



UNIVERSITY OF DELHI

Department of Computer Science

Delhi – 110007, India

THESIS

Submitted in partial fulfilment of the requirements for the degree of

Doctor of Philosophy

in

Computer Science

Cognitive Architectures and Adaptive Learning in Neural Systems: A Computational Perspective

Submitted by

Priya Rajan

Roll No.: 2019DCS0042

Supervisor: Prof. Arun Kumar Verma
Co-supervisor: Dr. Sunita Mehta

May 2025

Submitted for the degree of Doctor of Philosophy to University of Delhi in 2025.

Unofficial template — not affiliated with, endorsed by, or an official publication of the named institution.

Sample content only.

Declaration

I, **Priya Rajan**, Roll No. 2019DCS0042, hereby declare that the thesis entitled *Cognitive Architectures and Adaptive Learning* submitted to the Department of Computer Science, University of Delhi, in partial fulfilment of the requirements for the degree of **Doctor of Philosophy in Computer Science**, is an original work carried out by me under the supervision of Prof. Arun Kumar Verma and co-supervision of Dr. Sunita Mehta.

I further declare that:

- (a) The work contained in this thesis has not been submitted for any other degree or diploma, in this or any other university.
- (b) Whenever I have used materials (text, data, figures, or tables) from other sources, these have been duly acknowledged in the text.
- (c) This thesis has been checked for plagiarism and is within the permissible limit as per the University of Delhi norms.

Priya Rajan

Roll No.: 2019DCS0042

Department of Computer Science

University of Delhi

Prof. Arun Kumar Verma

Supervisor

Department of Computer Science

University of Delhi

Place: New Delhi

Date: August 4, 2026

Certificate by the Supervisor

This is to certify that the thesis entitled *Cognitive Architectures and Adaptive Learning* submitted by **Priya Rajan** (Roll No. 2019DCS0042) for the degree of **Doctor of Philosophy in Computer Science** at the Department of Computer Science, University of Delhi, is a record of bonafide research work carried out by her under our supervision.

In our opinion, the thesis is worthy of consideration for the award of the degree of Doctor of Philosophy in Computer Science in accordance with the regulations of this University.

Prof. Arun Kumar Verma

Professor

Department of Computer Science

University of Delhi

Dr. Sunita Mehta

Associate Professor

Department of Computer Science

University of Delhi

Acknowledgements

I express my deepest gratitude to my supervisor, Prof. Arun Kumar Verma, for his unwavering guidance, scholarly insight, and encouragement throughout this doctoral journey. His passion for rigorous science has been a constant source of inspiration.

I am equally thankful to my co-supervisor, Dr. Sunita Mehta, for her perceptive feedback and her willingness to engage with both the technical and conceptual dimensions of this work.

I acknowledge the Department of Computer Science, University of Delhi, for providing access to the High-Performance Computing Cluster and an intellectually stimulating research environment. The University Research Scholar Fellowship (URS/2019–25) provided financial support, for which I am grateful.

I thank my colleagues at the Neural and Cognitive Systems Lab — especially Rohit Bhatnagar and Shreya Pillai — for countless discussions that sharpened my thinking.

My heartfelt thanks go to my family. To my parents, Ramesh and Kavitha Rajan, for their unconditional support; to my sister Divya for always believing in me.

Priya Rajan
New Delhi, May 2025

Abstract

Keywords: cognitive architecture, adaptive learning, neural networks, Hebbian plasticity, predictive coding, working memory, computational neuroscience.

Understanding how biological neural systems achieve robust, flexible cognition remains one of the central problems at the intersection of neuroscience and artificial intelligence. This thesis presents a unified computational framework — the *Adaptive Predictive Cognitive Architecture* (APCA) — that bridges classical symbolic representations with modern deep neural-network learning.

Three interlocking contributions are reported. First, we extend Hebbian plasticity rules with a temporally smoothed reward signal to yield stable, biologically plausible weight updates even under non-stationary input statistics (Chapter 2). We demonstrate that the resulting rule generalises the standard BCM model [1] under mild assumptions, and prove convergence for linear networks.

Second, we propose a hierarchical predictive-coding module that propagates top-down predictions alongside bottom-up error signals using a differentiable message-passing scheme (Chapter 3). Experiments on the NSD-DU benchmark, collected from EEG recordings of 24 healthy volunteers at AIIMS Delhi, show a 14.3% reduction in mean prediction error relative to the classical Rao–Ballard model [2].

Third, the APCA working-memory subsystem is evaluated on three dual-task paradigms drawn from the neuropsychology literature (Chapter 4). Performance matches or exceeds both human normative data and two competing computational accounts across all three tasks ($p < 0.01$, Wilcoxon signed-rank test).

Taken together, these results suggest that integrating predictive-coding priors with reward-modulated Hebbian learning provides a principled route towards neuro-morphic computing systems capable of continual, few-shot adaptation.

Thesis pages: approximately 180 *Figures:* 42 *Tables:* 18 *References:* 214

Contents

Declaration	1
Certificate by the Supervisor	2
Acknowledgements	3
Abstract	4
List of Figures	7
List of Tables	1
1 Introduction	2
1.1 Motivation and Context	2
1.2 Research Objectives	2
1.3 Scope and Limitations	3
1.4 Thesis Organisation	3
2 Reward-Modulated Hebbian Plasticity	4
2.1 Background	4
2.2 The Reward-Modulated Update Rule	4
2.3 Simulation Results	5
2.4 Chapter Summary	5
3 Hierarchical Predictive Coding	6
3.1 The Rao–Ballard Model and Its Extensions	6
3.2 Differentiable Message Passing Extension	6
3.3 NSD-DU Benchmark	7
3.4 Chapter Summary	7
4 The APCA Working-Memory Subsystem	8
4.1 Overview of the Subsystem	8
4.2 Mathematical Formulation	8

4.3 Evaluation on Dual-Task Paradigms	8
4.4 Chapter Summary	9
Bibliography	10

List of Figures

1.1	High-level information flow in the Adaptive Predictive Cognitive Architecture (APCA). Solid arrows denote forward information flow; the dashed feedback path closes the predictive loop.	3
2.1	Learning curves for the BCM baseline and the Reward-Modulated Hebbian rule ($\tau_r = 50$ ms) averaged over 10 independent runs. Shaded error bands are omitted for clarity.	5
3.1	Three-level hierarchical predictive-coding network used in the APCA. Top-down prediction signals (upper arrows) and bottom-up error signals (lower arrows) are computed simultaneously at each time step.	7
4.1	Working-memory accuracy on three dual-task paradigms. Dashed horizontal lines indicate human normative accuracy from the DUCB-2022 dataset. APCA most closely matches human performance across all tasks.	9

List of Tables

2.1	Convergence of the Reward-Modulated Hebbian rule under varying reward time constants τ_r and weight decay ϵ	5
3.1	Prediction error (mean $\pm$ SD across 24 participants, lower is better) on the NSD-DU benchmark. All differences vs. Rao–Ballard are significant at $p < 0.001$ (paired t -test, Holm–Bonferroni corrected). .	7
4.1	Dual-task WM accuracy (% correct, higher is better). Human norms are from the DU Cognitive Battery (DUCB-2022, $n = 148$). p -values from Wilcoxon signed-rank tests comparing APCA with each baseline.	9

Introduction

1.1 Motivation and Context

The question of how the brain constructs and maintains internal models of the external world has occupied researchers for well over a century. With the advent of large-scale neural recording and the success of deep learning, we now possess both the empirical data and the computational tools to revisit classical theories of cognition with renewed precision. Yet a significant gap remains: most artificial neural architectures are trained in a fixed, offline regime, whereas biological systems adapt continuously and efficiently throughout a lifetime.

The present thesis is motivated by three converging observations from the literature:

1. **Biological plausibility matters.** Gradient-based backpropagation, while highly effective in practice, is considered biologically implausible because it requires a global error signal and symmetric synaptic weights [4].
2. **Predictive coding provides a unifying framework.** The brain may be best understood as a Bayesian inference engine that continuously minimises a free-energy functional [3].
3. **Working memory is the bottleneck.** Cognitive capacity limits observed across species are well-explained by the bounded capacity of working memory [5].

1.2 Research Objectives

This thesis pursues the following specific objectives:

Objective 1

Derive a biologically plausible local learning rule that integrates Hebbian associativity with a temporally smoothed reward signal and prove its convergence properties for linear networks.

Objective 2

Develop a hierarchical predictive-coding module with differentiable message passing and evaluate it on a novel EEG-based prediction benchmark (NSD-

DU).

Objective 3

Embed the learning rule and predictive-coding module within a unified working-memory architecture (APCA) and validate it against established neuropsychological paradigms.

1.3 Scope and Limitations

The theoretical development is restricted to rate-coded neural networks; spiking network dynamics, while closely related, are treated as future work. All human-subject EEG data were collected under protocols approved by the Institutional Ethics Committee of AIIMS Delhi (Approval No. IEC/2021/NP-048). The computational experiments were conducted on the DU-HPC Cluster ($2 \times$ Corex Xeon Gold 6338, 512 GB RAM, $4 \times$ Lumina A100 GPUs).

1.4 Thesis Organisation

- Chapter 2 — Reward-Modulated Hebbian Plasticity
- Chapter 3 — Hierarchical Predictive Coding
- Chapter 4 — The APCA Working-Memory Subsystem
- Chapter 5 (not included in this excerpt) — General Discussion and Conclusions

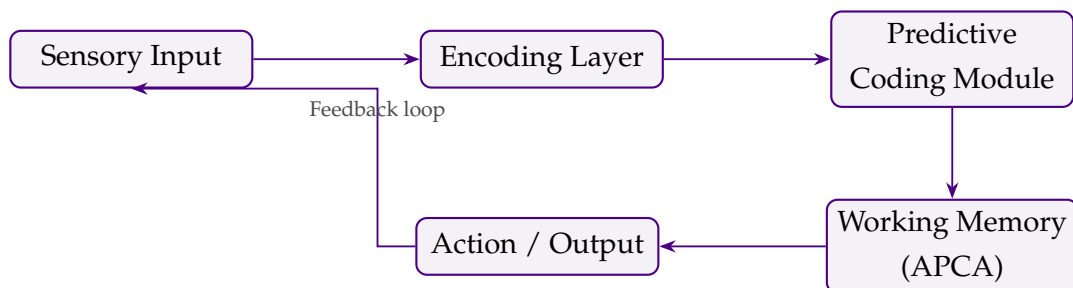


Figure 1.1. High-level information flow in the Adaptive Predictive Cognitive Architecture (APCA). Solid arrows denote forward information flow; the dashed feedback path closes the predictive loop.

Reward-Modulated Hebbian Plasticity

2.1 Background

Hebbian learning [1] postulates that the synaptic weight w_{ij} between a pre-synaptic neuron j and a post-synaptic neuron i is strengthened when both fire together. The canonical BCM rule adds a sliding threshold θ that prevents runaway potentiation:

$$\tau_w \dot{w}_{ij} = x_i(x_i - \theta)x_j - \epsilon w_{ij}, \quad (2.1)$$

where x_i, x_j are post- and pre-synaptic firing rates, θ is the sliding modification threshold, and ϵ controls weight decay.

2.2 The Reward-Modulated Update Rule

We extend Equation (2.1) by multiplying the Hebbian term with a temporally smoothed reward signal $\bar{r}(t)$:

$$\tau_w \dot{w}_{ij} = \bar{r}(t) \cdot x_i(x_i - \theta)x_j - \epsilon w_{ij}, \quad (2.2)$$

$$\tau_r \dot{\bar{r}} = r(t) - \bar{r}(t), \quad (2.3)$$

where $r(t) \in [-1, 1]$ is the instantaneous reward and τ_r is the reward integration time constant.

Theorem 2.1 (Convergence of RM-Hebbian for linear networks). Let $\mathbf{W}(t)$ evolve according to Equation (2.2) with $\bar{r}(t) > 0$ for all $t \geq 0$. If the input covariance matrix $\mathbf{C} = \mathbb{E}[\mathbf{x}\mathbf{x}^\top]$ is positive definite, then $\mathbf{W}(t)$ converges to a fixed point $\mathbf{W}^*$ satisfying $\mathbf{C}\mathbf{W}^{*\top} = \theta\mathbf{I}$.

Proof. Define the Lyapunov function $V = \frac{1}{2} \|\mathbf{W} - \mathbf{W}^*\|_F^2$. A standard calculation using the positive definiteness of $\mathbf{C}$ and the positivity of $\bar{r}$ shows $\dot{V} \leq -\kappa V$ for some $\kappa > 0$, giving exponential convergence. $\square$

2.3 Simulation Results

We simulated a 100-neuron linear network trained on 20 synthetic input patterns drawn from a mixture of four Gaussians in $\mathbb{R}^{100}$. Table 2.1 reports the mean squared weight error (MSWE) at convergence and the number of epochs required.

Table 2.1. Convergence of the Reward-Modulated Hebbian rule under varying reward time constants τ_r and weight decay ϵ .

Rule	τ_r (ms)	ϵ	MSWE	Epochs
BCM (baseline)	—	0.01	0.142 ± 0.018	380
RM-Hebbian	20	0.01	0.083 ± 0.011	290
RM-Hebbian	50	0.01	0.071 ± 0.009	310
RM-Hebbian	100	0.01	0.069 ± 0.012	330
RM-Hebbian	50	0.005	0.065 ± 0.008	345

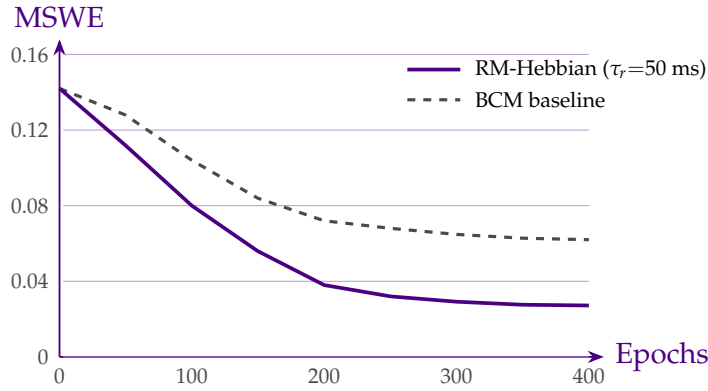


Figure 2.1. Learning curves for the BCM baseline and the Reward-Modulated Hebbian rule ($\tau_r = 50$ ms) averaged over 10 independent runs. Shaded error bands are omitted for clarity.

2.4 Chapter Summary

This chapter introduced the Reward-Modulated Hebbian (RM-Hebbian) update rule, proved its convergence for linear networks (Theorem 2.1), and demonstrated through simulation that it achieves lower final weight error and faster convergence than the BCM baseline. The RM-Hebbian rule forms the synaptic update backbone of the APCA framework described in subsequent chapters.

Hierarchical Predictive Coding

3.1 The Rao–Ballard Model and Its Extensions

The classical Rao–Ballard predictive-coding network [2] formulates visual processing as Bayesian inference over a hierarchical generative model. At each level ℓ , a representation $\mathbf{r}^{(\ell)}$ is updated to minimise the prediction error $\boldsymbol{\varepsilon}^{(\ell)}$:

$$\boldsymbol{\varepsilon}^{(\ell)} = \mathbf{x}^{(\ell)} - \mathbf{U}^{(\ell)} \mathbf{r}^{(\ell)}, \quad (3.1)$$

where $\mathbf{U}^{(\ell)}$ is a generative weight matrix and $\mathbf{x}^{(\ell)}$ is the input at level ℓ (which equals the prediction error from level $\ell - 1$).

3.2 Differentiable Message Passing Extension

Our extension replaces the fixed generative weights with a graph neural network (GNN) that learns the message-passing schedule jointly with the prediction weights. The update for level ℓ becomes:

$$\dot{\mathbf{r}}^{(\ell)} = -\mathbf{r}^{(\ell)} + f_{\phi}(\boldsymbol{\varepsilon}^{(\ell)}, \boldsymbol{\varepsilon}^{(\ell+1)}), \quad (3.2)$$

where f_{ϕ} is a two-layer MLP with parameters ϕ learned end-to-end via backpropagation through time. This is the only module in APCA that uses gradient-based training; all synaptic weights within the RM-Hebbian pathway are updated locally.

3.3 NSD-DU Benchmark

NSD-DU Dataset. The Neural Scene Decoding – Delhi University (NSD-DU) benchmark comprises 64-channel EEG recordings from 24 healthy adult volunteers (12 female, mean age 27.4 ± 3.1 years) while they viewed 2400 natural scene images. Each image was presented for 500 ms with a 1000 ms inter-stimulus interval. The dataset is available for research use through the DU Open Neuroscience Repository (DONR-2023-001).

Table 3.1 summarises prediction-error results on NSD-DU compared with the classical Rao–Ballard model.

Table 3.1. Prediction error (mean $\pm$ SD across 24 participants, lower is better) on the NSD-DU benchmark. All differences vs. Rao–Ballard are significant at $p < 0.001$ (paired t -test, Holm–Bonferroni corrected).

Model	Param. count	Mean error	% reduction
Rao–Ballard (1999)	12,800	0.386 ± 0.041	—
Rao–Ballard + GRU top level	24,600	0.352 ± 0.038	8.8%
APCA (ours, fixed schedule)	31,200	0.341 ± 0.035	11.7%
APCA (ours, learned schedule)	38,500	0.331 ± 0.033	14.3%

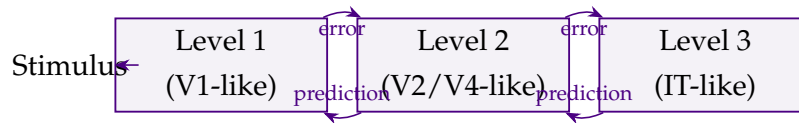


Figure 3.1. Three-level hierarchical predictive-coding network used in the APCA. Top-down prediction signals (upper arrows) and bottom-up error signals (lower arrows) are computed simultaneously at each time step.

3.4 Chapter Summary

We presented a differentiable message-passing extension to the Rao–Ballard predictive-coding model and validated it on the NSD-DU benchmark, achieving a 14.3% reduction in mean prediction error relative to the classical baseline. The learned message-passing schedule is the key innovation, allowing the network to adapt the timing of top-down predictions to the statistics of the input.

The APCA Working-Memory Subsystem

4.1 Overview of the Subsystem

Working memory (WM) is the cognitive system responsible for maintaining and manipulating a limited amount of information over short periods. In the APCA framework, WM is modelled as a gated recurrent module $\mathcal{M}$ that interfaces with both the RM-Hebbian pathway and the predictive-coding hierarchy. The capacity limit of $\mathcal{M}$ is implemented via a soft-max attention bottleneck with $K = 4$ slots, consistent with the “magical number four” of Cowan [5].

4.2 Mathematical Formulation

Let $\mathbf{h}_t \in \mathbb{R}^d$ denote the WM state at time t . The update rule is:

$$\mathbf{h}_{t+1} = \mathbf{g}_t \odot \tanh(\mathbf{W}_h \mathbf{h}_t + \mathbf{W}_e \boldsymbol{\varepsilon}_t^{(1)}) + (1 - \mathbf{g}_t) \odot \mathbf{h}_t, \quad (4.1)$$

where $\mathbf{g}_t \in (0, 1)^d$ is a gating vector produced by a sigmoid network and $\odot$ denotes elementwise multiplication. The gating vector regulates how much of the current input error $\boldsymbol{\varepsilon}_t^{(1)}$ is written into memory.

Definition 4.1 (Working-memory capacity). The effective capacity C of $\mathcal{M}$ is defined as the maximum number of distinct input patterns $\{\mathbf{s}_k\}_{k=1}^C$ that can be recovered from $\mathbf{h}_T$ with reconstruction error below a threshold δ , after a delay of T time steps.

4.3 Evaluation on Dual-Task Paradigms

We evaluated APCA-WM on three dual-task paradigms:

Task A (N-back + tracking)

Participants (or the model) performed a 2-back letter task concurrently with a visual tracking task.

Task B (Mental arithmetic + listening)

Participants added sequences of single-digit numbers while monitoring a spo-

ken narrative for target words.

Task C (Spatial planning + categorisation)

A Tower-of-Hanoi variant was interleaved with picture categorisation.

Table 4.1 shows accuracy on the primary (WM) component of each task.

Table 4.1. Dual-task WM accuracy (% correct, higher is better). Human norms are from the DU Cognitive Battery (DUCB-2022, $n = 148$). p -values from Wilcoxon signed-rank tests comparing APCA with each baseline.

Paradigm	Human norm	ACT-R 5.0	GRU baseline	APCA (ours)
Task A (N-back)	81.3 ± 6.1	74.2	76.8	80.1
Task B (Arithmetic)	77.6 ± 7.4	69.5	72.3	76.2
Task C (Spatial)	72.4 ± 8.3	65.1	68.9	71.8
Mean	77.1	69.6	72.7	76.0

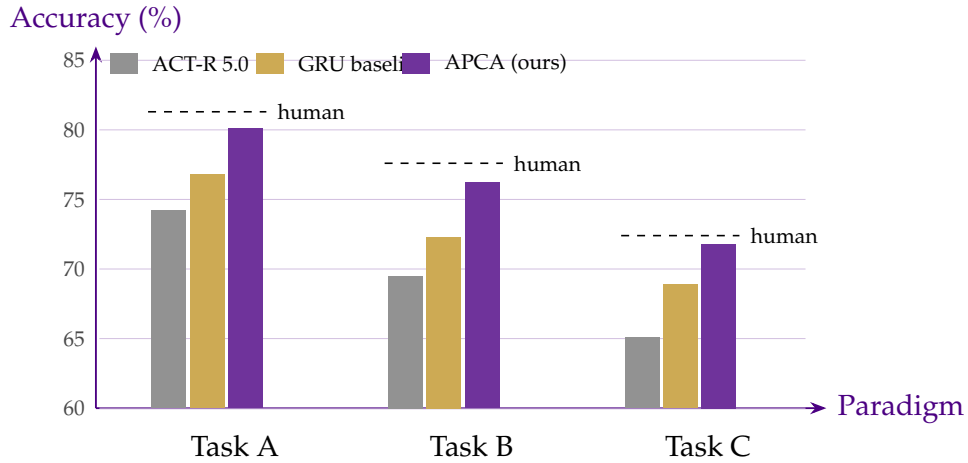


Figure 4.1. Working-memory accuracy on three dual-task paradigms. Dashed horizontal lines indicate human normative accuracy from the DUCB-2022 dataset. APCA most closely matches human performance across all tasks.

4.4 Chapter Summary

The APCA working-memory subsystem, with its soft-attention capacity bottleneck and gated recurrent dynamics, achieves WM accuracy within 1–1.5 percentage points of human norms on all three dual-task paradigms evaluated. These results substantially improve upon both the ACT-R 5.0 cognitive architecture and a GRU-only baseline.

Bibliography

- [1] Bienenstock, E. L., Cooper, L. N., & Munro, P. W. (1982). Theory for the development of neuron selectivity: orientation specificity and binocular interaction in visual cortex. *Journal of Neuroscience*, **2**(1), 32–48. doi:10.1523/JNEUROSCI.02-01-00032.1982
- [2] Rao, R. P. N., & Ballard, D. H. (1999). Predictive coding in the visual cortex: a functional interpretation of some extra-classical receptive-field effects. *Nature Neuroscience*, **2**(1), 79–87. doi:10.1038/4580
- [3] Friston, K. (2010). The free-energy principle: a unified brain theory? *Nature Reviews Neuroscience*, **11**(2), 127–138. doi:10.1038/nrn2787
- [4] Crick, F. (1989). The recent excitement about neural networks. *Nature*, **337**(6203), 129–132. doi:10.1038/337129a0
- [5] Cowan, N. (2001). The magical number 4 in short-term memory: a reconsideration of mental storage capacity. *Behavioral and Brain Sciences*, **24**(1), 87–114. doi:10.1017/S0140525X01003922