

The Considerations of Biological Plausibility in Deep Learning

By James Campbell

Department of Physics, College of Arts and Sciences at Cornell University

Abstract

In this literature review, we examine several deep learning algorithms in the context of biological plausibility and, in turn, argue that a backprop-like algorithm is the most likely candidate for how learning operates in the brain. Although there are numerous difficulties in how the backpropagation algorithm might be implemented in neural circuitry, we note that slight variations of the algorithm have been found to circumvent biological constraints and that seemingly unrelated algorithms can often be theoretically related to it. In particular, we examine the literature behind feedback alignment, target propagation, and equilibrium propagation, after giving some general background on learning in biology, AI, and their intersection. Ultimately, we acknowledge that there is no true consensus as to which learning algorithm the brain actually uses, but suspect that the answer is backprop-like in nature.

Introduction

With the recent successes of deep learning algorithms, it has become a topic of interest to see whether such algorithms (or variations of them) could potentially be the same as those used by biological systems. This has led to the question of biological plausibility, which measures the extent to which a learning algorithm could hypothetically be realized in the brain's hardware (Illing, 2019).

It is widely known that deep neural networks, consisting of layers of neurons and synaptic connections, are loosely based on the structure and function of the brain. Nevertheless, backpropagation, the central algorithm used to train deep neural networks, is considered by many to be incompatible with leading theories in neuroscience (Lillicrap et al., 2020).

In the following section, we provide some

background on learning in the brain and in artificial neural networks. We then examine the implementational constraints imposed by neural hardware and why the backpropagation algorithm violates them. In response to these constraints, several learning algorithms, such as feedback alignment, target propagation, and equilibrium propagation, have been devised, each of which attempts to overcome some of the difficulties encountered by backprop. The majority of this review consists of an analysis of these methods, including their successes and failures. Some of these successes are rather surprising and suggest that backprop-like algorithms are not as infeasible for the brain as previously thought. It is for this reason that we contend that the true function of the brain is likely similar in nature to backpropagation.

Background

Learning in Brains: Synaptic Plasticity

Today, it is widely believed that learning in

the brain happens at the synapses between neurons. This idea was first supported in experiments by Bliss and Lomo (1973) where they demonstrated that rapid stimulation of pre-synaptic neurons can lead to long-term changes in the activity of post-synaptic neurons. However, while they presented a mechanism (long-term potentiation) for learning via synapses, the exact details of synaptic plasticity remain largely elusive today. Part of the reason for this is the sheer complexity of biochemical processes present at any synapse; it is well known that synaptic behavior can depend on countless properties of neurons' vesicles, ion concentrations, receptors, proteins, and other structures.

Another prominent, though not necessarily contradictory, theory of learning is that of Hebbian learning. Propounded by Donald Hebb (1949), the general principle is that neurons that spike together most frequently will result in stronger synaptic connections. Since its origin, this notion has matured into what is now known as spike-timing-dependent plasticity (STDP), which is supported by a body of work better detailing the relationship between synaptic strength and the relative timing of pre- and post-synaptic neuron firings (Markram et al., 2011). Though it has come to inspire some neural network models (with limited success), STDP seems more likely to be an emergent phenomenon than a fundamental learning rule (Shouval et al., 2010).

Though synaptic plasticity represents the neuroscientific consensus with respect to learning in the brain, it is worth noting that there is some work challenging this idea or at least suggesting that it is incomplete. For instance, animals are able to perform one-shot learning over time periods much longer than those dictated by STDP (Gallistel and Balsam, 2014) and researchers were able to induce long-term learning through epigenetic means (Bedecarrats et al., 2018). For

the purposes of this review, however, the main takeaway from this section is that it is largely sufficient to think of the brain as a network of neurons that learn via synaptic changes and that this theory is well supported in the neuroscience literature from both old and recent experiments alike (Bliss and Lomo, 1973; Nabavi et al., 2014).

Learning in Artificial Neural Networks: The Backpropagation Algorithm

Whereas our knowledge of learning in the brain remains amorphous, there is one learning procedure that is ubiquitous in deep learning today: the backpropagation algorithm. Despite having traces in the early twentieth century, backpropagation first gained prominence in artificial intelligence when it was discussed in a paper by Rumelhart and Hinton (1986) as a means of training multi-layer perceptrons. The algorithm relies on the notion of automatic differentiation, in which the chain rule of calculus is used to differentiate along a computation graph (Millidge et al., 2020). Used in conjunction with iterative optimizers such as gradient descent, backprop gives a closed-form expression for the gradient of the loss function with respect to every weight in the network in terms of other weights and activations. In other words, gradient descent states that

$$\Delta W_{ij}^l = -\alpha \frac{\partial L}{\partial W_{ij}^l} \quad (1)$$

and backpropagation uses the chain rule to obtain a recursive definition for the derivative,

$$\frac{\partial L}{\partial W_{ij}^l} = a_j^{l-1} \delta_i^l \quad (2)$$

$$\delta^l = (W^{l+1})^T \delta^{l+1} \odot \sigma'(z) \quad (3)$$

Here there are 'k' training examples and 'l' layers, and 'z' is simply the matrix-vector multiplication of each layer's activations with its weight matrix. Moreover, the dotted circles represent element-wise multiplication and L is the loss function. Of particular importance for biological

plausibility is the delta term in the last layer of the network:

$$\delta^{last} = \nabla_a L \odot \sigma'(z^{last})_{(4)}$$

This is because it illustrates that the delta terms in fact correspond to error signals, where “error” is interpreted as the gradient of the loss function. In the case of a squared loss, for example, this quantity is the difference between the model’s prediction and the labeled target (Lillicrap et al., 2020). That is to say, backpropagation relays information about how different a model’s perception of the world is from reality.

Backpropagation’s Biological Incompatibility

While backpropagation has been extraordinarily successful in deep learning, there are several reasons why it is considered biologically implausible.

One major reason is that it is hard to come up with neural circuits that can implement the feedback computations specified by backpropagation. As can be seen in equation (4), backprop says that feedback error signals must be multiplied by the transpose of their layers’ weight matrix. Because feedback connections in the brain can either be implemented by sending error signals back along the original network or by allocating a separate feedback network, the requirement that error signals be multiplied by the transpose matrix gives rise to the weight symmetry and weight transport problems, respectively. Specifically, if error is propagated back along the feedforward network, then in order to implement backprop, the brain would have to enforce a synaptic weight matrix that is equal

to its transpose, which is to say symmetric (Bengio et al., 2016; Lillicrap et al., 2016; Lillicrap et al., 2020; Whittington & Bogacz, 2019). On the other hand, if the brain were to use separate feedback networks, then the weight transport problem addresses how the feedback synapses could possibly gain access to the strengths of the feedforward synapses (because unlike in a computer, the brain can’t just copy the weights to some memory address) (Grossberg, 1987).

Beyond the use of the transposed weight matrix, Bengio et al. (2016) note (as can be seen in equation (4)) that backpropagation also requires multiplication by the derivative of the activation function and that implementing this computation biologically is non-trivial. More generally, activation functions represent another source of biological constraint: neurons in the brain are known to fire discretely whereas the backpropagation algorithm, which contains derivative terms, benefits from continuous functions (Bengio et al., 2016; Whittington & Bogacz, 2019). Indeed, the very fact that backpropagation communicates error signals with floating-point precision casts major doubt on its biological plausibility (Neftci et al., 2017).

Another contention is that backpropagation learning contains two temporally alternating phases: one for the forward pass and another for the backward pass (Neftci et al., 2017). Moreover, whereas the forward pass performs nonlinear calculations, the equations of backpropagation are linear. This disparity would likely require distinct biological mechanisms that are unlikely to be found in the brain (Bengio et al., 2016). In addition, in deep learning, performing the backwards pass doesn’t influence the feedforward activations, though this does not seem to be the case in the brain (Lillicrap et al., 2020). Even more basic is the question of how the brain’s

neural networks could obtain targets or labels in order to compute the original error signal in the final layer.

Overall, many of the above restrictions stem from the simple principle of locality: that neurons in the brain can only interact with neighboring neurons (Whittington & Bogacz, 2019). As we will see, finding algorithms that can learn while maintaining the physical proximity of interacting components can go a long way towards obtaining biological plausibility.

Feedback Alignment

The method of feedback alignment offers to resolve many of the difficulties encountered in trying to implement backpropagation in the brain. The breakthrough paper for feedback alignment came in 2016 by Lillipap et al. (2016). The main idea is that one can avoid the weight symmetry and weight transport problems by using a completely separate and unrelated weight matrix for feedback. In other words, the authors attempt to perform backpropagation but replace the transposed feedforward matrix with a random one. For the l th layer, call the random matrix B_l . Then equation (4) is transformed to

$$\delta^l = B^l \delta^{l+1} \odot \sigma'(z)_{(5)}$$

In addition to feedback alignment, another method, termed direct feedback alignment was introduced by Arild Nøkland (2016). In this variation, the final error gets used as the one and only error signal for each layer. In other words, direct feedback alignment, in making yet another tweak to the equations for backpropagation, is able to remove the need to propagate error back through multiple layers of the network in the first place.

Direct feedback alignment is described by

$$\delta^l = B^l \nabla_a L \odot \sigma'(z^l)_{(6)}$$

It is of note that both direct and indirect feedback alignment seem downright counter-intuitive. Indeed, the authors themselves acknowledge that their results are surprising. Why, one may wonder, should using a random matrix result in a performance that rivals the precisely computed gradients of traditional backprop? In short, as the authors prove, even though a random matrix is not the same as the feedforward transposed matrix, it can still become aligned with it.

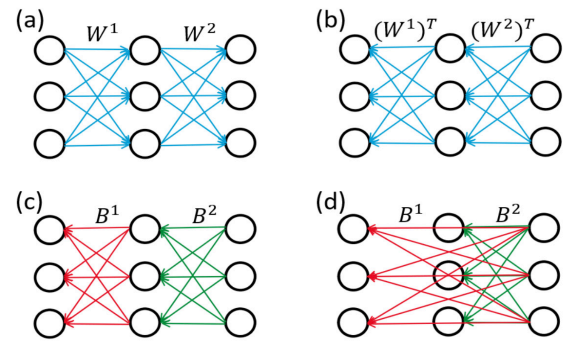


Figure 1: Diagram (a) depicts a feedforward neural network with the blue arrows representing the standard matrix-vector multiplication. Diagram (b) portrays the backpropagation algorithm where error signals are sent back through the network. The arrows are blue to represent the fact that backpropagation makes use of the transpose of the feedforward weights. Diagram (c) displays feedback alignment. As can be seen, the 'B' matrices have no relation to the feedforward weights, but error signals are still propagated back through the network. Diagram (d) shows direct feedback alignment. Here the error found in the output layer is combined with random matrices to compute the weight updates for all prior layers.

One may also wonder what the purpose of direct feedback alignment is considering the relative success of indirect feedback align-

ment. Arild Nøkland, the author of the direct feedback alignment paper, provides two reasons. The first is that while feedback alignment demonstrated high performance on a variety of complex tasks, its success plummeted as it was applied to deeper networks. Direct feedback alignment, by contrast, does much better, even when used in networks as large as 100 layers.

The second advantage of direct feedback alignment is that it does not require error signals to be propagated through many layers. Accordingly, it can be implemented via local interactions, which alleviates many biological constraints. In fact, direct feedback alignment has an incredible degree of flexibility; as Nøkland notes, under direct feedback alignment, a neuron can receive its error signal from a post-synaptic neuron, from a reciprocally connected neuron, from a pre-synaptic neuron, or from any location further upstream in the informational pathway.

In this sense, in spite of the fact that the specifics of the brain's error propagation mechanism remain unknown, a range of possible alternatives are all compatible with direct feedback alignment. Still, as Nøkland remarks, this is far from saying that direct feedback alignment is the exact learning algorithm employed by the brain; there is a lot that simply remains unknown, though direct feedback alignment is a step in the right direction.

Target Propagation

Pioneered by Yoshua Bengio (2014) and Yann LeCun (1986), target propagation is a method that uses auto-encoders to facilitate local feedback connections. Auto-encoders are unsupervised models that seek to predict their own input and will often have a hidden layer of reduced dimension. In target propagation, auto-encoders are stacked atop one

another in parallel to the main network. In short, if an auto-encoder can be trained to act like an inverse function, then the labeled targets can be propagated back through the auto-encoders, forming “hidden” targets to be compared to the hidden activations. By taking the difference between these “hidden” targets and their corresponding activations, one arrives at a local form of error that can be used to direct learning in the network (Lillicrap et al., 2020).

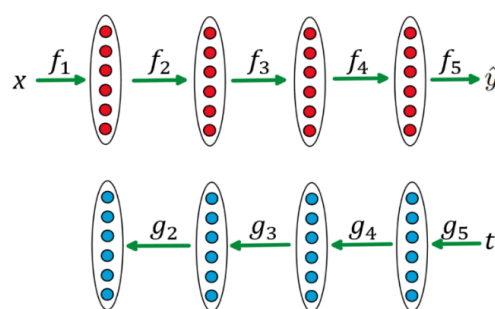


Figure 2: Target propagation utilizes auto-encoders to generate feedback that doesn't involve any gradient computations. It does this by propagating the labeled targets back through a series of operations that approximately invert the forward operations of the network. In this sense, a specialized target (in blue) is created for each hidden activation (in red), the effect of which is a feedback mechanism that is local in nature. For each layer in Figure 2, the error is given by the difference between the activations (in red) and the “hidden” targets (in blue). This error value can be used to nudge a given layer's synapses in a way that facilitates learning. Unlike a gradient, subtraction is a computation that can be more easily implemented biologically.

For some time, the main problem with target propagation was the inability to obtain auto-encoders that act as perfect inverses without resorting to backpropagation. This, however, was solved by difference target propagation (Lee et al., 2015), which is a simple linear correction to target propagation that is achieved by turning the network itself into a stack of autoencoders. For a

more in-depth explanation of difference target propagation, we direct our readers to Lillicrap et al. (2020).

As it turns out, difference target propagation performs decently on a variety of tasks and in many cases approaches the results of backpropagation (Lee et al., 2015). It therefore makes a formidable case for a biologically plausible learning algorithm. In terms of how it could be implemented by the brain, there are two main options; one would be to allocate a separate feedback network, in which the auto-encoders live. Again, this is a questionable hypothesis within neuroscience (Stork, 1989).

Another possibility is that all the various targets, activations, and reconstructions for a node dwell within a single neuron. This idea derives from recent research that neurons possess numerous spatial compartments that can perform far more nuanced calculations than the simplified model of the point-particle neuron would suggest (Körding & König, 2001; Urbanczik & Senn, 2014; Naud & Sprekeler, 2018). Ultimately, though, even if it cannot be implemented in neural circuitry, difference target propagation demonstrates how relatively mild architectural modifications can make backprop-like algorithms comply with many more biological constraints than originally thought.

Equilibrium Propagation

Another prominent biologically plausible algorithm is equilibrium propagation. It was introduced by Scellier and Bengio (2017) and uses an energy function to train a Hopfield model. More specifically, they define an “energy function” E , a “cost function” C , and a “total energy function” F as follows

$$E(u) := \frac{1}{2} \sum_i u_i^2 - \frac{1}{2} \sum_{i,j} W_{ij} p(u_i) p(u_j) - \sum_i b_i p(u_i) \quad (7)$$

$$C := \frac{1}{2} (y^* - t)^2 \quad (8)$$

$$F := E + \beta C \quad (9)$$

By varying the “clamping factor,” β , one is able to influence the effect that the cost function has on the system. In fact, equilibrium propagation occurs in two phases. In the “free phase,” $\beta=0$ and the system is allowed to establish equilibrium. In the second, “weakly clamped phase,” $\beta>0$ and the system is again allowed to settle to equilibrium. Because the potential energy will exert a force on the state of the system, the weights within the network will be nudged in a meaningful direction and moreover, in a local manner. In this way, learning is able to take place, even as equilibrium propagation offers an entirely new framework for thinking about the system.

The dynamical systems framework also results in behavior that deeply resembles that of spike-timing-dependent plasticity. At this point, one would think that equilibrium propagation is largely unrelated to backpropagation; interestingly, however, it can be shown that equilibrium propagation approximates the gradients computed in backpropagation-through-time. Unlike backprop, though, equilibrium propagation only requires one type of computation and one neural circuit.

But while it may check many of the boxes of biological plausibility, equilibrium propagation still requires two phases, which alternate temporally. As for performance, the original equilibrium propagation algorithm presented by Scellier and Bengio was not all that impressive, although recent work has added slight modifications that enable more competitive accuracy (Laborieux et al., 2021).

Conclusions

Having reviewed three major attempts at enforcing biologically plausibility

in deep learning, we now take a moment to reflect on some general trends.

For one, it is astonishing how relatively minor tweaks to existing algorithms have resulted in major progress toward biological plausibility: going from feedback alignment to direct feedback alignment enabled local connectivity and dramatically improved performance in deeper networks; going from target propagation to difference target propagation allowed for the natural approximation of inverse functions; modifying the architecture marginally in equilibrium propagation resulted in major performance boosts (Laborieux et al., 2021).

It is moreover curious just how rudimentary many of these tweaks were at heart. Feedback alignment simply fixed a variable in the equations for backprop; direct feedback alignment merely fixed a different quantity. The central insight of target propagation was to substitute gradient computations for auto-encoders; difference target propagation simply added more auto-encoders in different places. Equilibrium propagation, while a less trivial modification to backprop, was itself enhanced tremendously by a comparably trivial design alteration. All this is in spite of the fact that many of these adjustments addressed problems that were previously considered to be unyielding.

In this sense, we suspect that the learning algorithm actually employed by the brain is indeed rather similar to backpropagation. It is, after all, the starting point of many of the other algorithms. Additionally, even for algorithms that initially seem unrelated to backprop, such as equilibrium propagation, it is almost always the case that they can be related mathematically to it (Lillicrap et al., 2020). Furthermore, backprop is universally accepted as the algorithm with the best performance, and one would imagine that natural selection would favor such efficacy in humans.

It is worth acknowledging that, in spite of these suspicions, the brain's true learning algorithm remains utterly unknown (Stork, 1989). Nevertheless, after considering a sizable number of algorithms, we can say that the general idea behind backprop, which is that computed error signals inform the update of synaptic strengths, is almost certainly the basic mechanism that allows the brain to learn. One can only wonder, then, what other tweaks to the backpropagation algorithm might exist out there; perhaps at the core of human cognition lays something surprisingly simple.

References

- Bedecarrats, A., Chen, S., Pearce, K., Cai, D., & Glanzman, D.L. (2018). RNA from Trained *Aplysia* Can Induce an Epigenetic Engram for Long-Term Sensitization in Untrained *Aplysia*. *eNeuro*, 5(3). <https://doi.org/10.1523/ENEURO>
- Bengio, Y. (2014). How auto-encoders could provide credit assignment in deep networks via target propagation. ArXiv preprint. <https://arxiv.org/abs/1407.7906>
- Bengio, Y., Lee, D., Bornschein, J., Mesnard, T., & Lin, Z. (2016). Towards Biologically Plausible Deep Learning. ArXiv preprint. <https://arxiv.org/abs/1502.04156v3>
- Bliss, T.V. & Lomo, T. (1973). Long-lasting potentiation of synaptic transmission in the dentate area of the anaesthetized rabbit following stimulation of the perforant path. *The Journal of Physiology*, 232(2), 331-356. <https://doi.org/10.1113/jphysiol.1973.sp010273>
- Gallistel, C.R., & Balsam, P.D. (2014). Time to rethink the neural mechanisms of learning and memory. *Neurobiology of*

Learning and Memory, 108(1), 136-144.
<https://doi.org/10.1016/j.nlm.2013.11.019>

Grossberg, S. (1987). Competitive learning: from interactive activation to adaptive resonance. *Cognitive Science*, 11(1), 23-63. <https://doi.org/10.1111/j.1551-6708.1987.tb00862.x>

Hebb, D. (1949). *The Organization of Behavior: A Neuropsychological Theory*. Psychology Press.

Illing, B., Gerstner, W., & Brea, J. (2019). Biologically plausible deep learning—But how far can we go with shallow networks? *Neural Networks*, 118(1), 90-101. <https://doi.org/10.1016/j.neunet.2019.06.001>

Körding, K., & König, P. (2001). Supervised and unsupervised learning with two sites of synaptic integration. *Journal of Computational Neuroscience*, 11(1), 207-215. <https://doi.org/10.1023/A:1013776130161>

Laborieux, A., Ernoult, M., Scellier, B., Bengio, Y., Grollier, J., & Querlioz, Q. (2021). Scaling Equilibrium Propagation to Deep ConvNets by Drastically Reducing Its Gradient Estimator Bias. *Frontiers in Neuroscience*, 15(1), 129-130. <https://doi.org/10.3389/fnins.2021.633674>

LeCun, Y. Learning processes in an asymmetric threshold network. (1986). *Disordered Systems and Biological Organization*, 1(1), 233-240.

Lee, D.H., Zhang, S., Fischer, A., & Bengio, Y. (2015). Difference Target Propagation. *Joint European Conference on Machine Learning and Knowledge Discovery in Databases*, 498-515. http://dx.doi.org/10.1007/978-3-319-23528-8_31

Lillicrap, T.P., Cownden, D., Tweed, D.B., & Akerman, C.J. (2016). Random synaptic feedback weights support error back-propagation for deep learning. *Nature Communications*, 7(1), 13276. <https://doi.org/10.1038/ncomms13276>

Lillicrap, T.P., Santoro, A., Marris, L., Akerman, C.J., & Hinton, G. (2020). Backpropagation and the brain. *Nature Reviews Neuroscience*, 21(1), 335-346. <https://doi.org/10.1038/s41583-020-0277-3>

Markram, H., Gerstner, W., & Sjöström, P.J. (2011). A history of spike-timing-dependent plasticity. *Frontiers in Synaptic Neuroscience*, 3(1). <https://doi.org/10.3389/fn-syn.2011.00004>

Millidge, B., Tschantz, A., & Buckley, C.L. (2020). Predictive coding approximates backprop along arbitrary computation graphs. *ArXiv preprint*. <https://arxiv.org/abs/2006.04182v5>

Nabavi, S., Fox, R., Proulx, C.D., Lin, J.Y., Tsien, R.Y., & Malinow, R. (2014). Engineering a memory with LTD and LTP. *Nature*, 511(1), 348-352. <https://doi.org/10.1038/nature13294>

Naud, R. & Sprekeler, H. (2018). Sparse bursts optimize information transmission in a multiplexed neural code. *Proceedings of the National Academy of Sciences*, 115(27), 6329-6338. <https://doi.org/10.1073/pnas.1720995115>

Neftci, E.O., Augustine, C., Paul, S., & Detorakis, G. (2017). Event-Driven Random Back-Propagation: Enabling Neuromorphic Deep Learning Machines. *Frontiers in Neuroscience*, 11(1). <https://doi.org/10.3389/fnins.2017.00324>

Nøkland, A. (2016). Direct feedback alignment provides learning in deep neural networks. *Advances in Neural Information Processing Systems*, 29(1), 1037–1045. <https://proceedings.neurips.cc/paper/2016/file/d490d7b4576290fa60eb31b5fc917ad1-Paper.pdf>

Rumelhart, D., Hinton, G., & Williams, R. (1986). Learning representations by back-propagating errors. *Nature*, 323(1), 533–536. <https://doi.org/10.1038/323533a0>
Scellier, B., & Bengio, Y. (2017). Equilibrium propagation: Bridging the gap between energy-based models and backpropagation. *Frontiers in Computational Neuroscience*, 11(1), 24–37. <https://doi.org/10.3389/fncom.2017.00024>

Shouval, H.Z., Wang, S., & Wittenberg, G.M. (2010). Spike timing dependent plasticity: a consequence of more fundamental learning rules. *Frontiers in Computational Neuroscience*, 4(1). <https://doi.org/10.3389/fncom.2010.00019>

Stork, D. (1989). Is backpropagation biologically plausible? *International Joint Conference on Neural Networks*, 2(1), 241–246. <https://doi.org/10.1109/IJCNN.1989.118705>

Urbanczik, R. & Senn, W. (2014). Learning by the dendritic prediction of somatic spiking. *Neuron*, 81(3), 521–528. <https://doi.org/10.1016/j.neuron.2013.11.030>

Whittington, J., & Bogacz, R. (2019). Theories of Error Back-Propagation in the Brain. *Trends in Cognitive Science*, 23(3), 235–250. <https://doi.org/10.1016/j.tics.2018.12.005>