

UDC 519.651

Analyzing steady state Variance in Hebbian Learning: A Moment Closure Approach

Edgar A. Vardanyan

Russian-Armenian University, Yerevan, Armenia
e-mail: edgarvardanyan1999@gmail.com

Abstract

Hebbian learning, an important concept in neural networks, is the basis for various learning algorithms that model the adaptation of neural connections, also known as synapses. Among these models, Oja's rule stands out as an important example, giving valuable insights into the dynamics of unsupervised learning algorithms. The fact that the final steady-state solution of a single-layer network that learns using Oja's rule equals the solution of Principal component analysis is well known. However, the way in which the learning rate can affect the variance of the final parameters is less explored. In this paper, we investigate how different learning rates can influence the variance of parameters in Oja's rule, utilizing the moment closure approximation. By focusing on the variance, we offer new perspectives on the behavior of Oja's rule under varying conditions. We derive a closed-form equation that connects the parameter variance with the learning rate and shows that the relationship between these is linear. This gives valuable insights that may help to optimize the learning process of Hebbian models.

Keywords: Oja's rule; Hebbian learning; learning rate; neural networks; moment closure approximation.

Article info: Received 13 October 2024; sent for review 31 October 2024; accepted 21 November 2024.

1. Introduction

In recent decades, intensive research has been conducted on synaptic plasticity and learning. Much of this work was inspired by Hebb's postulate [1]. The main idea of Hebbian learning is that changes in synaptic transmission efficiency are driven by correlated firing activities of neurons connected by the synapse. Hebbian theory postulates that connections between neurons become stronger when they are activated at the same time.

Synaptic wiring processes are widely believed to be an integral part of the encoding of memories in the brain [2]. As a result, Hebbian learning has been studied as a biologically plausible algorithm for extracting patterns from different types of data. Unlike backpropagation, Hebbian learning does not require any labeled data and is an unsupervised learning

algorithm. It is believed that this unsupervised approach is the most common way the brain learns. This makes Hebbian Learning particularly interesting because of our desire to understand the human brain and because of the scarcity of labeled data in many problems [3]. As a result of this, Hebbian learning has found numerous applications in various fields such as computer vision and modeling of human memory [4, 5].

In a single-layer architecture, Hebb's rule can be formally expressed using the weight update equation:

$$y(t) = \sum_{i=1}^N w_i(t) \cdot x_i(t)$$

$$w_i(t+1) = w_i(t) + \alpha F(w_i(t), x_i(t), y(t))$$

where $w_i(t)$ is the synaptic coupling strength of the i -th input neuron at time step t , $x_i(t)$ is the i -th input neuron value, y is the output neuron value, and α will be referred to as the learning rate of the system.

This is the general form of Hebbian learning. F here is an undetermined function with an important limitation being the exclusion of any argument other than the existing synaptic coupling strength and the values of pre-synaptic and post-synaptic neurons [6]. Building on Hebb's rule, different specific forms of learning rules have been developed over time [7, 8].

The analysis provided in this paper will focus on studying the Oja's rule. Oja's rule solves stability problems encountered in other learning rules. It projects high-dimensional data into lower dimensions while preserving the maximal variance, thus generalizing the Principal Component Analysis. The updating rule for the weights in Oja's Rule is given by:

$$w_i(t+1) = w_i(t) + \alpha[x_i(t)y - y^2w_i(t)] \quad (1)$$

where $w_i(t)$ is the weight of the i -th variable at time step t , $x_i(t)$ is the i -th input variable, y is the output, and α is the learning rate. The term $y^2w_i(t)$ in the update rule ensures that the weights do not grow indefinitely, overcoming the stability limitation frequently encountered in Hebb's rule [9].

When creating artificial neural networks with Oja's rule or other similar rules, learning rate becomes one of the most important parameters. High learning rates may cause divergence, while low learning rates may cause slower training time. This creates a tradeoff, which can be controlled by adjusting the learning rate.

In backpropagation-based neural networks, learning rate schedulers that adapt based on the loss function are commonly used to enhance the convergence of the network [10]. However, in the context of Oja's Rule, there is no explicit loss function, and thus, traditional learning rate schedulers cannot be employed. This necessitates alternative approaches for learning rate adjustment in Hebbian learning models. Further research is needed to explore these possibilities and to understand the impact of learning rate on the convergence and stability of Hebbian-based networks.

This paper concentrates on analyzing the impact of the learning rate on final variance of the parameters in Oja's Rule. A closed-form formula is derived that connects the final variance of parameters with the learning rate of the system for a bivariate normal distribution data using the moment closure approximation [11, 12, 13]. This formula is validated using a comparison with numerical values derived from computer simulations. Understanding these variance relations can help in establishing metrics on how well converged is the lossless network that can be controlled simply by adjusting the learning rate.

2. Problem Setup

In this work, we consider a two-variable case, wherein the two variables are denoted as x_1 and x_2 . The data for these variables is assumed to be generated from a bivariate normal distribution. Without loss of generality, we focus on normalized data. This assumption is crucial as it simplifies the covariance matrix and aids in further analysis.

The data (x_1, x_2) is modeled as a bivariate normal distribution with the following properties:

1. The mean of the distribution is 0 for both variables, i.e., $\mu_{x_1} = \mu_{x_2} = 0$.
2. The data is normalized, so the variances $\sigma_{x_1}^2$ and $\sigma_{x_2}^2$ are both 1.

Given the above, the covariance matrix Σ of the distribution is:

$$\Sigma = \begin{pmatrix} 1 & \rho \\ \rho & 1 \end{pmatrix}$$

where ρ is the correlation coefficient between x_1 and x_2 . The value of ρ lies in the interval $[-1, 1]$, where $\rho = 1$ indicates a perfect positive correlation, $\rho = -1$ indicates a perfect negative correlation, and $\rho = 0$ indicates no correlation.

Given the above properties, the distribution of (x_1, x_2) is denoted as:

$$(x_1, x_2) \sim \mathcal{N}(\mathbf{0}, \Sigma) \quad (2)$$

where $\mathbf{0}$ is a vector of zeros representing the mean, and Σ is the covariance matrix as defined previously.

The specific structure of the covariance matrix has significant implications for learning in neural networks employing Oja's Rule. It has been shown that Oja's rule extracts principal components from the data, trying to create the signal with the highest variance [9]. As the data comes from a normalized distribution with known correlations, the learning dynamics and, as we will see in Section 4, the final variance of the parameters in Oja's Rule can be analyzed as a function of the learning rate α and the correlation coefficient ρ .

3. Steady State Solution

It is known that the stable steady state solution of Oja's rule matches with the solution of the Principal Component Analysis(PCA), meaning that our weight vector will be an eigenvector of the covariance matrix that corresponds to the biggest eigenvalue (other eigenvectors are non-stable solutions, which means that if you move w a little away from this solution, it won't come back)[14, 15, 16, 9].

A sketch of the proof will be provided, and the exact stable solutions for the bi-variate case will be calculated in this section. Let's first define $\mathbf{x}$ and $\mathbf{w}$ as:

$$\mathbf{x} = \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_n \end{bmatrix}, \quad \mathbf{w} = \begin{bmatrix} w_1 \\ w_2 \\ \vdots \\ w_n \end{bmatrix}$$

Now we can derive the steady state solution by asserting that the expected value of the change of the weights is equal to zero.

$$y = w^T x = x^T w$$

$$\begin{aligned}
\frac{1}{\alpha} E[\Delta w(t)] &= E[yx - y^2 w] \\
&= E[xy - y^2 w] \\
&= E[(xx^T)w - (w^T x)(x^T w)w] \\
&= E[(xx^T)]w - (w^T E[xx^T]w)w \\
&= \Sigma w - (w^T \Sigma w)w = 0
\end{aligned}$$

From this, we have that at the steady state w is an eigenvector of Σ , whose eigenvalue λ is equal to $(w^T \Sigma w)$. From this we can derive the L^2 norm of the steady state solution.

$$\begin{aligned}
\lambda = (w^T \Sigma w) &= w^T \lambda w = \lambda w^T w \\
w^T w &= 1 \\
\|w\|_2^2 &= 1
\end{aligned}$$

This will allow us to find the steady state solution for our bi-variate case. If our correlation ρ is positive, the largest eigenvalue will be $\lambda = 1 + \rho$, and the steady state solution will be

$$w = \pm \begin{bmatrix} \frac{\sqrt{2}}{2} \\ \frac{\sqrt{2}}{2} \end{bmatrix}.$$

If ρ is negative, we will have $\lambda = 1 - \rho$ and

$$w = \pm \begin{bmatrix} \frac{\sqrt{2}}{2} \\ -\frac{\sqrt{2}}{2} \end{bmatrix}.$$

The sign of the steady state solution will depend solely on the initial values of the weights.

4. Variance of the weights at the Steady State

Let's define the mean and the variance of the weights during the steady state solution.

$$\begin{aligned}
\mu_i(t) &= E[w_i(t)] \\
V_i(t) &= E[(w_i(t) - \mu_i(t))^2] \\
V_{ij}(t) &= E[(w_i(t) - \mu_i(t))(w_j(t) - \mu_j(t))] \\
\hat{\mu}_i &= \lim_{t \rightarrow \infty} \mu_i(t) = \frac{\sqrt{2}}{2} \\
\hat{V}_i &= \lim_{t \rightarrow \infty} V_i(t), \quad \hat{V}_{ij} = \lim_{t \rightarrow \infty} V_{ij}(t)
\end{aligned} \tag{3}$$

From this point on, we will do the calculations only for the positively correlated pairs. The same calculations will hold if ρ is negative as well. Let's introduce auxiliary variables $\bar{w}_1$ and $\bar{w}_2$ for describing the state of our system as the difference between the weights and their steady state solutions.

$$\begin{aligned}
\bar{w}_1 &= w_1 - \frac{\sqrt{2}}{2} \\
\bar{w}_2 &= w_2 - \frac{\sqrt{2}}{2}
\end{aligned} \tag{4}$$

This notation is natural as the variances whose steady state solution we are trying to calculate can be expressed by these variables using (4) and (3) as $V_1 = E[\bar{w}_1^2]$, $V_2 = E[\bar{w}_2^2]$ and $V_{12} = E[\bar{w}_1\bar{w}_2]$. According to (1), those variables will be updated at each step by the following rule.

$$\begin{aligned}\bar{w}_1(t+1) &= \bar{w}_1(t) + \alpha [x_1(t)y(t) - y(t)^2 w_1(t)] \\ \bar{w}_2(t+1) &= \bar{w}_2(t) + \alpha [x_2(t)y(t) - y(t)^2 w_2(t)]\end{aligned}\quad (5)$$

If we substitute the values of y , w_1 and w_2 in this by their representations through x_1 , x_2 , $\bar{w}_1$ and $\bar{w}_2$, we get the following update rule for our auxiliary variables after the expansion of (5)¹.

$$\begin{aligned}\bar{w}_1 \leftarrow \bar{w}_1 &+ \alpha \left(\left(\frac{1}{2\sqrt{2}} - \frac{\bar{w}_1}{2} - \frac{3\bar{w}_1^2}{\sqrt{2}} - \bar{w}_1^3 \right) x_1^2 \right. \\ &+ \left(-\frac{1}{2\sqrt{2}} - \frac{\bar{w}_1}{2} - \bar{w}_2 - \sqrt{2}\bar{w}_1\bar{w}_2 - \frac{\bar{w}_2^2}{\sqrt{2}} - \bar{w}_1\bar{w}_2^2 \right) x_2^2 \\ &+ \left. \left(-2\bar{w}_1 - \sqrt{2}\bar{w}_1^2 - 2\sqrt{2}\bar{w}_1\bar{w}_2 - 2\bar{w}_1^2\bar{w}_2 \right) x_1x_2 \right)\end{aligned}\quad (6)$$

$$\begin{aligned}\bar{w}_2 \leftarrow \bar{w}_2 &+ \alpha \left(\left(\frac{1}{2\sqrt{2}} - \frac{\bar{w}_2}{2} - \frac{3\bar{w}_2^2}{\sqrt{2}} - \bar{w}_2^3 \right) x_2^2 \right. \\ &+ \left(-\frac{1}{2\sqrt{2}} - \frac{\bar{w}_2}{2} - \bar{w}_1 - \sqrt{2}\bar{w}_2\bar{w}_1 - \frac{\bar{w}_1^2}{\sqrt{2}} - \bar{w}_2\bar{w}_1^2 \right) x_1^2 \\ &+ \left. \left(-2\bar{w}_2 - \sqrt{2}\bar{w}_2^2 - 2\sqrt{2}\bar{w}_2\bar{w}_1 - 2\bar{w}_2^2\bar{w}_1 \right) x_1x_2 \right)\end{aligned}\quad (7)$$

From the symmetry of the problem, it is obvious that $\hat{V}_1 = \hat{V}_2$. This means, that to calculate the steady state variances tracking the expected values of $\bar{w}_1^2$ and $\bar{w}_1\bar{w}_2$ is sufficient. Rules of their update can be calculated by multiplying (6) with itself and with (7). After these multiplications, we will have the following update rule.

$$\begin{aligned}\bar{w}_1^2 \leftarrow &\bar{w}_1^2 + \alpha^2 \left(\frac{x_1^4}{8} - \frac{x_1^2x_2^2}{4} + \frac{x_2^4}{8} \right) + \alpha^2 \bar{w}_2^2 \left(-\frac{1}{2}x_1^2x_2^2 + \frac{3x_2^4}{2} \right) \\ &+ \bar{w}_1\bar{w}_2 \left(-2\alpha x_2^2 + \alpha^2 (-2x_1^3x_2 + 6x_1x_2^3 + 2x_2^4) \right) \\ &+ \bar{w}_1^2 \left(\alpha (-x_1^2 - 4x_1x_2 - x_2^2) \right. \\ &+ \left. \alpha^2 \left(-\frac{5x_1^4}{4} + x_1^3x_2 + 6x_1^2x_2^2 + 3x_1x_2^3 + \frac{x_2^4}{4} \right) \right) \\ &+ \bar{w}_1f_1(x_1, x_2) + \bar{w}_2f_2(x_1, x_2) + \sum_{i+j \geq 3} \left(w_1^i w_2^j f_{ij}(x_1, x_2) \right)\end{aligned}\quad (8)$$

¹For simplicity from now on we will use $F(w_1, w_2) \leftarrow G(w_1, w_2, x_1, x_2)$ to notate update rules of the form $F(w_1(t+1), w_2(t+1)) = G(w_1(t), w_2(t), x_1(t), x_2(t))$.

$$\begin{aligned}
\bar{w}_1 \bar{w}_2 \leftarrow & \bar{w}_1 \bar{w}_2 + \alpha^2 \left(-\frac{x_1^4}{8} + \frac{x_1^2 x_2^2}{4} - \frac{x_2^4}{8} \right) \\
& + \bar{w}_1^2 \left(-\alpha x_1^2 + \alpha^2 \left(x_1^4 + \frac{5x_1^3 x_2}{2} - \frac{x_1 x_2^3}{2} \right) \right) \\
& + \bar{w}_2^2 \left(-\alpha x_2^2 + \alpha^2 \left(x_2^4 + \frac{5x_1 x_2^3}{2} - \frac{x_1^3 x_2}{2} \right) \right) \\
& + \bar{w}_1 \bar{w}_2 \left(\alpha \left(-x_1^2 - 4x_1 x_2 - x_2^2 \right) \right. \\
& \quad \left. + \alpha^2 \left(-\frac{x_1^4}{4} + 2x_1^3 x_2 + \frac{13x_1^2 x_2^2}{2} + 2x_1 x_2^3 - \frac{x_2^4}{4} \right) \right) \\
& + \bar{w}_1 g(x_1, x_2) + \bar{w}_2 g(x_1, x_2) + \sum_{i+j \geq 3} \left(w_1^i w_2^j g_{ij}(x_1, x_2) \right) \quad (9)
\end{aligned}$$

Here f_1 , f_2 , f_{ij} , g_1 , g_2 , and g_{ij} are polynomial functions of two variables. To estimate the steady state variances we must calculate the expected values of both sides of the above equations at the limit of $t \rightarrow \infty$. In order to complete these calculations we must take into account the following.

- At the limit of $t \rightarrow \infty$ we conclude from the symmetry of our multivariate Gaussian data distribution that $\hat{V}_1 = \hat{V}_2$. We will notate this variance using $\hat{V}$.
- Since $\bar{w}_i(t)$ and $x_j(t)$ are independent for any i and j

$$E \left[F \left(x_1(t), x_2(t) \right) G \left(w_1(t), w_2(t) \right) \right] = E \left[F \left(x_1(t), x_2(t) \right) \right] E \left[G \left(w_1(t), w_2(t) \right) \right]$$

- $x_1(t)$ and $x_2(t)$ are coming from a multivariate normal distribution, their moments can be calculated using Isserlis' theorem [17, 18].

$$\begin{aligned}
E[x_1(t)] &= E[x_2(t)] &= 0 \\
E[x_1^3(t)] &= E[x_2^3(t)] &= 0 \\
E[x_1^2(t)x_2(t)] &= E[x_2^2(t)x_1(t)] &= 0 \\
E[x_1^2(t)] &= E[x_2^2(t)] &= 1 \\
E[x_1(t)x_2(t)] &= \rho \\
E[x_1^2(t)x_2^2(t)] &= 1 + 2\rho^2 \\
E[x_1^3(t)x_2(t)] &= E[x_2^3(t)x_1(t)] &= 3\rho \\
E[x_1^4(t)] &= E[x_2^4(t)] &= 3
\end{aligned}$$

For calculating the expected values of polynomials involving $\bar{w}_1$ and $\bar{w}_2$ we will use the second-order moment closure approximation [11, 12, 13]. We will also keep in mind that at the limit of $t \rightarrow \infty$ expected values of $\bar{w}_1$ and $\bar{w}_2$ are equal to zero.

$$E[F(\bar{w}_1 \bar{w}_2)] \approx \frac{1}{2} \left(\frac{\partial F(\bar{w}_1 \bar{w}_2)}{\partial^2 \bar{w}_1} + \frac{\partial F(\bar{w}_1 \bar{w}_2)}{\partial^2 \bar{w}_2} \right) V + \left(\frac{\partial F(\bar{w}_1 \bar{w}_2)}{\partial \bar{w}_1 \partial \bar{w}_2} \right) V_{12}$$

Thus, we may obtain the following equations for the expected values of polynomials involving $\bar{w}_1$ and $\bar{w}_2$

$$\begin{aligned} E[\bar{w}_1] &= E[\bar{w}_2] = 0 \\ E[\bar{w}_1^2] &= V_1 \\ E[\bar{w}_2^2] &= V_2 \\ E[\bar{w}_1 \bar{w}_2] &= V_{12} \\ \text{If } i + j \geq 3, \text{ then } E[\bar{w}_1^i \bar{w}_2^j] &= 0 \end{aligned}$$

Now we can finally calculate the expected values of both the right-hand sides and the left-hand sides of the update rules described in (8,9). We will do all calculations for the limit of $t \rightarrow \infty$.

$$\begin{aligned} \alpha\left(\frac{1}{2} - \frac{\rho^2}{2}\right) + (-2 - 4\rho + \alpha(7 + 12\rho + 11\rho^2))\hat{V} + (-2 + \alpha(6 + 12\rho))\hat{V}_{12} &= 0 \\ \alpha\left(-\frac{1}{2} + \frac{\rho^2}{2}\right) + (-2 + \alpha(6 + 12\rho))\hat{V} + (-2 - 4\rho + \alpha(5 + 12\rho + 13\rho^2))\hat{V}_{12} &= 0 \end{aligned}$$

These two equations will allow us to calculate the steady state solutions $\hat{V}$ and $\hat{V}_{12}$. Since $\alpha \ll 1$, $\hat{V} \ll 1$ and $\hat{V}_{12} \ll 1$, we can neglect all terms that include $\alpha\hat{V}$ or $\alpha\hat{V}_{12}$. Thus, we obtain the following closed-form formula for calculating the variances at the steady state.

$$\hat{V} = \alpha \frac{1 - \rho^2}{8|\rho|}, \quad (10)$$

$$\hat{V}_{12} = -\hat{V}. \quad (11)$$

As we can see after sufficiently long iterations the correlation between w_1 and w_2 is equal to -1 . This means that they jump around the steady state solutions, always being on the different sides of it. Their individual variance is proportional to the learning rate, which means that the decaying learning rate once the steady state is reached will also proportionally decrease the variance of the parameters, thus attributing to the better convergence of the model.

5. Experiments

To validate the results of (10) and (11), we have created a simple experimental setup. Initially, we set $w_1 = 0$ and $w_2 = 1$. At each iteration, we generate a new data point from the bi-variate Gaussian distribution (2). Then we train for sufficiently long iterations until the steady state distribution is reached. We repeat this training process from scratch 500 times, save the final weights after each training process and calculate the variance of these 500 weights². Then we repeat this same process for different learning rates to capture the relation between the final variance of the weights and the learning rate α .

²Since the sign of the steady state solution depends on the initial weights every time we set the same value for the weights at the beginning. The same variance will be obtained for other initial conditions as well, while the mean steady state value may differ in sign.

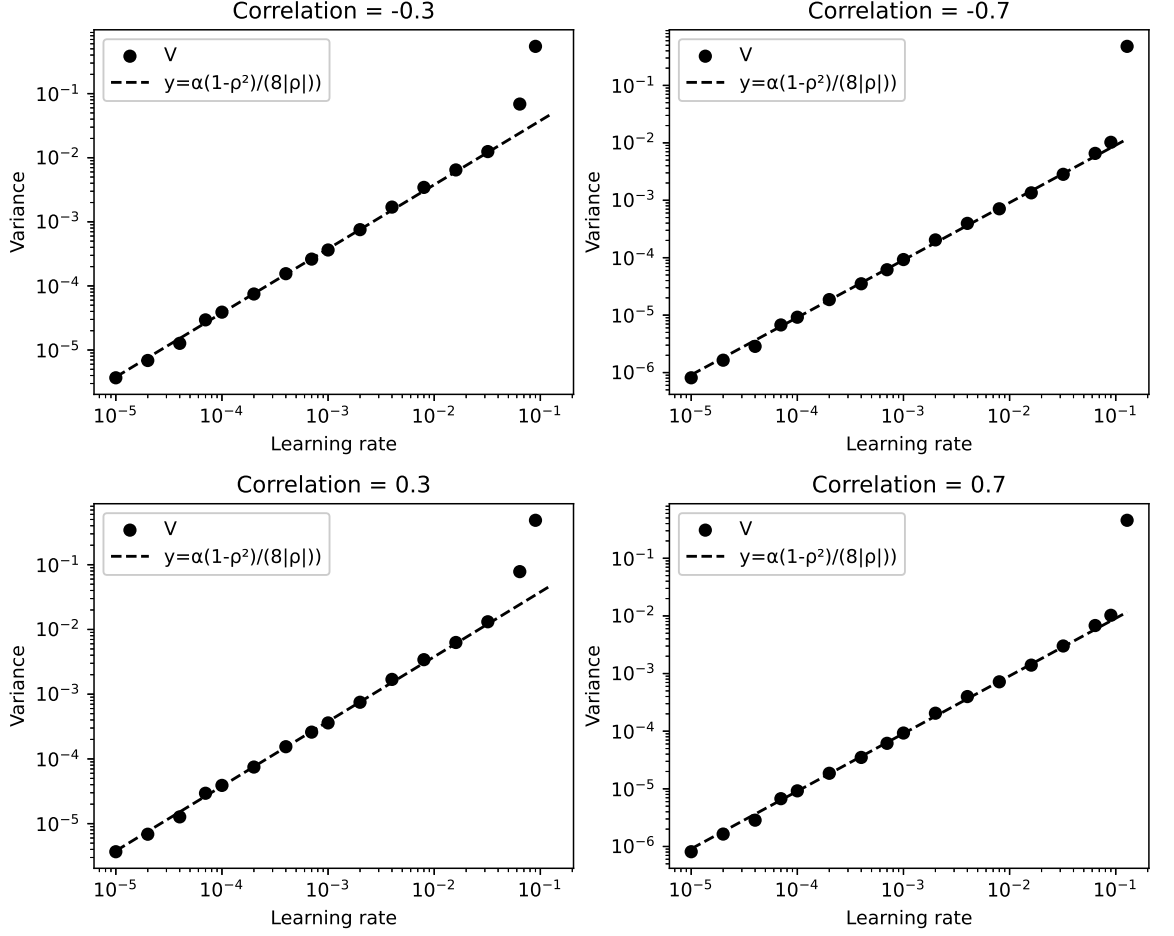


Fig. 1. If we train the system from scratch many times and calculate the variance of the final weight w_1 , it will be very close to the analytically calculated value of (10) for learning rates that are sufficiently small for converging. Here it is checked for 4 different correlation coefficients, both positive and negative.

This same process is repeated for different sets of data points generated from multivariate normal distribution with different correlation coefficients ρ to check the dependence of variance on correlation found in (10). Comparisons of the variances with the analytic values obtained in the previous section are represented in Fig. 1.³

6. Conclusion

This article proposes an analysis of parameter variance in Oja's Rule using moment closure approximation. The steady state variance is studied, leading to a closed-form equation connecting variance to the learning rate by a linear relation.

³The results of those experiments can be reproduced by following the steps at <https://github.com/edgarvardanyan/oja-variance>.

A key finding is the linear relationship between parameter variance and learning rate, showing that variance measures convergence. For small learning rates, variance is directly proportional to the learning rate, derived from a simple closed-form equation and validated through simulations across various input correlations.

These results have potential applications in optimizing learning rate schedulers and algorithms by controlling variance, thus improving convergence efficiency without extra computational cost. This can guide unsupervised learning models based on the Ojas Rule in achieving better results without a loss function.

As we can see, once the learning rate is small enough for the convergence of the model (usually achieved with $\alpha < 0.1$ for the provided synthetic data), our closed-form formula is able to estimate the final variance of parameters with good enough accuracy.

References

- [1] D. O. Hebb, *The Organization of Behavior: A Neuropsychological Theory*. Wiley, New York, 1949.
- [2] S. Nabavi, R. Fox, C. D. Proulx, J. Y. Lin, R. Y. Tsien, and R. Malinow, “Engineering a memory with LTD and LTP,” *Nature*, vol. 511, pp. 348–352, 2014.
- [3] E. Kuriscak, P. Marsalek, J. Stroffek, and P. G. Toth, “Biological context of Hebb learning in artificial neural networks: A review,” *Neurocomputing*, vol. 152, pp. 27–35, 2015.
- [4] G. Amato *et al.*, “Hebbian learning meets deep convolutional neural networks,” in *Image Analysis and Processing – ICIAP 2019*, E. Ricci *et al.*, Eds. Springer International Publishing, Cham, pp. 324–334, 2019.
- [5] J. P. Johansen *et al.*, “Hebbian and neuromodulatory mechanisms interact to trigger associative memory formation,” *Proceedings of the National Academy of Sciences*, vol. 111, no. 51, pp. E5584–E5592, 2014.
- [6] T. J. Sejnowski and G. Tesauro, “The Hebb rule for synaptic plasticity: Algorithms and implementations,” in *Neural Models of Plasticity*, J. H. Byrne and W. O. Berry, Eds. Academic Press, pp. 94–103, 1989.
- [7] E. L. Bienenstock *et al.*, “Theory for the development of neuron selectivity: Orientation specificity and binocular interaction in visual cortex,” *Journal of Neuroscience*, vol. 2, no. 1, pp. 32–48, 1982.
- [8] T. J. Sejnowski, “Storing covariance with nonlinearly interacting neurons,” *Journal of Mathematical Biology*, vol. 4, no. 4, pp. 303–321, 1977.
- [9] E. Oja, “A simplified neuron model as a principal component analyzer,” *Journal of Mathematical Biology*, vol. 15, no. 3, pp. 267–273, 1982.
- [10] L. N. Smith, “Cyclical learning rates for training neural networks,” in *Proceedings of the IEEE Winter Conference on Applications of Computer Vision (WACV)*, pp. 464–472, 2017.
- [11] E. Vardanyan and D. B. Saakian, “The analytical dynamics of the finite population evolution games,” *Physica A: Statistical Mechanics and its Applications*, vol. 553, p. 124233, 2020.

- [12] E. Vardanyan, E. Koonin, and D. B. Saakian, “Analysis of finite population evolution models using a moment closure approximation,” *Journal of the Physical Society of Japan*, vol. 90, no. 1, p. 014801, 2021.
- [13] V. Galstyan and D. B. Saakian, “Quantifying the stochasticity of policy parameters in reinforcement learning problems,” *Physical Review E*, vol. 107, p. 034112, 2023.
- [14] E. Oja, “Principal components, minor components, and linear neural networks,” *Neural Networks*, vol. 5, no. 6, pp. 927–935, 1992.
- [15] A. Hyvriinen and E. Oja, “Independent component analysis: Algorithms and applications,” *Neural Networks*, vol. 13, no. 4–5, pp. 411–430, 2000.
- [16] W. Gerstner, W. M. Kistler, R. Naud, and L. Paninski, *Neuronal Dynamics: From Single Neurons to Networks and Models of Cognition*. Cambridge University Press, 2014.
- [17] L. Isserlis, “On a formula for the product-moment coefficient of any order of a normal frequency distribution in any number of variables,” *Biometrika*, vol. 12, no. 1–2, pp. 134–139, 1918.
- [18] L. Isserlis, “On certain probable errors and correlation coefficients of multiple frequency distributions with skew regression,” *Biometrika*, vol. 11, no. 3, pp. 185–190, 1916.

ՀԵՐՔՅԱՆ ՈՍՏՈՑՄԱՆ ԿԱՅՈՒՆ ՎԻՃԱԿԻ ՂԻՍԱԿԵՐՍԻԱՅԻ ՎԵՐԼՈՒԾՈՒԹՅՈՒՆ՝ ՄՈՄԵՆՏՆԵՐԻ ՓԱԿՄԱՆ ՄՈՏԱՐԿՄԱՆՔ

Էդգար Ա. Վարդանյան

Հայ-Ռուսական համալսարան, Երևան, Հայաստան
e-mail: edgarvardanyan1999@gmail.com

Ամփոփում

ՀԵՐՔՅԱՆ ուսուցումը, որը ներդրումային ցանցերի կարևոր գաղափար է, հանդիսանում է տարբեր ուսուցման ալգորիթմների հիմքը, որոնք մոդելավորում են ներդրումային կապերի, այսպես կոչված սինապսների, ադապտացիան: Այս մոդելներից մեկն է Օյայի կանոնը, որը կարևոր օրինակ է՝ ապահովելով արժեքավոր պատկերացումներ ինքնասովորեցման ալգորիթմների դինամիկայի մասին: Լավ հայտնի է այն փաստը, որ միաշերտ ցանցի վերջնական կայուն վիճակի լուծումը, որը սովորում է Օյայի կանոնի կիրառմամբ, համընկնում է հիմնական բաղադրիչների վերլուծության լուծման հետ: Սակայն այն, թե ինչպես կարող է ուսուցման արագությունը ազդել վերջնական պարամետրերի տատանման վրա, քիչ է ուսումնասիրված: Այս հոդվածում մենք ուսումնասիրում ենք, թե ինչպես տարբեր ուսուցման արագություններ կարող են ազդել Օյայի կանոնի պարամետրերի ՂԻՍԱԿԵՐՍԻԱՅԻ վրա՝ օգտագործելով մոմենտների փակման մոտարկումը: Կենտրոնանալով տատանման վրա՝ մենք առաջարկում ենք նոր հեռանկարներ Օյայի կանոնի վարքագծի վերաբերյալ՝ տարբեր պայմաններում: Մենք ստանում ենք փակ տեսքով հավասարում, որը կապում է պարամետրերի տատանման մեծությունը ուսուցման արագության հետ և ցույց է տալիս,

որ դրանց միջև հարաբերությունը գծային է: Սա արժեքավոր պատկերացումներ է տալիս, որոնք կարող են օգնել օպտիմալացնել Հեբբյան մոդելների ուսուցման գործընթացը:

Բանալի բառեր` Հեբբյանի ուսուցում, ուսուցման արագություն, նեյրոնային ցանցեր, մոմենտների փակման մոտեցում:

Анализ устойчивой дисперсии в обучении Хебба: ПОДХОД МОМЕНТНОГО ЗАМЫКАНИЯ

Едгар А. Варданын

Институт проблем информатики и автоматизации НАН РА, Ереван, Армения

e-mail: edgarvardanyan1999@gmail.com

Аннотация

Обучение Хебба, важное понятие в нейронных сетях, является основой для различных алгоритмов обучения, моделирующих адаптацию нейронных соединений, также известных как синапсы. Среди этих моделей выделяется правило Ойи, которое предоставляет ценные сведения о динамике алгоритмов обучения без учителя. Хорошо известно, что итоговое устойчивое состояние однослойной сети, обучающейся по правилу Ойи, совпадает с решением метода главных компонент. Однако влияние скорости обучения на дисперсию конечных параметров изучено недостаточно. В данной работе мы исследуем, как различные скорости обучения влияют на дисперсию параметров в правиле Ойи, используя подход моментного замыкания. Сфокусировавшись на дисперсии, мы предлагаем новые взгляды на поведение правила Ойи в разных условиях. Мы выводим аналитическое уравнение, связывающее дисперсию параметров со скоростью обучения, и показываем, что эта зависимость является линейной. Это предоставляет ценные данные, которые могут помочь в оптимизации процесса обучения моделей Хебба.

Ключевые слова: правило Ойи, обучение Хебба, скорость обучения, нейронные сети, подход моментного замыкания.