



A bio-inspired model towards vocal gesture learning in songbird

Silvia Pagliarini, Xavier Hinaut, Arthur Leblois

► To cite this version:

Silvia Pagliarini, Xavier Hinaut, Arthur Leblois. A bio-inspired model towards vocal gesture learning in songbird. 2018 Joint IEEE International Conference on Development and Learning and Epigenetic Robotics (ICDL-EpiRob), Sep 2018, Tokyo, Japan. hal-01906459v2

HAL Id: hal-01906459

<https://inria.hal.science/hal-01906459v2>

Submitted on 22 Apr 2021

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

A bio-inspired model towards vocal gesture learning in songbird

Silvia Pagliarini

Inria Bordeaux Sud-Ouest, Talence, France.
LaBRI, UMR 5800, CNRS, Bordeaux INP,
IMN, UMR 5293, CNRS,
Université de Bordeaux, France.
silvia.pagliarini@inria.fr

Xavier Hinaut*

Inria Bordeaux Sud-Ouest, Talence, France.
LaBRI, UMR 5800, CNRS, Bordeaux INP,
IMN, UMR 5293, CNRS,
Université de Bordeaux, France.
xavier.hinaut@inria.fr

Arthur Leblois*

IMN, UMR 5293, CNRS,
Université de Bordeaux, France.
arthur.leblois@u-bordeaux.fr

Abstract—The paper proposes a bio-inspired model for an imitative sensorimotor learning, which aims at building a map between the sensory representations of gestures (sensory targets) and their underlying motor pattern through random exploration of the motor space. An example of such learning process occurs during vocal learning in humans or birds, when young subjects babble and learn to copy previously heard adult vocalizations. Previous work has suggested that a simple Hebbian learning rule allows perfect imitation when sensory feedback is a purely linear function of the motor pattern underlying movement production. We aim at generalizing this model to the more realistic case where sensory responses are sparse and non-linear. To this end, we explore the performance of various learning rules and normalizations and discuss their biological relevance. Importantly, the proposed model is robust whatever normalization is chosen. We show that both the imitation quality and the convergence time are highly dependent on the sensory selectivity and dimension of the motor representation.

I. INTRODUCTION

Imitative sensorimotor learning can be thought as a control problem aiming to map the sensory input into a motor output. For example, humans and songbirds learn to produce species-specific vocalizations as juveniles by imitating surrounding adults. The vocal learning process displays distinct (although partially overlapping) processes. First, during a sensory learning phase, young subjects memorize adult vocalizations, and build neuronal representations of their species vocal gestures. Then, in a sensorimotor phase, they start vocalizing and progressively converge to a good imitation of previously experienced vocalizations. It is believed that during the early phase of this sensorimotor process, called babbling, the subject maps representation of sensory (auditory) targets to the corresponding motor commands. In other words, the subject learns an inverse model. In this study, we assume the auditory selective responses to be already in place and investigate biological plausible mechanisms to learn an inverse model enabling this sensorimotor mapping.

An interesting property of some neurons in the brain is the ability to respond similarly during the production or observation (or listening for vocal gestures) of a given movement. This property has been linked with imitative sensorimotor learning

and internal models (inverse or forward model). These neurons are referred to as mirror neuron [1], [2], [3]. While forward models describe a causal relationship between the sensory input and the motor system, inverse models have the aim to provide an appropriate motor command to a given state of the motor system. Although building an inverse model is easier when a forward model is available as proposed in [4], [5], inverse models could be enough to bootstrap the development of a simple computational mechanism, describing a memory system composed by the motor plan and the sensory stimulus [6]. The importance of computational models involving mirrors systems and learning processes has been stressed by Oztog et al. [1], [6], and the "Mirror-system hypothesis" stated by Arbib in [7] links mirror neurons with the emergence of human language during evolution. Songbirds learn their vocalization by imitation through a vocal learning process that very much resembles speech learning in human babies [8]. In the brain of songbirds, a part of the basal ganglia-thalamo-cortical circuitry is devoted to song learning in juveniles and plasticity in adults. This circuit is homologous to the basal ganglia circuits responsible for motor learning in mammals, and involved in speech learning in humans [9], [10]. Moreover, this song-related BG-thalamo-cortical circuit in birds receives input from mirror neurons [2]. Therefore, the brain circuits responsible for avian song learning represent an ideal framework to study the neural mechanisms underlying imitative learning.

A theoretical model describing the implementation of an inverse model between auditory and motor areas through associative learning has been proposed recently [11], [12]. The model is based on a simple Hebbian learning rule driven through random motor exploration and auditory feedback responses to this motor exploration. The proposed model assumes linearity for mathematical simplicity. As auditory responses in the brain are rather sparse and nonlinear we aim to extend the theoretical framework to a more realistic scenario [13].

We used an inverse model inspired by Hahnloser and Ganguli [11] to describe the interaction between two populations of neurons, one formed by motor neurons and another formed by sensory neurons. We first show that replacing the learning rule by a simple normalized Hebbian learning rule allows

* Corresponding authors that co-supervised the study.

rapid convergence in a simple non-linear model. We apply different normalizations in the learning rule. We then explore the influence of the sharpness of auditory selectivity in relation with the learning error after convergence and convergence time of the learning. Finally, we show how changing the number of sensory or motor dimensions modifies learning.

II. METHOD

A. Network and goal

The model includes two neural populations as shown in Fig. 1. The first layer is composed by afferent neurons and represents the sensory area. The second layer is composed by motor neurons, which represent the starting point for muscle activation, and thereby for movement production. The synaptic weights describing the strength of the connection between neurons are defined by matrix W .

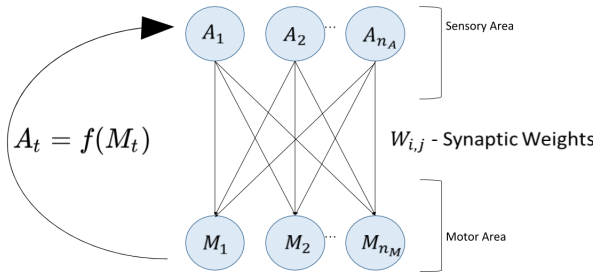


Fig. 1. **Neural network schema.** The network includes two neural populations: the first layer is composed by n_a afferent neurons and represents the sensory area, the second layer is composed by n_m motor neurons and represents the motor area. $W_{i,j}$ represents the synaptic connections between sensory and motor neurons. Below the network schema we highlight how we define the sensory response. At each time step t , the sensory response is a function of the motor output, that is $A_t = f(M_t)$.

A model describing sensorimotor phase of learning based on a network as in Fig. 1 has been previously proposed by Hahnloser and Ganguli [11]. Neurons are linear units and at each time t motor and auditory activity are defined as a n_m -dimensional vector M_t and a n_a -dimensional vector A_t , where n_m and n_a represent respectively the number of motor and auditory neurons in the network.

B. Auditory response

At each time t the auditory activity A_t is defined as a function of the motor pattern M_t at that particular time, that is $A_t = f(M_t)$. Hahnloser and Ganguli in [11] define the auditory activity as a linear function of M_t , that is

$$A_t = QM_t, \quad (1)$$

where M_t represents the motor pattern at time t and Q the linear map defining the auditory activity due to the auditory feedback driven by current motor activity.

We then introduce non-linearity in the sensory response to auditory feedback. To represent selective responses as observed in various high sensory brain areas (e.g. auditory

regions of the pallium in birds display responses selective to tutor syllables or to the bird's own syllables), we define the auditory activity A_t for each $j = 1, \dots, n_a$ neurons as a bell-shaped function around a target motor pattern:

$$A_{j_t} = \exp\left(\frac{-||M_j^* - M_t||^2}{2\sigma^2 n_m}\right), \quad (2)$$

where σ represents the selectivity tuning width, n_m the number of motor neurons belonging to the network, M_t the motor pattern at time t and M^* the center of the auditory activity.

C. Learning process

Learning is driven by the Hebbian learning rule

$$\Delta W_t \propto \eta M_t A_t, \quad (3)$$

where W_t represents the synaptic weights between sensory and motor neurons, η the learning rate, M_t the motor pattern at time t and A_t the auditory activity at time t .

Synaptic weights $W_{t=t_0}$ between sensory and motor neurons are initially weak and increase according with (3) during learning until a certain time $t = t_f$, as

$$W_t = W_{t-1} + \Delta W_t, \quad (4)$$

where W_t represents the synaptic connections between sensory and motor neurons and ΔW_t is defined by (3).

However, synaptic weights have an upper boundary due to biological limitations (maximal number of synaptic receptors or neurotransmitters released). This can be introduced by a normalization either on the synaptic weights $W_{i,j}$ or on their variation $\Delta W_{i,j}$.

Hahnloser and Ganguli in [11] proposed a postdictive Hebbian learning rule given by

$$\Delta W_t = \eta(M_t - W_{t-1}A_t)A_t^T, \quad (5)$$

where W_t represents the synaptic weights between sensory and motor neurons, η the learning rate, M_t the motor pattern and A_t is defined by Eq. (1). Here the apex T indicates the transpose of the vector A_t .

In our model, we kept the basic Hebbian rule and we tested three normalizations in two different cases: a normalization over motor neurons (over all targets of one postsynaptic neuron) and a normalization over auditory neurons (over all inputs of one presynaptic neuron).

In practice, the two types of normalization are respectively implemented by normalizing over the lines or columns of the weights matrix W . The aim of the normalization is to bound either the mean of each column or line of $W_{i,j}$ or the euclidean norm of each column or line of $W_{i,j}$ to a maximum of 1. Considering the case of normalizing with respect to auditory neurons and pushing the mean of every column of W to a maximum of 1, the applied normalizations are the following:

- Maximum weights normalization

$$W_{i,j} = \frac{n_m W_{i,j}}{\sum_i W_{i,j}}, \quad (6)$$

- Supremum weights normalization

$$W_{i,j} = \begin{cases} W_{i,j} & \frac{\sum_i W_{i,j}}{n_m} < 1, \\ \frac{n_m W_{i,j}}{\sum_i W_{i,j}} & \text{otherwise,} \end{cases} \quad (7)$$

- Decreasing factor normalization

$$\Delta W_{i,j} = \eta M_t A_t \left(1 - \frac{\sum_i W_{i,j}}{n_m} \right). \quad (8)$$

Here $W_{i,j}$ represents the synaptic connections between sensory neuron j and motor neurons $i = 1, \dots, n_m$, where n_m is the number of motor neurons.

To obtain the normalization over the motor neurons and pushing the mean of each line of W to a maximum of 1, it is enough to change the index i for the index j and use the auditory dimension n_a . In order to use the norm instead of the mean in the definition of the normalizations it is enough to introduce the norm on the column or the norm on the line of W in place of the mean.

At the same time, normalizing synaptic weights forces us to also normalize the motor target M^* , which represents what the model would learn at $t = t_f$. This is equivalent to a reduction of the target motor space (as it introduces a constraint on the final output of the model).

D. Simulation details

Each sensory neuron contains the response to a motor performance, and this is represented by the motor target M^* . We did not consider any strategy for the exploration. That is, at each time step t we simply considered the case of a random motor exploration M_t on which the auditory selectivity depends.

In the Hahnloser-Ganguli model there exist a direct map which defines the motor activity from the auditory activity as $M_t = Q^d A_t$. For the inverse model we need to define A_t in dependence on M_t as in Eq. (1). That is the inverse matrix of Q , i.e Q^d represents the motor target M^* , that is the ideal motor activity which the model should have learned once learning phase has ended. The goal then is to activate each sensory neuron can drive M_j^* through W , such that

$$WA^* \longrightarrow M^*, \quad (9)$$

where $A^* = A_0 I$ defines the ideal auditory activity. We fixed $A_0 = 1$.

Given the motor targets M^* , at each time step t , the distance between what the model actually learned and what it should have learned is defined as

$$d_t = \frac{\|M^* - W_t A^*\|}{n_m}, \quad (10)$$

where M^* represents the motor target, W_t the synaptic weights matrix, A^* the ideal auditory activity and n_m the number of motor neurons.

We defined the convergence time τ as the number of time steps at which the updates of the weights are small enough. That is, the distance between what the model targets should have learned and what he effectively learned reaches a plateau.

After have chosen $\epsilon = 1$ at $t = t_0$, given an interval of time $\Delta t = [k, k + 2N]$ of measure $2N$, we defined

$$\epsilon = \frac{1}{2N} \left(\sum_{k+N}^{k+2N} d_t - \sum_k^{k+N} d_t \right), \quad (11)$$

and we used it as threshold, in a way that a particular experiment stops either if ϵ reaches the value $\epsilon^* = 10^{-9}$ either if it goes until a fixed time $t = t_f$. We tested several values for the tuning selectivity width, i.e. $\sigma = [0.02, 0.05, 0.1, 0.2, 0.3, 0.5, 0.7]$, varying in this way the auditory selectivity. We used for almost every value an interval of length $N = 400$. However, since the distance evolves very slowly, for small values of σ this interval is not large enough. For instance, an interval with $N = 500000$ has been used for $\sigma = 0.02$.

III. RESULTS

A. Simple model with linear auditory responses

Fig. 2 shows the evolution in time of the smooth average distance between WA^* and the target motor pattern M^* , defined by Eq. (10) for each simulation. In the linear version of the model (blue line), the auditory activity is a linear function of the motor production, as in Eq. (1), and the postdictive Hebbian rule defined by Eq. (5) guides learning. As expected by the theory [11], the distance between WA^* and the target motor pattern M^* (which is the inverse of Q) converges exponentially to zero. The learning rule therefore ensures proper learning of the inverse model. In contrast, when auditory feedback is non-linear (orange line in Fig. 2), where the auditory activity is defined by Eq. (2), the postdictive Hebbian learning rule does not allow convergence, and the distance between WA^* and the target motor pattern M^* rather diverges.

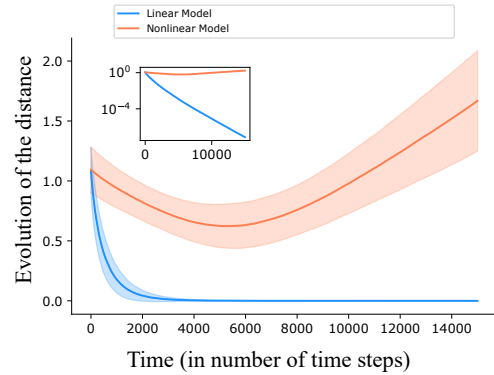


Fig. 2. **Ganguli-Hahnloser linear and nonlinear model.** Evolution in time of the smooth average distance between WA^* and the target motor pattern M^* , computed over 50 simulations. Comparison between the Hahnloser-Ganguli linear model (in blue) where the auditory response is defined by Eq. (1) and the nonlinear version of Hahnloser-Ganguli model (in orange) where the auditory response is defined by Eq. (2). In both cases learning is driven by the postdictive Hebbian learning rule in Eq. (5) and weights are updated following Eq. (4). To highlight the behavior of the linear model, the same comparison using a log scale is shown in the box. Parameters value: $t_f = 1.5 \times 10^4$, $n_m = n_a = 3$, $\eta = 0.01$, $\sigma = 0.1$, $C = 20$.

Define the auditory activity as in Eq. (1) means that the matrix Q needs to be invertible to reach convergence. We also don't know how to define the exploration space starting from the map Q . We assumed as given the space of the inverse of Q , which means that we solved at each time step t the linear system $Q^{-1}A_t = M_t$, where $Q^{-1} = M^*$. At the same time to solve a linear system means to face the problem of invertibility of matrix M^* . Indeed, if the matrix M^* is singular or close to be singular, then the numerical algorithm is not longer working. To avoid an ill-posed problem we added a condition by computing the condition number¹ of the matrix M^* , forcing it to be such that

$$k(M^*) < C, \quad (12)$$

where M^* represents the motor target and C is a positive constant belonging to $[1, +\infty]$. In this way we simulated all the simulations in Fig. 2 without having ill-posed problems. Without the application of this condition, simulations were often diverging because of the divergence of M^{*-1} . A nonlinear auditory response defined by Eq. (2) enables to avoid ill-posed problems. At the same time (as underlined in the small box with logarithmic scale in Fig. 2) it leads to a divergence in the distance between the target motor pattern M^* and WA^* . That is, the model does not learn anymore after the introduction of nonlinearity. So far, instead of keeping the postdictive Hebbian rule proposed by Ganguli and Hahnloser in [11], we used a traditional Hebbian rule to drive learning and tried to face the problem by applying other types of normalization.

B. A nonlinear auditory response

Fig. 3 shows the evolution of the distance between WA^* and the target motor pattern M^* for one example neuron. The initially weak synaptic connections evolves following the Hebbian learning rule given by Eq. (3) and finally approaches M^* . To obtain these results we applied the normalization defined by Eq. (8). By introducing nonlinearity, as shown by Fig. 2, the Ganguli-Hahnloser model does not converge anymore.

We tested three different normalizations, given by Eq. (6), Eq. (7) and Eq. (8). The upper panel of Fig. 4 shows the comparison between the three normalizations with respect to auditory neurons. That is, with respect to the columns of W . Normalizations given by Eq. (6) and Eq. (7) are applied directly to the weights matrix, which gives a faster convergence but a lost in smoothness. Normalization given by Eq. (8) is applied to the variation of the weights, by multiplying its classical definition by a decreasing factor. This means that the variation is smaller and smaller as the weights approaches the target, which results in a smooth trend of the distance curve. To highlight better the difference between normalization

¹ Given a general linear system $Ax = b$, its condition number is defined as $k(A) = \xi_{\max}(A)/\xi_{\min}(A)$, where $\xi_{\max}(A)$ and $\xi_{\min}(A)$ are respectively the maximal and minimal singular values of A . The value $k(A)$ represents the variability of the solution, so how much the solution x changes consequently to a change in b . The lower bound for the condition number is $k(A) = 1$, whereas it can reaches the value $k(A) = \infty$ if the matrix A is singular.

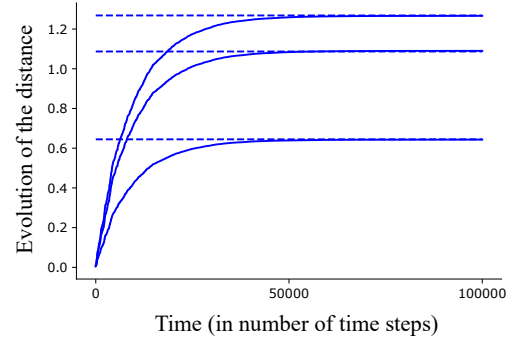


Fig. 3. **Evolution of synaptic weights in time relative to a single auditory neuron and distance from the target motor pattern (three neurons example).** Evolution in time of the synaptic weights W (continuous blue line) and the target motor pattern M^* (dashed blue line) for one example neuron. Each auditory neuron is composed by three components represented by the three lines. Here the auditory activity is defined by Eq. (2) and learning is driven by the Hebbian learning rule in Eq. (3). At each time steps weights are updated following Eq. (4) and the normalization defined by Eq. (8) has been applied. Parameters value: $t_f = 1 * 10^5$, $n_m = 3$, $n_a = 1$, $\eta = 0.01$, $\sigma = 0.1$.

over auditory and motor neurons, the bottom panel of Fig. 4 shows the comparison between the normalization given by Eq. (8) applied in its mean and norm version. As it is shown, a normalization over auditory neurons works better than the same normalization over motor neurons in the sense that the distance between WA^* and M^* is lower if the normalization is applied over the auditory neurons than if the normalization is applied over the motor neurons.

C. Auditory selectivity effect

Auditory selectivity impact on the learning can be observed by varying its value and by observing both the convergence time τ both the distance at $t = \tau$ between WA^* and M^* . Fig. 5 shows how the mean convergence time τ and the mean distance d_t at $t = \tau$ depends on the auditory selectivity. As the tuning selectivity width σ increases, a decreasing in the mean convergence time and an increasing in the mean distance can be observed. In particular, for the value $\sigma = 0.02$ convergence time is not fully correct because many simulation reached a fixed time $t_f = 2 * 10^7$ before having reached convergence. This is displayed on the plot by the first dashed part of the red line.

D. Varying network dimensions

The quality of learning in terms of the distance at the convergence time $t = \tau$ and how slow learning develops can be investigated by varying the network dimension. Firstly, the number of auditory neurons has been kept fixed at the value $n_a = 3$, and the number of motor neurons varied as $n_m = [2, 3, 4, 5, 6, 7]$. Then, viceversa, the number of motor neurons has been kept fixed at $n_m = 3$ and the number of the auditory neurons varied using the same values as before. Fig. 6 shows the effect of changes in the network dimension respectively on the mean convergence time τ and on the mean distance at $t = \tau$, computed over 50 simulations. The upper panel shows how, keeping fixed the motor dimension,

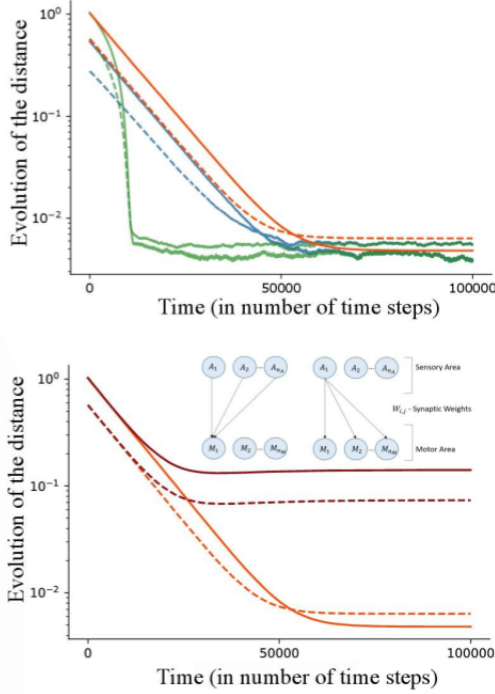


Fig. 4. **Comparison between different types of normalization: evolution in time of the distance.** Evolution in time of the smooth average distance between WA^* and the target motor pattern M^* , computed over 50 simulations. (Top) Comparison between the model normalized by the maximum weights normalization in Eq. (6) and the corresponding norm version (respectively the continuous blue line and the dashed blue line), the model normalized by the supremum weights normalization in Eq. (7) and the corresponding norm version (respectively the continuous green line and the dashed green line), the model normalized by the decreasing factor normalization in Eq. (8) and the corresponding norm version (respectively the continuous red line and the dashed red line). All the normalizations have been taken with respect to auditory neurons. (Bottom) Comparison between the model normalized by the decreasing factor normalization in Eq. (8) with respect to auditory neurons (red lines) and with respect to motor neurons (dark red lines). Comparison between the normalization applied using the mean of W (continuous lines) and the norm of W (dashed lines). Here the auditory activity is defined by Eq. (2) and learning is driven by the Hebbian learning rule in Eq. (3). At each time steps weights are updated following Eq. (4). Parameters value: $t_f = 1 * 10^5$, $n_m = n_a = 3$, $\eta = 0.01$, $\sigma = 0.1$.

there is not an evidence of network dimension effect on the mean convergence time. Viceversa, keeping fixed the auditory dimension the learning slows down as the motor dimension increases. However, the mean distance at $t = \tau$ is not affected by any change in the neural network dimensions, as shown in the bottom panel of Fig. 6. More details are available at <https://github.com/spagliarini/2018-ICDL-EPIROB>.

IV. DISCUSSION

Hahnloser and Ganguli [11] proposed a simple mathematical framework to approach the sensorimotor learning problem in songbirds. It is based on a linear auditory activity and a postdictive Hebbian learning rule. Linearity in the auditory activity makes the theoretical investigation possible but is not biologically realistic. To be invertible, the matrix Q for auditory response must be squared, which means that auditory

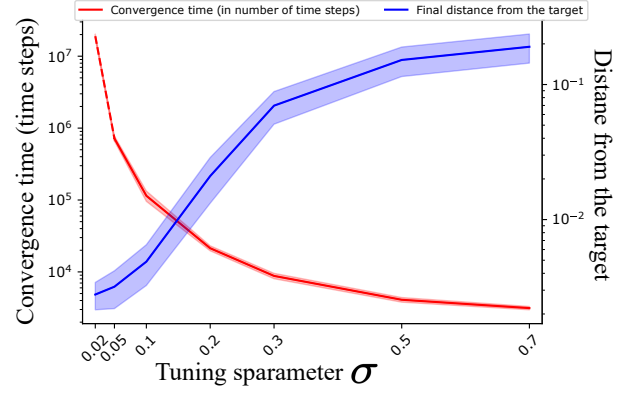


Fig. 5. **Effect of the auditory selectivity on convergence time and distance.** Auditory selectivity impact on the convergence time (in red) and on the distance between WA^* at the convergence time τ and the target motor pattern M^* (in blue), computed over 50 simulations. The first dashed part of the red line underlines the fact that for $\sigma = 0.02$ not all the simulations converges before having reach a fixed simulation exit time $t_f = 2 * 10^7$. Here the auditory activity is defined by Eq. (2) and learning is driven by the Hebbian learning rule in Eq. (3). At each time steps weights are updated following Eq. (4). Parameters value: $\sigma = [0.02, 0.05, 0.1, 0.2, 0.3, 0.5, 0.7]$, $n_m = n_a = 3$, $\eta = 0.01$, $\epsilon^* = 10^{-9}$. We applied the decreasing factor normalization given by Eq. (8). To exit the simulations we compute ϵ as in Eq. (11). We used for almost every value an interval with $N = 400$. However, since the distance evolves very slowly, for small values of σ this interval is not large enough. For instance, an interval with $N = 500000$ has been used for $\sigma = 0.02$.

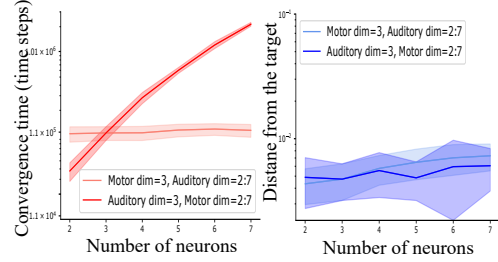


Fig. 6. **Effect of network dimension on the convergence time and the distance.** (Left) Network dimensions effect on the mean convergence time τ and (Right) on the mean distance between WA^* at time $t = \tau$ and the target motor pattern M^* , computed over 50 simulations. The dark red line refers to a network where the number of auditory neurons has been kept fixed at $n_a = 3$, whereas the light red line refers to a network where the number of motor neurons has been kept fixed at $n_m = 3$. In both cases, the second dimension has been varied as $[2, 3, 4, 5, 6, 7]$. Parameters value: $\eta = 0.01$, $\sigma = 0.1$. Here we applied the decreasing factor normalization given by Eq. (8). To exit the simulations we compute ϵ as in Eq. (11) with $N = 400$.

and motor dimensions have to be equal. Moreover, numerical implementation of the learning algorithm proposed by Hahnloser and Ganguli requires to invert the auditory response matrix (Q) to determine the range of motor output required for proper exploration and learning. As there is no general method to invert a random matrix Q , a numerical implementation of the model requires to set a specific Q that can be inverted. Finally, the postdictive learning rule only works for the linear model and it is not clear whether biologically realistic learning rules can still lead to convergence or near-convergence in the case of non-linear auditory feedback.

As Hebbian or associative learning rules are imple-

mented through activity-dependent synapse-specific increases in synaptic weights (synaptic potentiation), that must be augmented by global processes that regulate overall levels of neuronal and network activity to prevent explosion of synaptic weights [14]. Regulatory processes are often as important as the more intensively studied Hebbian processes in determining the consequences of synaptic plasticity for network function. Setting an upper bound on the total synaptic weights to or from a given neuron may also reflect the biological limitation of synaptic connections: limits are imposed on their growth due to the limited quantity of available material (receptors, neurotransmitters, ...). The introduction of normalization on the weights or on their variations is a simple solution to this problem. Several forms of normalization were considered here to take into account this biological limitation. While the linear model of Hahnloser and Ganguli [11] converges for the postdictive learning rule described there, we show that near-convergence can be achieved with multiple normalization rules added to a simple and typical associative learning rule given by Eq. (3). Convergence time, final distance from motor output to target, and smoothness of the distance evolution through time all depend on the specific normalization used. However, it is important to notice that the final "error" (distance from motor final weights to target motor pattern) does not vary much with normalization, assuming it is applied on all synaptic outputs from a given presynaptic (auditory) neuron.

In most of the simulations we focused on normalization given by the decreasing factor normalization in Eq. (8). We kept this normalization for all our analyses because it gives better performance (i.e. low "error"). We noticed that when this normalization is applied, it gives better performances over auditory neurons (presynaptic) than over motor neurons (post-synaptic), despite the fact that a normalization with respect to motor neurons (regulated at the level of the post-synaptic neuron) may be more biologically plausible. Although other forms of plasticity exist, including presynaptic modulation of synaptic strength, classic long-term potentiation/depression (LTP/LTD) mechanisms mostly involve postsynaptic receptor reorganization.

A remaining open question is related to the final value (after convergence) of the distance between the target motor pattern and what the model actually learned. Future work is needed to determine the factors that determine this final error and how it can be reduced. One possibility is that various motor targets (one for each auditory neuron) may interfere during learning, leading to imperfect copies. However, our preliminary experiments didn't show that interference had an influence on the final error.

We investigated how learning depends on the auditory selectivity observing that as the tuning selectivity width σ increases as the final distance between weight matrix and motor target (error) increases while convergence time decreases. There is therefore a trade-off between learning speed and accuracy that can be balanced through the selectivity of auditory neurons. One way to make learning both fast and accurate could be to start the learning process with a large tuning width (low

selectivity) and to decrease it progressively as learning goes on. Interestingly, in many songbird species (including the well-studied zebra finches), sensory learning overlaps with sensorimotor learning, and the auditory selectivity therefore develops during the early sensorimotor phase.

Finally, taking into account the influence of the dimensions (i.e. number of units) of the sensory and motor layers we noticed that motor dimension has a strong influence on convergence time. This strong influence comes from the fact that for lower values of σ the distance between the target and WA^* tends to decrease much more slowly. In our network we are considering a motor output that does not distinguishes muscle control and sound production. It is not clear which of these two components is responsible for the high increase in convergence time. A model displaying a motor output and a sound generating system is needed to resolve this question. Future work could include (1) the addition of an artificial syrinx model as motor output and more auditive like feature selectivity in the auditory layer, and (2) the influence of different exploration methods such as goal-directed exploration.

ACKNOWLEDGMENT

We would like to thank Camille Soetaert and Jean-Baptiste Zacchello for preliminary work done. We also thank Inria for the CORDI-S PhD fellowship grant.

REFERENCES

- [1] E. Oztop, M. Kawato, and M. Arbib. Mirror neurons and imitation: A computationally guided review. *Neural Networks*, 19(3):254–271, 2006.
- [2] J. F. Prather, S. Peters, S. Nowicki, and R. Mooney. Precise auditory-vocal mirroring in neurons for learned vocal communication. *Nature*, 451(7176):305, 2008.
- [3] N. Giret, J. Kornfeld, S. Ganguli, and R. HR Hahnloser. Evidence for a causal inverse model in an avian cortico-basal ganglia circuit. *PNAS*, 111(16):6063–6068, 2014.
- [4] M. I. Jordan and D. E. Rumelhart. Forward models: Supervised learning with a distal teacher. *Cognitive science*, 16(3):307–354, 1992.
- [5] A. K. Philippsen, R. F. Reinhart, and B. Wrede. Learning how to speak: Imitation-based refinement of syllable production in an articulatory-acoustic model. In *ICDL-Epirob, 2014*, pages 195–200. IEEE, 2014.
- [6] E. Oztop, M. Kawato, and M. A. Arbib. Mirror neurons: functions, mechanisms and models. *Neuroscience letters*, 540:43–55, 2013.
- [7] M. A. Arbib. fo the mirror system, imitation, and the evolution of language. *Imitation in animals and artifacts*, 229, 2002.
- [8] A. J. Doupe and P. K. Kuhl. Birdsong and human speech: common themes and mechanisms. *Annual review of neuroscience*, 22(1):567–631, 1999.
- [9] A. J. Doupe, D. J. Perkel, A. Reiner, and E. A. Stern. Birdbrains could teach basal ganglia research a new song. *Trends in neurosciences*, 28(7):353–363, 2005.
- [10] R. Mooney. Neural mechanisms for learned birdsong. *Learning & Memory*, 16(11):655–669, 2009.
- [11] R. Hahnloser and S. Ganguli. Vocal learning with inverse models. *Principles of Neural Coding*, pages 547–564, 2013.
- [12] A. Hanuschkin, S. Ganguli, and R. Hahnloser. A hebbian learning rule gives rise to mirror neurons and links them to control theoretic inverse models. *Frontiers in neural circuits*, 7:106, 2013.
- [13] R. HR Hahnloser and A. Kotowicz. Auditory representations and memory in birdsong learning. *Current opinion in neurobiology*, 20(3):332–339, 2010.
- [14] L. F. Abbott and S. B. Nelson. Synaptic plasticity: taming the beast. *Nature neuroscience*, 3(11s):1178, 2000.