

Human-like Learning Framework for Frequency-Skewed Multi-level Classification

Amarjot Singh (as2436@stanford.edu)

James L. McClelland (jlmcc@stanford.edu)

Department of Psychology, 450 Jane Stanford Way
Stanford, CA 94305 USA

Abstract

Contemporary deep neural network based classification systems are typically designed to learn information at a single level of granularity from datasets in which all items occur with equal frequency. Humans, on the other hand, acquire information at several different levels of granularity from experiences that contain some items more frequently than others. This allows us to learn and differentiate frequent items better from other items. We investigate the consequence of learning from a natural frequency/multi-level dataset in a deep neural network designed to model the human neocortex, complemented in some simulations with a replay buffer, playing the role of the human hippocampus. The NC network, when trained on its own, is able to learn more frequent items relatively quickly and differentiate them better from other items, as human learners do. However, the network's performance on infrequent and unseen examples pays a price in generalization performance compared to a standard training regime. The replay buffer serves to ameliorate these deficiencies, and we introduce a computationally and psychologically motivated replay weighting scheme that performs better than two alternatives.

Keywords: Multi-level classification; frequency effects; complementary learning systems

Introduction

Humans continuously acquire information from profoundly skewed distributions. We have far more experience with some items than others, and our performance is frequency sensitive – frequency affects accuracy and reaction times in word and object identification and many other tasks (Patterson et al., 2006). Not only do we recognize familiar things with higher probability, but we also differentiate them from other items better (Shiffrin, Ratcliff, & Clark, 1990). We also learn both general and specific information – for example, we learn what things are cars, which cars are Subarus, and which Subaru is our own specific Subaru.

In contrast to humans, in most machine learning systems, all items are presented equally frequently, and each item is generally assigned only to a single class or category at a single level of generality. Interest focuses only on the question of how well we have learned to classify items in a held-out test set, ignoring how well we know the things we have seen. A further issue is that most of this research ignores the question of how rapidly we learn. Machine learning research focuses on presenting all the training data repeatedly until we reach an optimum on held out items. For humans, there is no train-test split – each presentation of an item is potentially a test of our knowledge of it, as well as an opportunity to learn. And

it seems natural to think that we want to learn new things as quickly as possible – we want to know what we have been told (like people's names or word meanings) after as few presentations as possible.

Our approach to these issues starts with the idea that humans continuously acquire information using two complementary learning systems (Marr, 1971; McClelland, McNaughton, & O'Reilly, 1995). A fast learning system relying on the hippocampus in the medial temporal lobes (MTL) acquires new information quickly, while a slower learning system based in other neocortical (NC) brain areas gradually builds up knowledge in a form that no longer depends on the MTL. Integration into the NC is thought to depend in part on replay of information stored in the MTL. A great deal of evidence now supports the view that memory replay occurs during sleep in humans and in other animals (Wilson & McNaughton, 1994). David Marr, an early proponent of replay, and most empirical studies have emphasized replay events occurring during the night immediately after exposure to an item. In humans, however, evidence from the effects of brain lesions supports the view that memory depends on the fast learning system for a period that extends over many years or even decades (Mackinnon & Squire, 1989). Thus, we treat the fast-learning system as a limited capacity system that learns rapidly but forgets over an extended time scale. In this way, replay is more probable for recent items, but older items still have some opportunity for replay. We build from this starting place to propose a human-like learning framework that uses a complementary learning architecture to learn information at multiple levels of granularity about items whose frequency varies over a wide frequency range. Within this framework, we explore several issues:

- What consequences do learning with unequal frequency distributions and learning to classify at several levels of granularity have for learning outcomes and learning dynamics, both for items presented during training, and items held out of the training set for testing?
- Does learning with an unequal frequency distribution lead to greater differentiation of frequent items compared with less frequent items?
- What is the best policy for replaying items for facilitating learning in the neocortex-like deep neural network?

We have taken four steps aimed at the development and assessment of this framework, and to aid in answering these questions.

First, we have focused on creating a data set with a frequency distribution like that encountered in natural human experience and including items to be categorised at three levels of granularity. We call this a *natural frequency / multi-level* (NFML) Dataset. This data set is composed of items that appear with a Zipfian frequency distribution and encode its class labels at a coarse or superordinate level, a finer or more basic category level, and, for some frequent items, at an item specific level. For this work, we assume that all classes appear with the same frequency and only the items within a class appear with different frequencies. We also assume that all the data is available together avoiding the continuous acquisition of information. We hope to extend this study to include the two factors.

Second, we have aimed at designing a complementary learning architecture inspired by human learning mechanisms, composed of a fast learning network that idealises the human medial temporal lobe memory and a slow learning network that represents the neocortical brain areas. The slow learning NC system is represented by a standard deep network. This network is the primary focus of our study as we view our integrated neocortical learning as the form of learning that informs our most automatic and deeply entrenched expectations, reactions, and intuitions. Although we believe information stored in the MTL guides our behavior while still available in that system, at this stage of our work the role of the fast learning system is only to provide extra exposures to the NC network through replay. Both networks receive experience in the form of images sampled from the proposed Zipfian frequency dataset each ‘day’. The items in the fast learning system are replayed to the slow learning using different replay schemes during the following ‘night’.

Third, we have examined the learning in the NC system on its own, and we have developed alternative replay schemes to explore how they affect integration of knowledge into the NC system. We consider human replay capacity to be a finite resource, and model this by restricting the number of replay events that the rehearsal buffer can provide during each ‘night’ or replay cycle. We propose a replay weighting scheme that prioritizes items for replay weighted by the network’s error on the item and the item’s recency of occurrence. Weighting by the network error enhances the efficiency of learning by focusing learning where it can do the most good (Schaul, Quan, Antonoglou, & Silver, 2015). Weighting by recency advantages frequent items (frequent items will have been seen more recently, on average, than infrequent items), helping to minimize the overall loss because these items contribute to the loss more often. Recency weighting also reflects human forgetting – the tendency of an item to become less available in memory with the passage of time. As we shall see, this weighting scheme enhances overall learning compared to two baseline replay policies.

Fourth, we propose several measures of the NC system’s performance on the NFML dataset. The results are compared to the results to the standard uniform frequency distribution used in most machine learning classification task settings. Our work also allows us to explore whether deep neural network models trained with the proposed NFML dataset show greater differentiation of frequent items from other items, as human learners do. We also assess how well the different replay schemes support classifying items of differing frequencies at the item-specific level.

The approach that we have taken involves considerable simplification relative to the brain and the details of natural experience. Yet it allows, we hope, the prospect of beginning to understand more about how learning occurs with the highly skewed data distributions and multiple levels of classification that humans experience, in a setting where we seek to maximize overall learning success summed over the entire course of the learning process. In future work we plan to assess the combined performance of the fast and slow-learning systems at test time, as well as considering the performance of the slow-learning system on its own. We also aim to extend the current Zipfian datasets so that the data is available incrementally as well as the classes are presented with different frequencies in addition to the items.

Human-like Learning Framework

Our human-like learning framework is based on a computational systems that learns multi-level information from the proposed NFML frequency dataset as shown in Fig. 1. We first describe the development of the proposed learning regime has been on creating datasets that are in accordance with human experiences. Next, we present the complementary learning architecture, composed of a slow and fast learning system, that can learn this multi-level information from the Zipfian datasets. Finally, we describe the measures we have adopted to evaluate the learning performance of the slow learner on the proposed datasets.

Natural frequency/multi-level (NFML) Datasets

Traditional machine learning systems are designed to gradually acquire information at a single level of granularity from stationary batches of well-balanced training data with many repeated exposures. These assumptions are improbable for humans as detailed in the introduction section. We have adapted the standard CIFAR-100 (Krizhevsky & Hinton, 2009) dataset to include the frequency effect and encode label information at multiple levels.

This dataset is constructed of 100 classes which are composed of 20 classes (coarse-level) with each class further divided into 5 sub-classes (fine-level). Each fine class consists of 500 training ($500 \times 100 = 50,000$ total) and 100 test images ($100 \times 100 = 10,000$ total) with each image (item-level) of size $32 \times 32 \times 3$. The training and test sets are merged together to create the complete dataset of 60,000 images with each fine class containing 600 images.

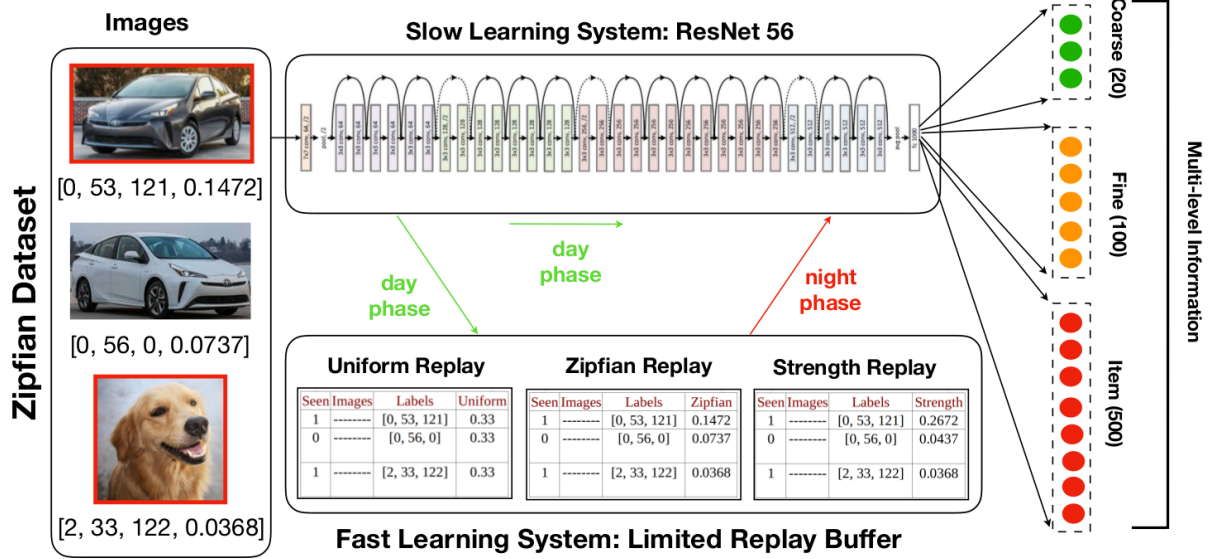


Figure 1: Our proposed human like learning framework, composed of the complementary learning architecture trained on a natural frequency/multi-level (NFML) data set. The architecture includes a slow learning system based on a ResNet 56 network and a fixed capacity replay buffer that idealizes the human fast learning system. The architecture consumes the NFML dataset that represents the skewed data distribution experienced by humans to learn to classify images at three levels: coarse, fine, and item. Three illustrative items with their labels are shown.

In our simulations, we have use three instances of a frequency-weighted data set. Each of the three instances (set-1, set-2, and set-3) is constructed by first randomly sampling 500 items for training and 100 items for test from each fine-level class. Next, we assign a random frequency rank to all 500 training set item within each class. This rank is used to produce a Zipfian probability distribution detailed below:

$$p(\text{item}_i) \propto 1/\text{rank}(\text{item}_i) \quad (1)$$

Probabilities are normalized to sum to one across all 50,000 items in the training set. To simplify the experiments, all classes appear with equal frequency. Once the probabilities are assigned to each item, batches of items to be used in training are sampled with replacement for presentation with the assigned probability. As a comparison to the standard CIFAR-100 dataset, if one million items are sampled from the Zipfian dataset, the most frequent items would be sampled 1472 times and the least frequent item just 3 times. This number is 20 $((1/50000) \times 1000000)$ presentations per item for the standard CIFAR-100 dataset. In addition to the coarse and fine labels, 5 of the most frequent items from each class are assigned an item-specific label. These are the first, second, fourth, eighth, and sixteen ranked items in each class, and they appear with the following frequencies per million presentations: 1472, 736, 368, 184, 92. Training of the slow-learning system (described below) occurs in alternating 'day' and 'night' phases. Each 'day' is thought of as arising from direct experience, consisting of 500 batches of 20 images each ($500 \times 20 = 10000$ images) sampled from the complete data set with replacement. Presentations during the night phase depend on the replay condition, described below.

Complementary Learning Architecture

The architecture consists of a commonly used deep neural network and a rehearsal buffer that idealizes the properties of a fixed-capacity fast-learning system.

ResNet: We use a ResNet-56 (He, Zhang, Ren, & Sun, 2015), a 56 layer convolutional architecture as the slow-learning deep neural network. The ResNet architecture employs 'Skip Connections' – connections that skip over sets of stacked hidden layers, making it possible to train very deep networks. Our only modifications to this architecture are its output layer, its activation function and the loss function. In our network, the next-to-last layer of the ResNet is fully connected to our modified output layer, which consists of 620 units, one for each of the 20 course class labels, one for each of the 100 fine class labels, and one for each of the 500 item-specific labels. Each output unit's activation is independently computed from the weighted input it receives from the next-to-last layer using the logistic function. For a given item, the target activation is 1 for the two units corresponding to its coarse and fine labels and for the unit corresponding to its item-specific label if it has one. The target activation is 0 for all of the other units. Note that this means that the target for all of the item-specific label units is 0 if the item does not have an item-specific label.

The network is trained using the binary (unit-wise) cross entropy (BCE) loss, defined for each item i as:

$$L_i = - \sum_u (y_u \log(\hat{y}_u) + (1 - y_u) \log(1 - \log(\hat{y}_u))) \quad (2)$$

Where u indexes the output units, y_u is the ground truth value for unit u and $\hat{y}_u$ is activation of output unit u after application

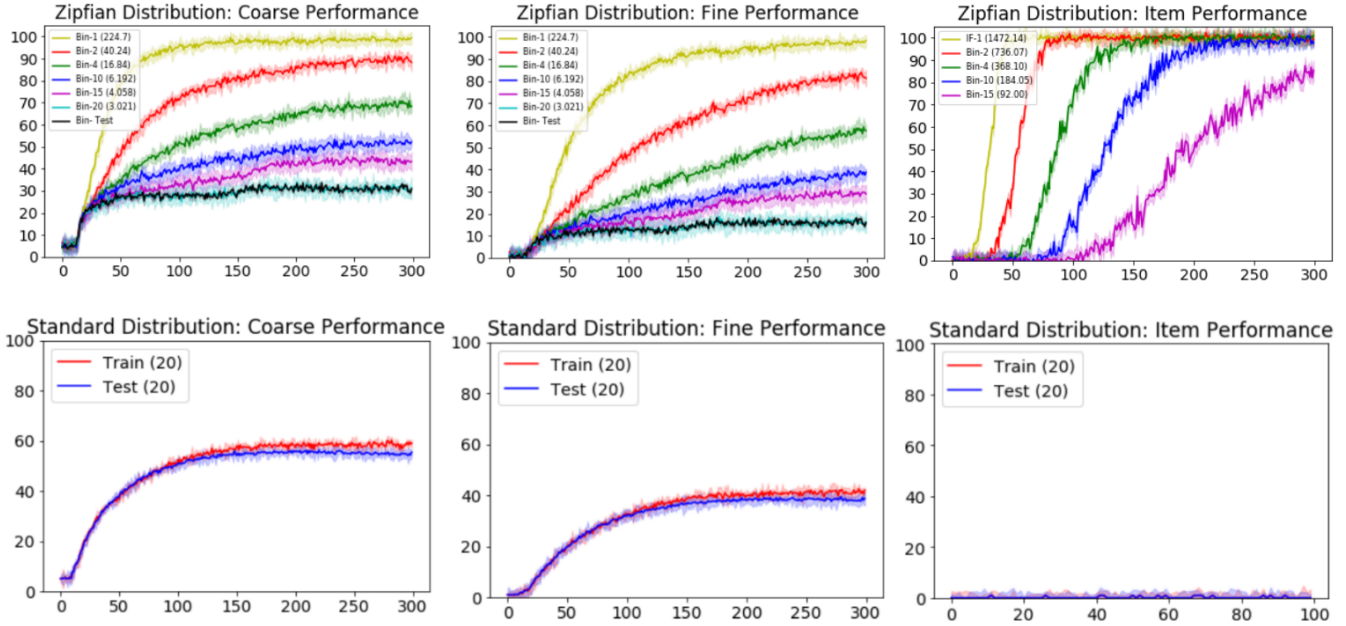


Figure 2: Frequency Effect: The illustration shows the coarse (20 classes), fine (100 classes), and item (500 classes) level performance comparison of the ResNet learner trained on the Zipfian and standard CIFAR-100 datasets. This analysis is performed on several bins which are composed of items with decreasing frequencies as detailed in the graphs. As expected the learner trained on the Zipfian dataset can rapidly learn the bins with frequent items, similar to human learners. Also, the frequent bins are learned better as their final asymptotic performance is higher. This frequent effect is not observed for the standard training regime as all the items are presented with the same frequency (20 per million presentations) and hence the learner has a similar performance on all items. The graphs show the average performance across the three (set-1 to set-3) experiments. The standard deviation is also shown at each point with the shaded region around the mean.

of the logistic function. The weights of the ResNet architecture are initialized by randomly sampling values from a uniform distribution with a range that depends on the number of inputs and outputs of each weight matrix (Glorot & Bengio, 2010). The loss is then averaged over the batch and passed to a stochastic gradient descent optimizer using a constant learning rate of 0.001 and no momentum.

Replay Buffer: The replay buffer is implemented as a list of entries, one for each item in the training set. Items that have occurred at least once during any training “day” are in the pool of items available for replay. Our experiments include a no-replay condition, as well as three replay conditions in which items are sampled from the replay buffer to present 5,000 training items to the slow-learning network during each simulated “night” following each simulated day. We propose what we call a *strength-based* replay policy which we compare with two simpler policies. According to the *Strength-based Replay Policy*, each item’s replay probability is proportional to its strength, which is the product of a regret factor and a recency factor, so that $strength = loss \times recency$. The regret factor is the summed loss at the output layer of the network the last time the item was presented. Quite simply, we imagine that the learner experiences regret in proportion to how poorly it produced the correct output, and therefore devotes more storage capacity (or more immediate rehearsal, increasing strength of the stored memory trace) to

the items that produce the most regret, promoting their replay and down-weighting the probability of replaying items that are already well known. Recency is a hyperbolic function of time t since last prior presentation as is typical of human forgetting curves. Specifically, $recency = (t)^{-0.16666}$ where one unit of time corresponds to 2,000 pattern presentations. One comparison policy is a *Zipfian* policy, such that items seen at least once in any previous day phase are chosen for replay in proportion to the item’s Zipfian probability. This scheme can be seen as representing a control for the extra presentations occurring during the night, using the same frequency-weighting scheme enforced by direct experience during the day. The other comparison policy is a *uniform* policy, such that items that have occurred at least once in any previous day are sampled from the buffer for replay according to a uniform distribution, independent of either recency or frequency.

Performance Measures

Zipfian Dataset Performance: The performance of the ResNet without replay is evaluated for the coarse and fine level classification performance, with the expectation that more frequent items should be learned better and more quickly than the infrequent items. To measure the frequency effect, the 500 training items in each fine class were sorted according to their sampling probability in the Zipfian dataset and then grouped into 20 bins of 25 items each. The frequency bins were combined across the 100 fine classes result-

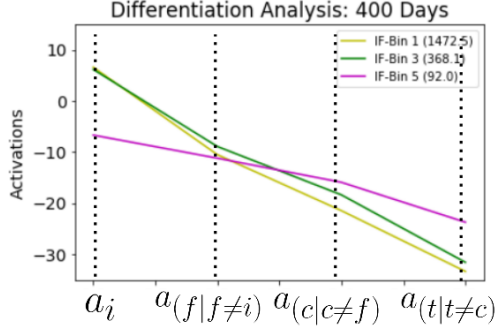


Figure 3: Differentiation of different target items with item specific labels from other items, as a function of the frequency of the target item. As explained in the text, greater differentiation corresponds high activation when the item itself is presented (left most point on each curve), and steeper fall off as other items less and less similar to it are presented. See *Differentiation* section in text for details.

ing in 2500 items per bin. Bin-1 contains the 25 most frequent items while Bin-20 contains the 25 least frequent items. Bin-1 is expected to be sampled at an average of 224 times per million presentations while this number is 3 for Bin-20. Bin-21 represents 10,000 test items from the held-out test set.

Differentiation: Humans differentiate more frequent items from less frequent items (Shiffrin et al., 1990). We looked for differentiation in the network’s performance on items with item specific labels. The learner’s differentiation ability is measured by evaluating the input activation a_i to the item-specific unit for the item, prior to the application of the logistic function. Following Criss (2006), differentiation is demonstrated by how sharply this value falls off as items that are less and less similar to the target are presented. In figure 3, we first plot the average activation a_i for the cases where, for each item in the bin, the target item itself is presented ($a_i|i$). Next we plot average values of a_i when any other item in the same fine class as the target is presented ($a(f|f \neq i)$), then the average activation ($a(c|c \neq f)$) for all the remaining items in the same coarse class as the target, and finally the average activation ($a(t|t \neq c)$) of unit for the item for all the remaining items in the Zipfian CIFAR-100 dataset.

Knowledge Integration: Different replay policies are explored for effective integration of the knowledge stored in the replay buffer at all three (coarse, fine, and item) levels. The best policy should produce faster acquisition of information with respect to no replay and is measured as:

$$M_g = \frac{M_{policy}}{M_{no-replay}} \quad (3)$$

Here M_g is either slope gained (S_g) or area gained (A_g) by the learning curve for a specific policy as compared to no replay. The slope for a learning curve is estimated by fitting a sigmoid curve. Area (A) corresponds to the area under the learning curve until the end of the training. This analysis is performed for the most, intermediate, and least frequent bins (Top, Middle, Bottom) for all three policies.

Table 1: The area and slope gained computed for the frequency weighted learning curves, as detailed in the performance measures section, is presented for the replay policies over three datasets. This is shown for all coarse, fine, and item levels.

	Uniform		Strength		Zipfian	
	S_g	A_g	S_g	A_g	S_g	A_g
Coarse-level						
Set-1	1.67	1.37	1.87	1.87	1.88	1.79
Set-2	1.71	1.41	1.99	1.91	1.93	1.75
Set-3	1.73	1.53	1.80	1.84	1.99	1.83
Fine-level						
Set-1	1.63	1.13	1.93	1.33	1.93	1.27
Set-2	1.55	1.19	1.99	1.44	1.95	1.31
Set-3	1.67	1.23	1.87	1.47	1.99	1.19
Item-level						
Set-1	1.79	1.32	1.97	2.04	1.81	1.59
Set-2	1.73	1.39	1.99	2.11	1.88	1.71
Set-3	1.81	1.44	2.01	1.97	1.79	1.67

In order to measure the overall effectiveness of each policy, the slope and area measures are also computed over the complete dataset using the frequency weighted curves obtained as shown below:

$$ovlc_{fw} = \frac{lc_1 \times freq_1 + lc_2 \times freq_2 + \dots + lc_N \times freq_N}{freq_1 + freq_2 + \dots + freq_N} \quad (4)$$

where $ovlc_{fw}$ represents the overall frequency weighted learning curve, lc_i represents the learning curve for a specific item and $freq_i$ represents the item’s frequency.

Experimental Results

Zipfian Classification Performance: The performance of the ResNet learner without replay was evaluated after each of 400 simulated ‘days’ on the three NFML datasets (set-1, set-2, set-3). The coarse and fine level classification performance for several bins with different item frequencies per million are shown in Fig. 2. It can be seen that higher frequency is associated with both faster learning and higher accuracy at the end of the training period, though accuracy is still increasing at the end of the simulation. Performance on held-out test items falls below that on even the lowest frequency training items, which have been seen only about 12 times each at the end of training, indicating that the network is sensitive to item-specific effects even for very infrequently trained items.

Comparing the results with identical ResNets trained on the same three data sets, but with the standard procedure of using equal frequencies for all of the trained items, we see that performance on held-out test items is far better with equal frequencies than with the Zipfian frequency distribution, both at the coarse and the fine levels. On the other hand, with equal frequencies, performance at the individual item level remains

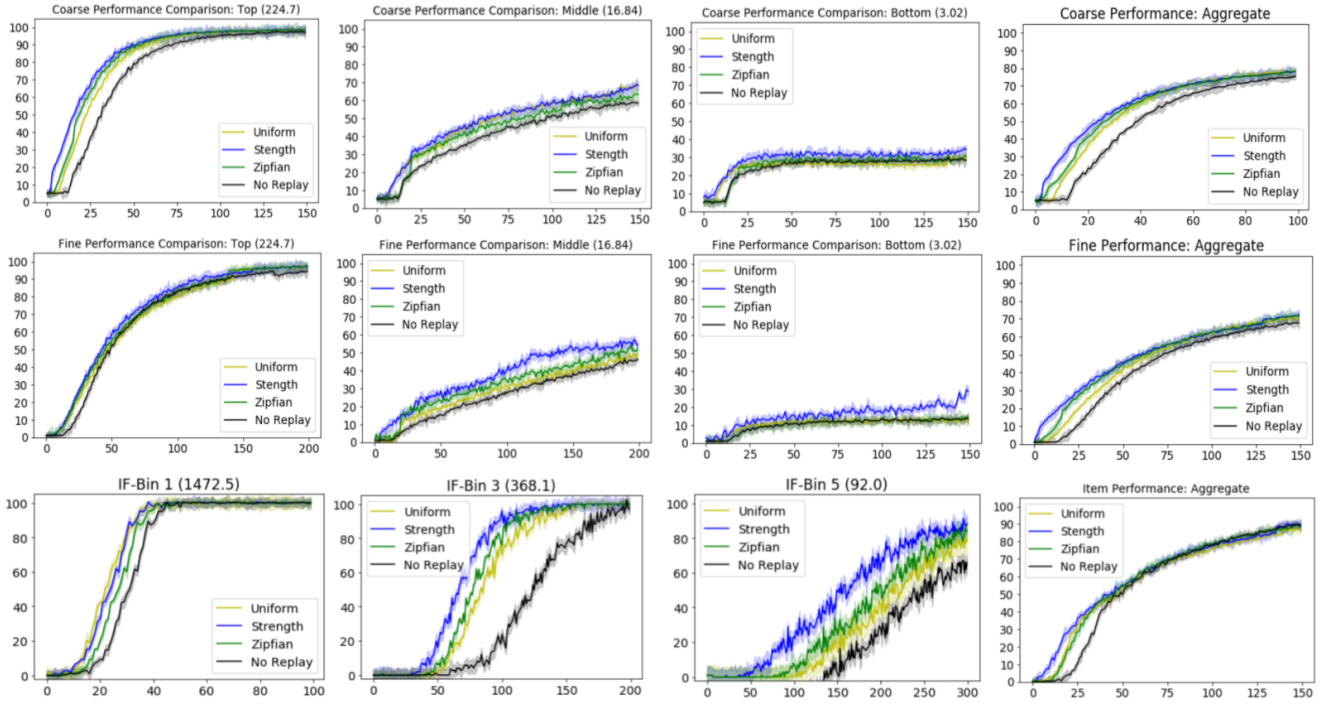


Figure 4: Replay Scheme Comparison: The illustration presents the learning curves for the most, intermediate, and least frequent bins (Top, Middle, Bottom) for the strength, Zipfian, and uniform replay schemes. The learning curves for no replay training regime are also shown. It can be seen from the curves that the proposed strength replay scheme outperforms the other two replay schemes as well as no replay as it results in faster consolidation for the items. The graphs show the average performance across the three (set-1 to set-3) experiments. The learning curves for frequency weighted aggregate performance computed for all the bins are also shown for all three replay policies as well as the no replay case. The standard deviation is also shown at each point with the shaded region around the mean.

at floor throughout training. Thus, the Zipfian distribution appears to favor specific over general information, while equal frequencies favors the general at the extreme expense of the specific.

Even with the standard, uniform frequency training regime, the ResNets trained to classify at the coarse, fine, and item level do not perform well. This reflects competition among the three classification tasks. When the ResNets are trained only to classify at the fine level, training accuracy reaches 93% correct, and test accuracy reaches 88% correct; when trained on both the coarse and the fine but not the item level, the train and test scores were 76% and 74% respectively.

Differentiation: Training with a Zipfian frequency distribution enables greater differentiation of frequent items from the other items. As seen in Fig. 3, the activation produced by the item itself (a_i) at its specific node is higher than the activation of that node produced by other items. The ResNet also learns to reduce the activation at the item-specific node for other items resulting in a cross-over effect, so that the three curves all cross each other, as expected if more frequent items are more strongly differentiated from other items.

Consolidation with different replay policies: our strength-based replay policy aids knowledge integration as compared to the uniform or Zipfian replay policies. As shown in Fig. 4, the strength-based policy has a larger advantage over the other two policies relative to no replay. This is seen

for all three coarse, fine, and item levels with particular advantage seen for the items with frequencies of 368 and 92 per million in Fig. 4. Replay itself is relatively unimportant for the highest frequency items which are learned very rapidly. The overall effectiveness of the replay policies is also computed using the frequency weighted analysis detailed above. The strength policy has an overall advantage as well over the remaining policies. This is detailed quantitatively as well using the slope (S_g) and area (A_g) gain measures in Table 1. It can be seen that the highest gains are for the strength policy on all three datasets.

Conclusion

This work proposed a human-like learning framework that used a deep neural network in combination with a replay buffer inspired by the neocortex and hippocampus respectively. Several aspects of our findings that are of note.

We found that both the multi-level classification task and frequency weighting greatly impacted the network’s learning. Considering first multi-level classification, when we compared the performance of the ResNet trained on the standard CIFAR 100-way classification task using uniform frequencies, it reached 88% correct performance on trained and held out test items, but when it was trained with the same uniform frequencies to classify items at the coarse, fine, and item levels, performance on fine level classification dropped to about

40% correct, while completely failing at item-level classification. This appears to be a weakness of these systems – it would be desirable for a learning system to be able to classify at multiple levels of granularity. Future work should explore neural network architectures that can accommodate classification at multiple levels more easily.

Given the challenge posed by multi-level classification, a Zipfian frequency distribution allowed the network to focus its efforts on items occurring more frequently, leading to higher aggregate accuracy on all three classification tasks than uniform frequency weighting (compare the aggregate ‘no replay’ curves in the right panels of Fig. 4 with the train set performance at the three levels shown in the bottom row of Fig. 2). This did come at the expense of performance on held out test items. Further research should focus on understanding how we as humans achieve frequency sensitivity, while also achieving good generalization.

Our strength-based replay policy builds on prioritized replay (Schaul et al., 2015) by introducing recency weighting and enhanced consolidation compared to the other replay policies considered. This supports the view that there are advantages to forgetting in addition to the loss produced by each item, since it allows focusing on recent items, which are more likely to be occur again (Anderson & Schooler, 1991).

We see this work as an initial foray into using slow-learning deep neural networks complemented by hippocampus-like fast learning systems to model human memory, where sensitivity to frequency and familiarity-driven differentiation are prominent aspects. An important next step is to use the hippocampus like system to work together with the cortex more fully, providing the initial basis for correct performance on new items and for retention of item-specific information, as it does in human learners (Knowlton & Squire, 1993). We are intrigued by the possibility that this might allow the deep network to specialize more in acquiring generalizable knowledge, thereby improving fine-level classification on held-out test items, which was severely impacted by the use of the NFML training set. The use of the hippocampus-like system to guide responding during direct experience might also enhance the impact of our strength-based hippocampal weighting scheme. We are particularly interested in exploring the impact of using such a system in the context of learning when the experience distribution is not completely stationary. Real natural experience involves relatively constant exposure over a life time to some items, but other items come and go, and a hippocampus-like system is likely to prove especially useful for performance on items that have recently occurred once or a few times, but that will never be experienced again. Here again this function may help buffer a cortex-like deep neural network from responsibility for specific items.

An even longer-term goal for a human-like learning system will be to replace the current replay buffer with truly MTL like learning system, which is unlikely to retain exact copies of training images and their targets and/or internal net-

work states, as in the rehearsal buffers used here and in many deep learning architectures. We look forward to emergence of large-scale systems that use the kinds of sparse distributed representations first envisioned by Marr (1971) that can encompass the kinds of data sets required to capture more fully the role of the human hippocampus in supporting our memory for the large number of items we know about and the wide range of different kinds of things we know about them.

Acknowledgment. This research was supported by the DARPA L2M program. The views expressed are the authors.

References

- Anderson, J. R., & Schooler, L. J. (1991). Reflections of the environment in memory. *Psychological science*, 2(6), 396–408.
- Criss, A. H. (2006). The consequences of differentiation in episodic memory: Similarity and the strength based mirror effect. *Journal of Memory and Language*, 55(4), 461–478.
- Glorot, X., & Bengio, Y. (2010). Understanding the difficulty of training deep feed forward neural networks. *Proc. 13th International Conf. on AI and Statistics*.
- He, K., Zhang, X., Ren, S., & Sun, J. (2015). Deep residual learning for image recognition. *arXiv:1512.03385*.
- Knowlton, B. J., & Squire, L. R. (1993). The learning of categories: Parallel brain systems for item memory and category knowledge. *Science*, 262(5140), 1747–1749.
- Krizhevsky, A., & Hinton, G. (2009). Learning multiple layers of features from tiny images. <https://www.cs.toronto.edu/~kriz/cifar.html>.
- Mackinnon, D. F., & Squire, L. R. (1989). Autobiographical memory and amnesia. *Psychobiology*, 17(3), 247–256.
- Marr, D. (1971). Simple memory: a theory for archicortex. *Philosophical transactions of the Royal Society of London. Series B, Biological sciences*, 262(841), 23–81.
- McClelland, J. L., McNaughton, B. L., & O’Reilly, R. C. (1995). Why there are complementary learning systems in the hippocampus and neocortex. *Psychological Review*, 102(3), 419–457.
- O’Reilly, R. C., Bhattacharyya, R., Howard, M. D., & Ketz, N. (2014). Complementary learning systems. *Cognitive science*, 38(6), 1229–1248.
- Patterson, K., Ralph, M. A. L., Jefferies, E., Woollams, A., Jones, R., Hodges, J. R., & Rogers, T. T. (2006). “pre-semantic” cognition in semantic dementia: Six deficits in search of an explanation. *Journal of Cognitive Neuroscience*, 18(2), 169–183.
- Schaul, T., Quan, J., Antonoglou, I., & Silver, D. (2015). Prioritized experience replay. *arXiv preprint arXiv:1511.05952*.
- Shiffrin, R. M., Ratcliff, R., & Clark, S. E. (1990). List-strength effect: II. theoretical mechanisms. *Journal of Experimental Psychology: LMC*, 16(2), 179.
- Wilson, M. A., & McNaughton, B. L. (1994). Reactivation of hippocampal ensemble memories during sleep. *Science*, 265(5172), 676–679.