 209–216. Springer (2012)
- [15] Roy, S., Pal, N.R.: An anatomy-based V1 model: Extraction of low-level features, reduction of distortion and a V1-inspired SOM. *arXiv preprint arXiv:2302.09074* (2023)
- [16] Santos-Mayo, A., Moratti, S., de Echegaray, J., Susi, G.: A model of the early visual system based on parallel spike-sequence detection, showing orientation selectivity. *Biology* **10**(8), 801 (2021). <https://doi.org/10.3390/biology10080801>
- [17] Shi, J., Wielaard, J., Smith, R.T., Sajda, P.: Perceptual decision making investigated via sparse decoding of a spiking neuron model of v1. In: *2009 4th International IEEE/EMBS Conference on Neural Engineering*. pp. 558–561 (2009). <https://doi.org/10.1109/NER.2009.5109357>
- [18] Soures, N., Kudithipudi, D.: Deep liquid state machines with neural plasticity for video activity recognition. *Frontiers in Neuroscience* **13**, 686 (2019). <https://doi.org/10.3389/fnins.2019.00686>
- [19] Wang, Q., Li, P.: D-LSM: Deep liquid state machine with unsupervised recurrent reservoir tuning. In: *2016 23rd International Conference on Pattern Recognition (ICPR)*. pp. 2652–2657. IEEE (2016). <https://doi.org/10.1109/ICPR.2016.7900035>