

The impact of Inter Stimulus Interval on Semantic Priming: hysteresis or adaptation? A SOM neural network model

Valentina Gliozzi (valentina.gliozzi@unito.it)

Department of Computer Science and Center for Logic, Language and Cognition (LLC)
Università di Torino, Italy

Abstract

Recent results show that 18 months old infants are sensitive to taxonomic relations and that, similarly to adults, these relations are modulated by Inter Stimulus Interval. A very influential proposal in the distributed representations literature explains the impact of ISI on semantic priming as the result of a phenomenon called *hysteresis*. Here we propose that the same results could also be explained by the opposite phenomenon of adaptation. The existence of two possible explanations calls for more experiments to understand if hysteresis or adaptation can explain the role of ISI on semantic priming.

Keywords: computational modelling; infants; semantic priming; conceptual change

Introduction

Willits, Wojcik, Seidenberg, and Saffran (2013); Delle Luche, Durrant, Floccia, and Plunkett (2014); Plunkett, Delle Luche, Hills, and Floccia (2022) have shown that when listening to a sequence of spoken words, 18- and 24-month old infants are sensitive - among others - to taxonomic relations between subsequent words (as in “dog, pig, cat, sheep, ...”)¹. Similarly to adults (Alario, Segui, & Ferrand, 2000), infant sensitivity to taxonomic relations is modulated by the Inter Stimulus Interval (ISI): semantic (here taxonomic) priming can be observed at short ISI (400 ms.) while it fades away at long ISI (800 ms.).

Within the Distributed Representations Tradition (in which semantic representations are explicitly represented as sets of features) the impact of ISI on taxonomic priming is usually explained by the phenomenon of *hysteresis* (Plaut, 1995; Plaut & Booth, 2000). By hysteresis, semantic representations elicited by word utterances become fully activated only after a lapse of time. When fully activated, semantic representations become attractors and prevent the transition to other semantic representations, even taxonomically related, thus inhibiting semantic (taxonomic) priming.

In this paper we suggest that within the Distributed Representations Tradition, there is another possible explanation of the impact of ISI on semantic priming. We propose a neural network model that explains the impact of ISI on taxonomic priming by the opposite phenomenon of *adaptation*,

and more precisely adaptation of some semantic features. Adaptation is a biologically sound mechanism used by several neuro-cognitive models (e.g., Lerner, Bentin, & Shriki, 2012; Huber & O’Reilly, 2003; Treves, 2005), and it explains equally well the impact of ISI on semantic priming.

Both our model and Plaut (1995); Plaut and Booth (2000)’s model assume that semantic priming between taxonomically related representations depends on the similarity, or feature overlap, between these representations (the higher the similarity, the higher the priming). Both our model and Plaut (1995); Plaut and Booth (2000)’s model suggest that the Inter Stimulus Interval impacts the similarity degree between taxonomically related representations. However, Plaut (1995); Plaut and Booth (2000)’s model suggests that at long ISI the similarity between taxonomically related representations decreases (due to the activation of representational details that possibly differentiate the two representations). On the contrary, here we propose that by adaptation it is the distance between taxonomically *unrelated* representations to be impacted, and to decrease. The moral of the paper is that further experiments are needed to understand whether hysteresis or adaptation are responsible for the impact of ISI on semantic priming.

To illustrate our alternative explanation on the role of ISI on semantic priming, we propose a model based on self-organising maps and Hebbian links. These are widely recognised as psychologically plausible tools (Miikkulainen, 1997; Miikkulainen, Bednar, Choe, & Sirosh, 1997; Li, Farkas, & MacWhinney, 2004; Li, Zhao, X., & MacWhinney, 2007). Our model is an extension of (Mayor & Plunkett, 2010) model with activation dynamics. It directly mimics infants listening to a sequence of words and successfully replicates results by Delle Luche et al.’s (2014) and Plunkett et al.’s (2022).

Infant Experimental Data

In a series of studies on early semantic priming using a semantically-based version of the Head-turn Preference Procedure, Jusczyk and Aslin (1995), Delle Luche et al. (2014) and Plunkett et al. (2022) show that sensitivity to taxonomic relations is already present by 18 months of age. Delle Luche et al. (2014) demonstrate that 18 months olds listen longer to lists of words coming from the same taxonomic category (for instance, all animals, e.g., “cow, dog, sheep, cat ...”) than

¹Plunkett et al. (2022) also show that infants are sensitive to associative/thematic relations. We do not consider these thematic relations here. Associative relations, together with developmental issues are considered in a richer version of this paper under journal revision.

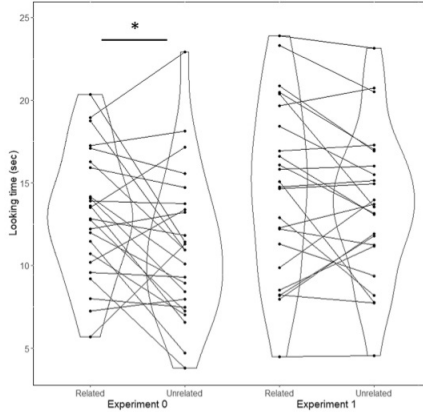


Figure 1: Results from Della Luche et al. (2014) replotted as Experiment Zero (SHORT ISI = 400 ms.), and from Plunkett et al. (2022), Experiment One (LONG ISI = 800 ms.). For each experiment, plots indicate the looking time distributions for the Related and Unrelated conditions (looking time is taken as a proxy of listening time). Significant differences ($p < 0.05$) between conditions are indicated as asterisks in each Experiment. Experiments Zero revealed significant differences in looking times, differently from Experiment One.

to lists of words coming from different categories (“nappy, boots, tummy, mouth, sock, ...”). Plunkett et al. (2022) investigate the role of ISI on taxonomic priming, and show that the effect disappears at long ISI, so that when ISI is extended to 800 ms. (instead of the 400 ms. considered by Della Luche et al., 2014) infant listening time was not significantly different between the taxonomically-related and taxonomically-unrelated conditions (see Figure 1: Experiment One).

The Model

Model Architecture

Our model is depicted in Figure 2². It is an extension of Mayor and Plunkett’s (2010) model of early word learning. As in Mayor and Plunkett’s (2010) model, our model contains two sub-networks which are self-organising maps (SOMs, Kohonen, 2001). These are an auditory map and a conceptual/semantic map. The auditory map receives and learns to represent auditory inputs. The conceptual/semantic map learns to represent semantic inputs.

As in Mayor and Plunkett’s (2010) model, each map independently learns to represent its input stimuli (auditory or semantic), giving rise to auditory representations and semantic representations, respectively. With training, similar stimuli are mapped onto close-by units of their map. In this way, taxonomically related semantic stimuli are mapped to close-by areas of the conceptual map, and auditory stimuli which are similar to each other occupy neighbouring regions of the auditory map.

²Our model is implemented in Matlab. We use `somtoolbox`

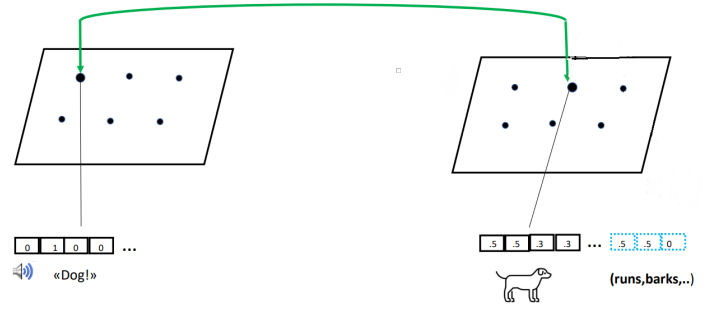


Figure 2: The model contains two pathways: auditory input provides activation to the auditory self-organising map, and semantic input provides activation to the semantic self-organising map. During the self-organising learning phase, the two maps separately learn to represent their respective stimuli. Auditory and semantic maps are associated through Hebbian connections, linking a word’s auditory representation to its semantic counterpart. These links are learnt from joint presentations of auditory and semantic stimuli (e.g., the word “dog” jointly presented with the semantic representation of a dog). In semantic representations we distinguish core features and extra features. We assume these two kinds of features adapt at a different speed.

Auditory and conceptual maps are linked by Hebbian connections, associating auditory units with their corresponding conceptual units e.g., the word ‘dog’ with the concept DOG (Figure 2).

With respect to Mayor and Plunkett’s (2010) model, we introduce an adaptation mechanism by which the activation of semantic representations diminishes with time.

As with the stimuli used in Della Luche et al. (2014) and Plunkett et al. (2022) infant studies, model stimuli (both auditory and semantic) belong to one of 4 superordinate categories: animals, clothes, food, and body parts. Each superordinate category has 8 basic sub-categories for a total of 32 basic categories.

In our model semantic stimuli are encoded by distributed representations made of 16-dimensional feature vectors. Our stimuli are artificially built in order to maximise inter-class distance and minimise intra-class distance. In our stimuli we differentiate between core features (the first 8 values) and extra features (the last 8 values). An examples of semantic stimulus (e.g., dog’s encoding) is:

0.34, 0.27, 0.41, 0.22, 0.38, 0.24, 0.36, 0.06, 0.7, 0.65, 0.12, 0.16, 0.06, 0.04, 0.02

For auditory stimuli we consider 32 stimuli corresponding to the 32 words considered. Each stimulus is encoded by a kind of one-hot encoding: each stimulus is characterised by a distinct set of three 1s and all other values set to 0. All word stimuli are equally distant from each other.

Training

In our model, both the semantic and the auditory maps are *self-organizing maps* (Kohonen, 2001), and consist of a set of neurons, or units, spatially organised in a grid, as in Figure 2. Each map unit u is associated with a weight vector w_u of the same dimensionality as its input vectors. This weight vector constitutes the representation (either semantic or auditory) associated to the unit. At the beginning of training, all weight vectors are initialised to small random values. During training, the input stimuli are presented in random order to the map. After each presentation of an input x , the best-matching unit (BMU_x) is selected: this is the unit i whose weight vector w_i is closest (in Euclidean Distance) to the stimulus x as in the following Equation 1:

$$i = \operatorname{argmin}_j \|x - w_j\| \quad (1)$$

The weights of the BMU and of its surrounding units are updated becoming closer to the stimulus x hence increasing the chances that the same unit (or the surrounding units) will be selected as the BMU for the same stimulus or for similar stimuli on subsequent presentations. As specified in Equation 2:

$$w_j(n+1) = w_j(n) + \eta(n) \cdot h_{BMU_x, j}(n) \cdot (x - w_j(n)) \quad (2)$$

where $\eta(n)$ is the learning rate, and $h_{BMU_x, j}$ is the neighbourhood function between the BMU for x and j . $h_{BMU_x, j}(n)$ is defined as in Equation 3:

$$h_{BMU_x, j}(n) = e^{-d^2(BMU_x, j)/2\sigma_n^2} \quad (3)$$

where $d(BMU_x, j)$ is the distance between BMU_x and unit j on the map's grid, and $\sigma(n)$ is the width of the gaussian. The neighborhood function plays a key role in the topological organisation of the map, by which similar inputs have close BMUs on the map. As standard in SOMs, the learning rate lowers with time, and the width of the gaussian shrinks with time, so that the first weight changes involve a large portion of the maps whereas the last weight changes concern fewer units. These properties of learning rate and neighborhood capture the assumption that major changes occur at the beginning of the formation of a category, whereas subsequent changes are minor refinements. In our simulations, the learning rate $\eta(n)$ starts at 0.1, and decreases to 0.001 in a manner inversely proportional to n ($\eta(n) = a/(n+b)$ with $b = 256, a = b \cdot 0.1$). The gaussian σ starts at 2 and shrinks with n to 1 ($\sigma(n+1) = 2 + (n-1) \cdot \varepsilon$ with $\varepsilon = -4^{-5}$). In our simulations we train the maps for 100 epochs (i.e., 100 overall presentations of the whole training set). With this training regime, semantic and auditory maps learn to topologically represent the semantic and acoustic stimuli, respectively.

Hebbian connections between the acoustic and the semantic map are learned after these two maps have separately learned to organise the semantic and acoustic stimuli, respectively. Hebbian connections are initially set to 0, i.e., there

are no connections between conceptual and auditory maps³. These Hebbian connections are then reinforced after each single joint presentation of an auditory and a semantic stimulus to the auditory and conceptual maps, respectively. The BMU for the semantic and the auditory stimulus is singled out in the conceptual and in the acoustic map, and the Hebbian connection between them is strengthened. As in the original Mayor and Plunkett (2010) model, this strengthening of connections involves not only BMUs but also surrounding, mostly activated units, as specified by Equation 4:

$$w_{u_a, u_s}(n+1) = w_{u_a, u_s}(n) + \omega \quad (4)$$

for $\omega = 0.1$, and u_a, u_s auditory and conceptual units such that:

1. either u_a is the BMU for auditory stimulus a (BMU_a) or its error for stimulus a is low enough ($\|a - w_{u_a}\| < 0.1 + \|a - w_{BMU_a}\|$).
2. either u_s is the BMU for semantic stimulus c (BMU_c) or its error for stimulus c is low enough ($\|c - w_{u_s}\| < 0.01 + \|c - w_{BMU_c}\|$).

In our current implementation once the maps have been trained, and Hebbian connections between these two maps have been learned, semantic representations are enriched with new features associated to the initial core representations. In brief, for each element x of the Semantic Training Set its best-matching BMU_x is singled-out, the dimension of w_{BMU_x} is increased: $w_{BMU_x} = w_{BMU_x}, [0, 0, \dots, 0]$, with as many 0 as there are extra features (in our simulations 8). Each extra j^{th} feature value is determined as $w_{BMU_x}(j)(t+1) = w_{BMU_x}(j)(t) + \mu \cdot (j - w_{BMU_x}(j)(t))$ for learning rate μ (in our simulations $\mu = 0.01$). For 100 epochs.

By this phase of semantic/conceptual enrichment we allow in our model new features to enrich semantic representations at all times. So for instance, the initial representation of “dog” may consist of just an initial prototypical representation of a dog, that might include its typical shape, the fact that it barks, that it has lots of fur, together with other simple properties. This initial representation may then be enriched with a variety of new features discovered little by little: that dogs usually run, that they usually like to play, that they can have certain kinds of roles in tales, etc. In our current implementation (although there is no formal constraint to do so) we take extra features to be more functional and related to expectations, such as those just exemplified. We also take them to adapt faster than basic features, as it will become clear in the next section.

Activation dynamics

The biggest difference between our model and Mayor and Plunkett (2010)’s model is in the activation dynamics. When

³An alternative initialisation strategy could be to initialise these connections randomly, and then prune those connections that are not reinforced through Hebbian Learning: the result, we think, would be equivalent.

a word is presented to the model, it is presented to its auditory sub-network. This generates the activation of a semantic representation by a cascade process: first, the acoustic best-matching unit i is activated. Then, the activation propagates to the conceptual sub-network through Hebbian connections. The maximally activated unit on the conceptual map is singled-out, and its associated semantic representation is activated. Over time, the activation of the semantic representation weakens through *adaptation*. Not all features need to adapt at the same speed. Here, we assume that extra features (that as said we take as being more related to expectations) adapt faster than basic features. For simplicity, we let $f_i(800ms.) = f_i/10$ for extra features, and $f_i(800ms.) = f_i$ for core features (that will adapt at a later moment, not considered by the model).

Testing

During test, Network Looking Time (*NLT*) is measured. *NLT* is a function of the *ease of transition* between subsequent semantic representations: networks attend longer to a sequence of words if they can easily update the corresponding semantic representations⁴. In turn, this ease of transition is a function of the Euclidean distance between consecutive semantic representations (a proxy of how many features need to be changed, and by what amount, to transiate from a semantic representation to the next). *NLT* is defined as follows:

$$NLT(j) = A + e^{-dist(r_j, r_{j-1})^2 / 2\sigma^2} + noise \text{ (ms.)} \quad (5)$$

where r_j and r_{j-1} are the semantic representations associated to word j and the preceding word $j - 1$; $NLT(j)$ is a Gaussian function of $dist(r_j, r_{j-1})$; σ is the spread of the Gaussian function, set to 1 in our simulations; *noise* is a random value from a uniform distribution in the range $[-1, +1]$; A is just a constant to bring looking time in the same range of values than infants' looking time in (Plunkett et al., 2022) experiments.

Network Looking Time for the whole sequence of words is the sum of all $NLT(j)$ for all words j of the list.

We have evaluated several simulations in order to study the behaviour of the model, and to determine whether it exhibits semantic priming in a manner similar to that observed in infants. Paralleling infant experiments as described by Delle Luche et al. (2014) and Plunkett et al. (2022), in each simulation we compare Network Looking Time when listening to an Unrelated List of Words (where subsequent words are taken from different taxonomic categories) with Network Looking Time when listening to a taxonomically Related List of Words (where words are taken from the same taxonomic category).

Furthermore, in order to parallel infant experiments, we have evaluated *NLT* in two conditions: SHORT ISI (ISI = 400 ms.) or LONG ISI (ISI = 800 ms.).

⁴This definition of Network Looking Time is similar to Plaut and Booth model's Reaction Time.

The difference between the two conditions appears when calculating *NLT* in Equation 5, and in particular when calculating the Euclidean distance between a semantic representation r_j and the previous one r_{j-1} ($dist(r_j, r_{j-1})$). In the SHORT ISI condition, for all words j (except the first) when its semantic representation r_j is activated, and the transition between a previous semantic representation r_{j-1} and r_j must be operated, r_{j-1} semantic representation is still fully activated and no feature has adapted yet. On the contrary, in the LONG ISI condition, when the new semantic representation must be activated extra features in r_{j-1} have started to adapt. We take therefore $f_i = f_i/10$ for all extra features.

Results

Simulation 1

In this simulation, we model Experiment 1 of Delle Luche et al. (2014) by comparing Network Looking Time at Taxonomically Related Lists (Taxonomically Related Condition) with Network Looking Time at Taxonomically Unrelated Lists (Unrelated Condition) in the SHORT ISI condition. Paralleling Experiment 1 of Delle Luche et al. (2014), for the Taxonomically Related Condition each model is tested with 6 randomly generated Taxonomically Related Lists. Likewise, for the Unrelated Condition, each model is tested with 6 randomly generated Unrelated Lists.

As in the infant experimental studies, we tested 24 models each with a different random start state. The random start state is intended to simulate individual variation.

Figure 3, left subplot, shows that at SHORT ISI, Network Looking Time is significantly higher for Taxonomically Related Lists (mean = 10.89, std = 0.32) than for the Unrelated Lists (mean = 10.64, std = 0.15). Significance is determined by a paired sample t-test, with $p < .001$.

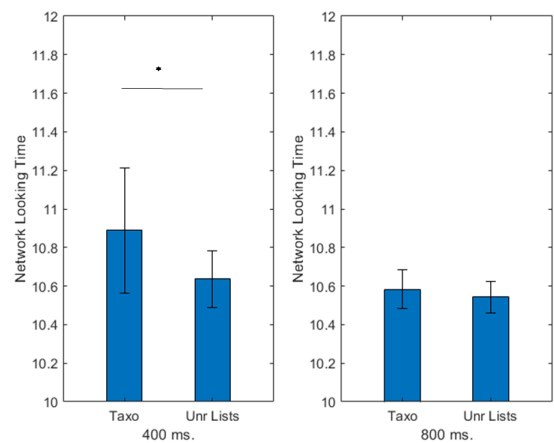


Figure 3: Results for Simulation 1 (left sub-plot), and for Simulation 2 (right sub-plot). At SHORT ISI (left plot) there is a significant difference between *NLT* at Taxonomically Related versus Unrelated Lists of Words. The difference is no longer significant at LONG ISI (right plot).

This matches infant behaviour, as assessed by Delle Luche et al. (2014, Experiment 1) (see Figure 1), where it was shown that 18-month old infants look longer at the Taxonomically Related versus Unrelated Lists of words. This points to a taxonomic priming effect.

Simulation 2

In this simulation, we address the same question addressed with infants by Experiment 1 in Plunkett et al. (2022), where it is shown that ISI has an impact on infant taxonomic priming, and taxonomic priming fades away at LONG ISI. The simulation follows the same steps as Simulation 1, but with LONG ISI (ISI = 800 ms.).

Figure 3, right subplot, shows that ISI has an impact on Network Looking Time, since the difference between the two conditions is no longer significant. Indeed, for Taxonomically Related Lists mean Network Looking Time is 10.58 (std. = 0.1), whereas for the Unrelated List the mean is 10.54 (std=0.08), and a paired sample t-test indicates that the difference is non-significant ($p = 0.13$). Recall that the only change from Simulation 1 to Simulation 2 is the longer ISI in the latter, i.e., 400 milliseconds vs. 800 milliseconds, respectively. The disappearance of the taxonomic priming effect at the LONG ISI is readily explained by the adaptation dynamics of the semantic representations in the network. By adaptation, the distance between taxonomically unrelated semantic representations becomes more comparable to the distance between taxonomically related semantic representations. Whence Network Looking Time becomes comparable for taxonomically related and for unrelated representations.

Discussion

We have proposed a model that aims to explain possible mechanisms underlying the impact of ISI on infant and adult semantic (taxonomic) priming. Contrary to the received view within the Distributed Representations Tradition (Plaut, 1995; Plaut & Booth, 2000), the model suggests that *adaptation* might be the mechanism responsible for the impact of ISI on semantic priming documented in 18 months old by Delle Luche et al. (2014) and Plunkett et al. (2022) (and in adults by (Alario et al., 2000)). Adaptation is the opposite of hysteresis postulated by Plaut and Booth (2000) to explain the impact of ISI on taxonomic priming. Interestingly, both adaptation and hysteresis change the geometry of similarity and distance among semantic representations. However, they do it in opposite ways. Hysteresis assumes that semantic representations at long ISI get richer of details. These extra details increase the distance between semantic representations that are taxonomically related, making it comparable to the distance between taxonomically unrelated representations. This comparable distance explains the fading away of taxonomic priming at long ISI. On the contrary, adaptation assumes that at long ISI semantic representations adapt, hence get poorer and lose details. The distance between taxonomically *unrelated* semantic representations becomes comparable to the distance between taxonomically related representa-

tions, and this explains the fading away of taxonomic priming at long ISI. Further research is needed to assess which mechanism, either hysteresis or adaptation, is responsible for the impact of ISI on semantic priming.

Plunkett et al. (2022) for infants (similarly to (Alario et al., 2000) for adults) shows that at LONG ISI priming is recovered only by the injection of associative links in the lists of words considered. Although associative links are outside the scope of the present papers, we acknowledge that they can be easily integrated in the model⁵.

In the current implementation of the model, the semantic and auditory stimuli are artificially created, and very simplified. Artificially created stimuli have proven very useful in studying category formation (see for instance Posner, Goldsmith, & Welton Jr, 1967 and Posner & Keele, 1968). These artificial stimuli have the advantage to facilitate the control of similarity relations within categories. Our stimuli have low inter-category similarity and high intra-category similarity.

In future work, we plan to consider more realistic stimuli. For what concerns semantic stimuli a candidate would be features as extracted by a convolutional neural networks (even if these are black boxes, and it is debated that features extracted by convolutional neural networks are similar to features humans use). As for acoustic stimuli, the ones we considered so far are place-holders for more realistic word representations. We leave for future work more realistic representations.

Conclusion

We have proposed a possible mechanistic account of infant semantic priming that replicates experimental results by Delle Luche et al. (2014) and Plunkett et al. (2022). The model extends Mayor and Plunkett (2010) model of early word learning, and provides a possible explanation of the role of ISI on semantic priming, as the result of adaptation. Further work is needed to understand if adaptation or hysteresis is responsible for the impact of ISI on semantic (taxonomic) priming.

References

- Alario, F. X., Segui, J., & Ferrand, L. (2000). Semantic and associative priming in picture naming. *The Quarterly Journal of Experimental Psychology: Section A* 53 (3), 741-764.
- Delle Luche, C., Durrant, S., Floccia, C., & Plunkett, K. (2014). Implicit meaning in 18-month-old toddlers. *Dev Sci*, 17(6), 948-955.
- Huber, D. E., & O'Reilly, R. C. (2003). Persistence and accommodation in short-term priming and other perceptual paradigms: temporal segregation through synaptic depression. *Cognitive Science*, 27, 403-430.
- Jusczyk, P. W., & Aslin, R. N. (1995). Infants detection of the sound patterns of words in fluent speech. *Cognitive psychology*, 29(1), 1-23.

⁵ Associative links are considered in a longer version of this paper under journal submission

- Kohonen, T. (2001). *Self-organizing maps (3rd ed.)*. Springer.
- Lerner, I., Bentin, S., & Shriki, O. (2012). Spreading activation in an attractor network with latching dynamics: Automatic semantic priming revisited. *Cognitive Science*, 36(8), 1339–1382.
- Li, P., Farkas, I., & MacWhinney, B. (2004). Early lexical development in a self-organizing neural network. *Neural Networks*, 17, 1345–1362.
- Li, P., Zhao, X., & MacWhinney, B. (2007). Dynamic self-organization and early lexical development in children. *Cognitive Science*, 31(4).
- Mayor, J., & Plunkett, K. (2010). A neurocomputational account of taxonomic responding and fast mapping in early word learning. *Psychological review*, 117 1, 1–31.
- Miikkulainen, R. (1997). Dyslexic and category-specific aphasic impairments in a self-organizing feature map model of the lexicon. *Brain and Language*, 59(2), 334–366.
- Miikkulainen, R., Bednar, J., Choe, Y., & Sirosh, J. (1997). *Computational maps in the visual cortex*. springer. springer.
- Plaut, D. C. (1995). Semantic and associative priming in a distributed attractor network. In *Proceedings of the 17th annual conference of the cognitive science society* (Vol. 17, pp. 37–42).
- Plaut, D. C., & Booth, J. R. (2000). Individual and developmental differences in semantic priming: Empirical and computational support for a single-mechanism account of lexical processing. *Psychological Review*, 107, 786–823.
- Plunkett, K., Delle Luche, C., Hills, T., & Floccia, C. (2022). Tracking the associative boost in infancy. *Infancy*, 1–18. doi: 10.1111/infa.12502
- Posner, M., Goldsmith, R., & Welton Jr, K. (1967). Perceived distance and the classification of distorted patterns. *Journal of Experimental Psychology*, 73(1), 28–38.
- Posner, M., & Keele, S. (1968). On the genesis of abstract ideas. *Journal of Experimental Psychology*, 77, 353–363.
- Treves, A. (2005). Frontal latching networks: a possible neural basis for infinite recursion. *Cognitive Neuropsychology*, 22, 276 - 291.
- Willits, J. A., Wojcik, E. H., Seidenberg, M. S., & Saffran, J. R. (2013). Toddlers activate lexical semantic knowledge in the absence of visual referents: Evidence from auditory priming. *Infancy*, 18(6), 1053–1075.