

Running head: UNILATERAL VS. BILATERAL DAMAGE

Why bilateral damage is worse than unilateral damage to the brain

Anna C. SCHAPIRO^{1,2}, James L. MCCLELLAND¹, Stephen R. WELBOURNE³, Timothy T. ROGERS⁴,

Matthew A. LAMBON RALPH³

¹Department of Psychology, Stanford University

²Department of Psychology and Neuroscience Institute, Princeton University

³Neuroscience & Aphasia Research Unit, School of Psychological Sciences,

University of Manchester

⁴Department of Psychology, University of Wisconsin-Madison

Corresponding author:

Anna C. Schapiro
Department of Psychology
Green Hall
Princeton University
Princeton, NJ 08540
+1-617-797-4555
schapiro@princeton.edu

or Prof. Matthew A. Lambon Ralph
Neuroscience & Aphasia Research Unit
School of Psychological Sciences
Zochonis Building, Brunswick Street
Manchester, M13 9PL
+44 161 275 2551
matt.lambon-ralph@manchester.ac.uk

Abstract

Human and animal lesion studies have shown that behavior can be catastrophically impaired after bilateral lesions but that unilateral damage often produces little or no effect, even controlling for lesion extent. This pattern is found across many different sensory, motor, and memory domains. Despite these findings, there has been no systematic, computational explanation. We found that the same striking difference between unilateral and bilateral damage emerged in a distributed, recurrent attractor neural network. The difference persists in simple feedforward networks, where it can be understood in explicit quantitative terms. In essence, damage both distorts and reduces the magnitude of relevant activity in each hemisphere. Unilateral damage reduces the relative magnitude of the contribution to performance of the damaged side, allowing the intact side to dominate performance. In contrast, balanced bilateral damage distorts representations on both sides, which contribute equally, resulting in degraded performance. The model's ability to account for relevant patient data suggests that mechanisms similar to those in the model may operate in the brain.

1. Introduction

There are indications from decades of research and clinical study that a unilateral lesion to a brain area can produce a very subtle – sometimes clinically insignificant – neurocognitive impairment, whereas a bilateral lesion may generate a profound deficit. Perhaps the best known case example arose in the 1950's, when Scoville produced a profound memory loss in patient HM by removing the hippocampus bilaterally. Scoville's motivation for publishing these findings was not only to warn other neurosurgeons of these clinically dramatic consequences but also to note that they ran counter to the expectations arising from patients after unilateral resections, in whom the observed memory deficit is relatively mild (Milner & Penfield, 1955; Scoville & Milner, 2000).

Why are there such different consequences of unilateral and bilateral lesions? One possibility is that a bilateral neural system is simply redundant enough across hemispheres to accommodate the amount of damage corresponding to complete unilateral removal. While this may be part of the story, we will review findings strongly suggesting that bilateral lesions are much more detrimental than unilateral lesions even when unilateral lesions involve the same total amount of damage as bilateral lesions. This leads us to consider the possibility that there is something special about unilateral relative to bilateral damage. The computational exploration reported here provides evidence for this possibility, demonstrating that bilateral damage is much more detrimental than unilateral damage in a range of distributed neural network models even when the amount of damage is equated.

In order to anchor the investigation with a concrete, non-trivial test case, we applied the computational work to the domain of semantic memory. This choice was motivated by the availability of a considerable body of relevant patient and neuroscience data. In addition, the new computational investigation built on earlier modelling work in which we employed a single pool of units to simulate effects of bilateral anterior temporal atrophy in semantic dementia (SD;

Rogers et al., 2004). The findings we obtain generalize across tasks and to a range of different network architectures. Our results may help, therefore, to explain the different effects of unilateral and bilateral damage that have been reported in many human and animal studies across a range of task domains and brain areas. Such differences apply not only to effects of damage to the human hippocampus (Scoville & Milner, 2000) and anterior temporal cortex as examined in detail here, but also to effects of damage in, for example, rat hippocampus (Li, Matsumoto, & Watanabe, 1999), primate anterior temporal cortex (Klüver & Bucy, 1939), primate frontal lobe (Warden, Barrera, & Galt, 1942), primate auditory cortex (Heffner & Heffner, 1986), human frontal lobe (D'Esposito, Cooney, Gazzaley, Gibbs, & Postle, 2006), and human amygdala (Adolphs & Tranel, 2004).

We note, in advance, that our results will apply to effects of unilateral vs. bilateral damage to the extent that there is a balance of involvement of brain structures across the hemispheres. In cases where a function is completely lateralized, damage to the hemisphere that specializes in the function will naturally be more disruptive to performance than balanced bilateral damage. But, as we shall show, a degree of asymmetry can still be consistent with a relatively milder deficit after unilateral damage, even to the more dominant hemisphere. We will demonstrate this by examining the effect of unilateral versus bilateral damage on two semantic tasks. Performance in one of these tasks (word-to-picture matching) appears to be subserved by the anterior temporal cortex approximately symmetrically, while performance in the other (naming) is underpinned to a greater degree by the left anterior temporal cortex (Lambon Ralph, Ehsan, Baker, & Rogers, 2012; Lambon Ralph, McClelland, Patterson, Galton, & Hodges, 2001; Seidenberg et al., 1998), which is captured in the model before lesioning through a mild structural asymmetry.

1.1 The role of anterior temporal regions in semantic representation

As noted above, we chose the domain of semantic memory as a test case for the

simulations and, in particular, the role of left and right anterior temporal lobe (ATL) regions in supporting semantic representations. The degradation of semantic memory in SD and herpes simplex virus encephalitis (HSVE) is associated with bilateral damage to and hypo-metabolism of the anterior temporal lobes (Mion et al., 2010; Nestor, Fryer, & Hodges, 2006; Noppeney et al., 2007; Rohrer et al., 2009). Behavioral data from these patients have suggested a model in which concepts are formed through the convergence of sensory, motor, and verbal experience in a transmodal representational hub in the ATL (Rogers et al., 2004).

Although previously overlooked, there is now a growing consensus that this transmodal ATL hub contributes critically to semantic cognition (Patterson, Nestor, & Rogers, 2007). This emerging view reflects a convergence of the established clinical data on SD, HSVE, etc., with contemporary neuroscience studies. Data from transcranial magnetic stimulation and functional neuroimaging indicate that both left and right ATL regions contribute to semantic memory. For example, the multimodal, selective semantic impairment of SD can be mimicked in neurologically-intact participants by applying rTMS to the left or right lateral ATL (Lambon Ralph, Pobric, & Jefferies, 2009; Pobric, Jefferies, & Lambon Ralph, 2007, 2010). Likewise, when using techniques that avoid (e.g., PET or MEG) or correct for the various methodological issues associated with successful imaging of the ATL (Devlin et al., 2000; Visser, Jefferies, & Lambon Ralph, 2010), studies find bilateral ATL activation for semantic processing across multiple verbal and nonverbal modalities (Binney, Embleton, Jefferies, Parker, & Lambon Ralph, 2010; Marinkovic et al., 2003; Sharp, Scott, & Wise, 2004; Vandenberghe, Price, Wise, Josephs, & Frackowiak, 1996; Visser, Embleton, Jefferies, Parker, & Lambon Ralph, 2010; Visser & Lambon Ralph, 2011).

Figure 1 - about here

This domain is a highly pertinent one for the current investigation given that unilateral and bilateral ATL damage have very different consequences on semantic performance (mirroring the

pattern found in other cognitive domains). This conundrum is summarized in Figure 1. In terms of clinical presentation, the bilateral atrophy of SD (Figure 1A) leads to a notable, selective semantic impairment (Patterson & Hodges, 1992; Snowden, Goulding, & Neary, 1989; Warrington, 1975). In contrast, patients with unilateral ATL damage (in this example, Figure 1B, a full *en bloc* ATL resection for intractable temporal lobe epilepsy; TLE) do not report comprehension impairment (Hermann, Seidenberg, Haltiner, & Wyler, 1995; Hermann et al., 1994) but do complain of amnesia and anomia, especially following left ATL resection (Glosser, Salvucci, & Chiaravalloti, 2003; Martin et al., 1998; Seidenberg et al., 1998). The two contrastive patients in Figure 1 are particularly pertinent given that the unilateral case represents a complete removal of ATL tissue (Figure 1B) but had little or no drop in accuracy on standard semantic tests (see below), whereas the bilateral case represents only partial tissue loss (Figure 1A) but exhibited considerable semantic impairment. Although an exact comparison of the extent of relevant tissue loss is not possible, the notable difference in the extent of deficit despite comparable lesion size suggests that size is only part of the story.

These clinical observations are striking but need to be treated with some caution given that clinical assessment of TLE patients tends to focus on naming and episodic memory, and rarely on comprehension (Giovagnoli, Erbetta, Villani, & Avanzini, 2005). Two recent case-series neuropsychological studies, therefore, assessed this clinically-apparent contrast via in-depth explorations of semantic function. The results not only supported the clinical observations but also provided a bridge to the neuroimaging and rTMS data noted above. The first study (Lambon Ralph, Cipolotti, Manes, & Patterson, 2010) investigated patients' accuracy on a range of tasks taken from a multimodal semantic battery, commonly used to assess patients with SD (Bozeat, Lambon Ralph, Patterson, Garrard, & Hodges, 2000). A case-series of patients with unilateral left or right ATL damage (of mixed aetiology) performed at or very close to normal levels of accuracy on the various comprehension tasks, whereas most SD patients were impaired

(Figure 2). In addition, like SD patients with asymmetrically left-biased ATL atrophy (Lambon Ralph et al., 2001), the left unilaterally-damaged patients exhibited significantly greater anomia than their right-sided counterparts. The second study focussed specifically on patients with left or right-sided ATL resections for the treatment of TLE (Lambon Ralph et al., 2012). Again, on standard accuracy-based measures of semantic function, the vast majority of the resected TLE patients performed within the normal range. However, mild semantic dysfunction was revealed when either the semantic demands of the tasks were increased (by probing specific-level concepts, low frequency, or abstract items) or by measuring reaction times as well as accuracy (the patients were 2-3 times slower than older controls even on the simpler semantic tasks). Performance on these more demanding semantic measures correlated significantly with the volume of resected tissue and, again, the left-resected patients exhibited significantly greater anomia than the right-sided cases. Together, these results provided us with a series of target phenomena to be captured in the computational model:

- i. Regions within the left and right ATL jointly support pan-modal semantic function in neurologically-intact participants;
- ii. Unilateral resection or stimulation via rTMS generates detectable but mild semantic impairment;
- iii. Bilateral lesions generate significantly greater semantic impairment than unilateral damage, even when overall amount of damage is similar;
- iv. Left ATL damage generates greater naming impairment than equivalent levels of right ATL damage.

In order to fully account for the range of lesion scenarios, we also explored the possibility that, in some patients, plasticity-related recovery after or during the course of brain damage may have affected post-lesion performance (Keidel, Welbourne, & Lambon Ralph, 2010; Welbourne & Lambon Ralph, 2005; Wilkins & Moscovitch, 1978). Patients continue to engage in semantic

processing after an acute lesion or, in degenerative cases, as their disease progresses, and this engagement may help to ameliorate the effects of damage (Welbourne & Lambon Ralph, 2005). We considered the effects of plasticity-related recovery in the model to demonstrate its applicability to these scenarios.

Figure 2 - about here

1.2 Overview of the model

The model we used to account for the stark contrast between the behavior of patients with unilateral and bilateral damage is an extension of a model used previously to explain the deterioration of semantic task performance in SD (Rogers et al., 2004). It consisted of processing units that were organized into layers and connected as shown in Figure 3. The connection weights between the units were established by training the network with patterns derived from descriptions and visual properties of familiar objects. The training allowed the network to map between different features of these objects and to capture the similarity relations of the objects to one another. The input to the network came from the *verbal descriptor* and *visual feature* layers, which are considered analogous to cortical areas that subserve high-level verbal and visual information, respectively. The verbal descriptors were divided into sublayers that represent names, verbal descriptions of perceptual features, functional descriptors, and encyclopedic information.

The model extended the previous simulation architecture by dividing what was previously a single pool of hidden units into two: *left demi-hub* and *right demi-hub*. Units in these pools had no pre-specified representations – rather, the pools developed distributed representations over the course of training that allowed the network to map between the features of the objects in the environment. The layers of the model were highly interconnected, with complete bidirectional connectivity (all the units in one layer connected bidirectionally to all the units in another layer) between *verbal descriptors* and *left demi-hub*, *visual features* and *left demi-hub*, and *visual*

features and *right demi-hub*. The projection from the *verbal descriptors* to the *right demi-hub* was complete but the return projection had about 80% of all the possible connections. The difference between density of projection from the *left demi-hub* to *verbal descriptors* and that from the *right demi-hub* to *verbal descriptors* caused the *left demi-hub*, over the course of training, to play a somewhat more important role in tasks like naming from visual input that require a verbal response, aligning this model with our previous one on laterality effects in SD (Lambon Ralph et al., 2001).

An important feature of the model was the stipulation that the density of the recurrent connections within each hidden layer was 90% while the density of the projection between the hidden layers was only 10%. This stipulation is consistent with the fact that there are more intra-hemispheric than inter-hemispheric connections in cortex (Braitenberg & Schüz, 1998; Gee et al., 2011; Lewis, Theilmann, Sereno, & Townsend, 2009). As we shall see, our model suggests that this difference may be an important part of the reason why bilateral damage is more disruptive than unilateral in networks (like human neocortex) with recurrent connections.

In what follows, we first demonstrate that our model simulates an advantage of unilateral compared to bilateral damage when overall lesion severity is equated. We proceed to consider this phenomenon in a range of network architecture variants, one of which is sufficiently simple to allow us to express mathematically the fundamental basis for the advantage of unilateral lesions. Our analysis shows why the advantage occurs in a wide range of neural network architectures and, therefore, why it is natural to expect that this advantage should occur in real biological neural networks such as those found in the brains of humans and other animals.

Figure 3 - about here

2. Materials and Methods

The reported simulations were conducted using the LENS neural network simulation environment (<http://www.stanford.edu/group/mbc/LENSManual/index.html>). A network very

similar to the one used by Rogers et al. (2004) was re-implemented within this environment. The following describes the main simulation reported first in the *Results* section. Other simulations used similar methods; where details differ from the main simulation, this is stated in the *Results* section.

2.1 Network initialization

In setting up the network, complete projections were created by simply connecting each unit in the projecting layer to each unit in the receiving layer. For sparser projections, for example the 10% connection between the two hidden layers, each unit in one hidden layer was connected to each unit in the other hidden layer with an independent probability of 0.1. As a result, each instantiation of a network had a different pattern of connectivity and a somewhat variable proportion of connections. Weights were also assigned randomly before training, with values chosen uniformly between -1 and 1.

2.2 Activation dynamics

Each unit in the verbal descriptors and visual features layers represented a unique verbal or visual property of an object. To present an object to the network, the units representing the features of the object took on the value 1, and otherwise had the value 0. The presentation of an item to the network was divided into seven time intervals; in each interval, activation of each unit i was calculated by first determining the net input, n_i , to the unit:

$$n_i = \sum_j a_j w_{ij}$$

where j indexes all units that project to unit i , a_j is the activation of unit j , and w_{ij} is the weight of the connection from unit j to unit i . The activation of unit i was then calculated based on the following logistic function, which bounds activation between 0 and 1:

$$a_i = \frac{1}{1 + e^{-n_i}}$$

All units were assigned a fixed, untrainable bias of -2, which deducts 2 from each unit's net input. This bias parameter, in the absence of other input, keeps unit activations on the low end of their range (at approximately 0.12).

2.3 Training

The network was trained and tested using 48 patterns organized into six natural taxonomic categories (birds, mammals, vehicles, household objects, tools, fruits). Patterns were generated from the visual feature and verbal descriptor prototypes used by Rogers et al. (2004), which were designed to reflect the similarity relations of typical objects in the six categories. The similarity of patterns within categories varies across categories as it does for real objects in the corresponding categories. Patterns were constructed exactly as in Rogers et al. (2004) with the exception that that object names were not given at different levels of specificity in the present work – each of the 48 objects was assigned its own unique name.

Training of the network involved teaching it to 'bring to mind' the full information associated with an item, either from the name, the visual pattern, or from non-name verbal descriptors of the item. Accordingly, in each training trial, the units corresponding to one of these three inputs (units in the name, other verbal descriptors, or visual features layer) were set (hard-clamped) to the specified input values for three intervals of processing. The clamps were then removed and the network cycled for four more intervals, so that all unit activations were then determined by the propagation of activation through the network's connections. During the last two intervals of the seven total, the target values for all visual and verbal patterns (including name) were applied and error derivatives for weights in the network were calculated. Training was conducted using the standard back propagation gradient descent learning algorithm for recurrent networks in LENS, with 200,000 sweeps through each of the 48 training patterns presented in a random order (each sweep is called an *epoch*). Weights were updated at the end of

each epoch, except during recovery, when a finer grain of detail is necessary and weights were updated after every pattern. The model was trained with a learning rate of 0.005, weight decay of 0.0002 per epoch during training (0.000001 per pattern during recovery), and no momentum. Zero-mean 0.5 standard deviation Gaussian noise was added to the net input of each unit at each activation update. Gaussian noise with the same standard deviation was also included during testing unless otherwise specified.

2.4 Testing

Testing for picture naming assessed the network's ability to activate the correct name unit from visual input. This was implemented by clamping the visual feature units to the values corresponding to a trained object, allowing activation to propagate for three intervals, then removing the input and allowing activation to continue to propagate for four more intervals. The unit with the highest activation was treated as the model's name response to the presented picture.

The test for word-to-picture matching compared the pattern across the hidden layers produced by an item's name with the patterns produced by several alternative visual input patterns. A trial was scored correct if the pattern produced by the name for an item was more similar to the pattern produced by the visual feature input pattern for the same item than it was to the pattern produced by the visual input pattern for seven distractor items. To accomplish this, the target name unit was clamped for the first three processing intervals; input was then removed and activation continued to propagate for four intervals. The resulting pattern of activation across all units in both hidden layers was recorded. This process was repeated with visual feature input for the target concept and each of the seven remaining items in the same semantic category (acting as the foils in the assessment). The resulting patterns of hidden unit activation can be thought of as nine vectors (one resulting from presentation of the name and eight from presentation of the visual features), where each unit's activation is an element of the vector. The cosine between the

vector from the name input and the vector from each of the visual feature inputs was used as the measure of similarity.

The cosine indicates the degree of similarity between two vectors, independent of their magnitudes, and is defined as the dot product of the two vectors divided by the product of their magnitudes:

$$\cos(V_1, V_2) = \frac{V_1 \cdot V_2}{\|V_1\| \|V_2\|}$$

where $V_1 \cdot V_2 = \sum_i V_{1i} V_{2i}$ and $\|V\| = \sqrt{\sum_i V_i^2}$. The set of visual features that resulted in the highest cosine with the name was interpreted as the model's choice of picture.

2.5 Simulation of behavioral deficits under damage

To investigate the effects of damage on task performance, the model was damaged at several levels of severity, by permanently removing connections into and out of a layer with a specified probability. Damage of, for example, 50% to a hidden layer consisted of removing connections into and out of each unit in that layer with probability 0.5. For progressive damage simulations, where lesions were applied incrementally and the network was tested as damage progressed, damage levels refer to the total cumulative proportion of connections removed up to that point. Bias weights were not lesioned, though results were the same when lesions to bias weights were allowed.

In one of the simulations reported below, damage was simulated by the removal of whole units rather than individual connections. In this case, damage to a layer consisted of unit removal with a specified probability. As with other forms of damage, removal was probabilistic, such that the actual number of units removed could vary.

Damage was applied to one or both of the hidden (demi-hub) layers. For a given severity level, a bilateral lesion was produced by applying this probability uniformly across the two hidden layers, while unilateral left and right lesions were produced by applying damage with

twice this probability to the connections into and out of only one of the two hidden layers. Following this procedure, the largest possible unilateral lesion corresponded to a 50% bilateral lesion. However, we extended bilateral lesions beyond this severity level to 100% in order to simulate more severe stages of SD and other bilateral diseases. In the case of bilateral damage, the connections between the two hidden layers were removed at the associated unilateral damage level, in order to ensure that the total number of connections was approximately equal in the unilateral and bilateral damage cases. For example, if 50% of the connections between hidden layers were removed for a unilateral lesion, 50% of those connections were again removed in the corresponding case of bilateral damage (as opposed to 25% removal twice – one for each hidden layer – which would result in only 44% of these connections being lesioned).

To examine how experience after damage or during the course of degeneration might affect performance in the model, we conducted simulations using varying amounts of recovery after performing a simulated lesion. During this progressive damage, retraining occurred after each round of lesioning, and the network was tested immediately before and after an episode of damage and the results of the two tests were averaged. This was done to approximate the more continuous mix of damage and exposure to the environment that patients with neurodegenerative disease experience.

Results reflect the performance of ten trained networks initialized with different random seeds. Each network was tested after each combination of lesion type (left, right, or bilateral), severity, and recovery amount (where these were varied) and the results were averaged over the ten separate networks.

Figure 4 - about here

3. Results

3.1 Differential effects of unilateral vs. bilateral damage

Our first simulation demonstrates a striking difference between the effects of unilateral

and bilateral lesions in the recurrent network architecture shown in Figure 3. Figure 4 summarizes simulated picture naming and word-to-picture matching performance as a function of different levels of unilateral left, unilateral right, and bilateral connection damage following training of the recurrent network. Results are presented separately for different amounts of recovery after each simulated lesion. Networks with bilateral damage performed much worse than those with unilateral damage on both tasks, even with the same total amount of damage and recovery (see Table 1 for statistics). Bilateral damage caused performance to fall relatively quickly, whereas even a complete unilateral lesion, especially with some recovery, had a relatively small effect. The difference between the effect of unilateral and bilateral damage became more pronounced with increasing lesion severity.

One additional feature of these results is that in none of these cases did performance drop off immediately as soon as there was any damage. The use of distributed representations in the model, encouraged by the presence of noise and weight decay in training and by retraining after damage, made the model quite robust. This *graceful degradation* (Farah & McClelland, 1991; Hinton & Sejnowski, 1983) is clear in patients as well – SD patients demonstrate a drop in semantic performance over time and have *some* remaining semantic ability even after extensive atrophy.

Table 1 - about here

There was also a difference, especially in naming, between left-sided and right-sided unilateral damage. Because of the somewhat greater role of the *left demi-hub* in mapping visual to verbal features, damage to the left side of the network was more detrimental to naming than damage to the right side. There was also a small difference between left-sided and right-sided unilateral damage in word-to-picture matching.

The simulation results presented so far employed cumulative damage with retraining after each stage of lesioning. Some of the unilateral patients, e.g., those with long-standing epilepsy

followed by resection, had damage that was only partially progressive and some, e.g., those with stroke, had damage that would not have been progressive at all. We ran additional simulations to test whether the model would perform differently if damage were applied all at once, followed by a period of recovery (to mimic an acute neurological event followed by spontaneous partial recovery) instead of interleaving damage and retraining. The results of the non-progressive damage simulations were very similar to those for the progressive damage simulations (see Appendix 1).

Figure 5 - about here

3.2 Effect of lesions on similarity structure

Why does a unilateral lesion have less effect on performance than the same severity of lesion distributed bilaterally? To understand these differential effects of unilateral and bilateral damage in the model, it is useful to examine how the similarity structure of its internal representations changed with damage, since the distortions to this similarity structure were the source of errors in word-to-picture matching and contributed to errors in naming (for naming, weights from the hidden layers to the name layer also contribute). For this analysis, we used a network architecture in which the differential connectivity from the left and right demi-hubs to verbal output was eliminated for simplicity, and no recovery was allowed, though similar results were obtained with an architecture that retained the differential connectivity and for networks with various levels of recovery.

We obtained the pattern of activation across both demi-hubs for presentations of each object's name and each object's visual input pattern, as in the word-to-picture matching task, then calculated the cosine similarity of the pattern produced by each name with the pattern produced by each visual input pattern, averaging the results across ten networks. This yielded a 48×48 matrix, with each item corresponding to a row of the matrix (visual feature input) and a column (name input), as shown in Figure 5. Objects in the same category are grouped together. High

cosine values along the diagonal of the matrix for the intact network (left-most panel of figure) indicated that the network was able to associate an object's visual features correctly with its name. While the cosines are generally higher within than between category, the correct (diagonal) entry is strongest in each row and each column before damage. As damage increased, the similarity structure of the unilaterally-lesioned network was fairly well preserved, allowing it to distinguish an item from the others in its category with fairly high accuracy, even after a complete lesion to one side of the network. The bilaterally damaged network, in contrast, lost the appropriate similarity structure rapidly; With a 50% lesion to each side of the network, the similarity structure seen with the intact network was nearly completely obliterated, accounting for its very poor performance.

Figure 6 - about here

3.3 Simpler networks and representation analyses

The preceding analysis shows that bilateral damage degrades the similarity structure of the network's internal representations, whereas unilateral damages leaves this largely intact. What causes this phenomenon? The bidirectional and recurrent connections in the current model make it difficult to answer this question, so we next considered a simpler network architecture in which all connections were feedforward (i.e., projecting in one direction, from input toward output).

Feedforward network. In the feedforward network, each of the layers that was previously used as both input and output was duplicated so that one copy could be used for the input and the other for the output (see Figure 6). In addition, there were no recurrent connections within or between the two hidden layers, no recovery was allowed, and the connections into and out of the two hidden layers were complete and symmetric. Although training affects both input and output connections to the hidden units in this model, only the input connections affect the internal representations, allowing a focused analysis of the fundamental basis for the unilateral vs. bilateral differences in word-to-picture matching in this version of model. We trained ten

networks initialized with different random seeds as in the main simulation, then applied unilateral and bilateral damage to this feedforward architecture by removing connections as in previous simulations. During training and testing, activations were calculated in a single feedforward pass, and noise was included in processing as per the main simulation. Even in this simplified network there was still an advantage for unilateral compared with bilateral damage, both in word-to-picture matching and naming (Appendix 2).

Figure 7 - about here

Representation fidelity. We chose a measure which we call *fidelity* to quantify the loss of similarity structure in this network. Since the relative values of the on- and off-diagonal cosines in the matrices shown in Figure 5 govern the network's ability to perform semantic tasks like naming and word-to-picture matching, fidelity is defined as the difference between mean diagonal and off-diagonal cosine terms, including both within- and between-category off diagonal entries (see Figure 7 caption for details). Although fidelity does not correspond directly to accuracy, greater fidelity is associated with more accurate performance. This measure is useful because it can be applied to the left and right semantic layers separately to provide an assessment of the extent to which each side carries the appropriate similarity information on its own. Figure 7A shows the fidelity of an intact semantic layer on its own, a damaged semantic layer, the average of those two, and compares this to the average of bilaterally damaged semantic layers (each of which produces fidelity similar to this average). We found that within a layer, fidelity was approximately equal to the fraction of connections remaining, with the result that the average fidelity was very similar in the unilateral and bilateral cases, equating for the total number of connections removed.

If average fidelity across the left and right semantic layers is as impaired after a unilateral lesion as it is after a bilateral lesion, why then is word-to-picture matching and naming more accurate after the unilateral lesion? The reason is that the fidelity calculated across both layers

remains much higher in the unilateral case. Figure 7B shows the fidelity measure calculated across both layers of the feedforward network, either with unilateral damage or bilateral damage. As in word-to-picture matching and naming performance, we found that the fidelity was always higher in the unilateral than bilateral case (see Table 2). Interestingly, fidelity followed a U-shaped function, decreasing over low to moderate levels of damage, then increasing again as the extent of damage approached its maximal value. This occurs because variation from item to item in the pattern of activation of the units in the damaged layer diminishes as the fidelity decreases, and so these units contribute less and less to the overall fidelity of the representations. By the time their contributions become very weak, the intact side of the network completely dominates performance. Thus, a layer with a large amount of unilateral damage carries little information about the items but also contributes little to the overall performance of the network because it does not provide a differential response across items. This is the essence of the reason why unilateral damage is so much more benign than bilateral damage: With more and more damage to just one side of the network, its own fidelity decreases, but the magnitude of its contribution to differences among the representations of different items also decreases, reducing its impact on overall fidelity, and therefore overall performance.

Table 2 - about here

Simplified mathematical analysis. To see this effect in mathematical terms, we considered the effect of lesions on the cosine of the vectors of net inputs (excluding bias) to units in the hidden layer in the simplified feedforward architecture. Although activations, and not net inputs, were used in the network testing and fidelity calculations reported thus far, the net input is the signal that produces the differences among representations of different items. Therefore, the analysis of net inputs captures the key elements of the difference between unilateral and bilateral lesions. The activation of each unit was a monotonic transformation of its net input that tended to preserve similarity relationships,

and fidelity calculated on net inputs instead of activations yielded very similar results to those in Figure 7.

To begin our examination of the net input cosines, let V_v be the net input across hidden layer units resulting from presentation of a particular object's visual features and let V_n be the net input across hidden layer units resulting from presentation of that same object's name. Each of these vectors can be broken up into the pattern of activation for the left half of the hidden units, V_{vl} and V_{nl} , and that for the right half, V_{vr} and V_{nr} . The cosine between the two complete vectors can be represented

$$\cos(V_v, V_n) = \frac{\sum_i V_{vli} V_{nli} + \sum_i V_{vri} V_{nri}}{\sqrt{(\sum_i V_{vli}^2 + \sum_i V_{vri}^2)(\sum_i V_{nli}^2 + \sum_i V_{nri}^2)}}$$

Note that each of the two sums in the numerator corresponds to the dot product of the corresponding vectors, and (from the definition of the cosine, in *Materials and Methods*) the dot product of two vectors is equal to the cosine between the vectors times the product of the vectors' magnitudes. Similarly, each summation in the denominator corresponds to the square of the magnitude of the corresponding vector. Thus the above equation can be rewritten:

$$\cos(V_v, V_n) = \frac{\|V_{vl}\| \|V_{nl}\| \cos(V_{vl}, V_{nl}) + \|V_{vr}\| \|V_{nr}\| \cos(V_{vr}, V_{nr})}{\sqrt{(\|V_{vl}\|^2 + \|V_{vr}\|^2)(\|V_{nl}\|^2 + \|V_{nr}\|^2)}}$$

We call this the *magnitude-weighted cosine equation*. We can now consider how unilateral vs. bilateral damage affects this cosine, by understanding how it affects the contributing magnitudes and cosines.

First, we consider in general terms the effect of damage to the weights coming into a common layer from two sources, corresponding to the name and the visual input to a pool of hidden units, on the net input cosines and magnitudes. The cosine reflects the degree to which the two net input vectors are similar. We consider the ideal case in which (before the

lesion) the cosines on the left and right are both equal to 1 (i.e., the visual and name net inputs are identical). We also assume the left and right vectors have the same number of elements (same number of units on the left and on the right). We further assume that a lesion affects incoming weights from the name and visual units with equal probability p , so that the expected fraction of remaining incoming name and visual weights is $R = 1 - p$. These assumptions hold for the lesions we have performed on our networks. We can now examine how a lesion degrades this similarity relationship. As shown in Figure 8A, the average value after a lesion sparing proportion R of the weights (denoted by $\langle \rangle_R$) of the cosine between two vectors is simply equal to R times the original value of the cosine:

$$\langle \cos(V_v, V_n) \rangle_R = R \cos(V_v, V_n)$$

This *cosine shrinkage equation* applies to the effect of a lesion to incoming weights on the cosine over either half of the hidden units, or to a uniform bilateral lesion over the weights coming to all of the hidden units.¹

To see how the cosine is affected by a unilateral or, more generally, an unequal lesion, we must consider the effect of damage to weights on the magnitudes of the contributing net input vectors, as well as the effect on the contributing vectors' cosines. The average effect of a lesion leaving R remaining weights on the magnitude of a net input vector is $\sqrt{R}\|V\|$, where $\|V\|$ represents the vector's magnitude before the lesion.² For a lesion leaving the same proportion R of the weights contributing to the net input vector produced by the name and the net input vector produced by the visual input pattern for a

¹ Note that the cosine shrinkage equation is an empirical characterization of the average effect of a lesion leaving a remaining fraction R of the existing connection weights. When the incoming weights to a given hidden unit from active input units are correlated, the

² As with the effect of a lesion on the cosine, this is an empirical observation. In this case, if the incoming weights to each receiving unit from each active input unit are highly correlated, the average effect of a lesion leaving R remaining weights on the length of a net input vector is $R\|V\|$, and not $\sqrt{R}\|V\|$. The observed value ($\sqrt{R}\|V\|$) is once again what is observed if there is just one active input unit, or if the weights coming into each hidden unit from the ensemble of active input units are uncorrelated.

given item, we see that the average value of the product of the two vectors' magnitudes after the lesion $\langle ||V_v|| ||V_n|| \rangle_R$ is equal to $R ||V_v|| ||V_n||$, as shown in Figure 8b.

Now, we can consider separately the effect of a unilateral lesion *versus* a bilateral lesion. The remaining fraction of weights on each side, R_r and R_l , determine the relative weight of the right-side and left-side cosine in the value of the overall cosine (again, on each side, we treat the lesion as affecting the incoming name and visual weights equally). We assume that, before the lesion, the name and visual vectors have the same magnitude, and that the magnitude of the left half of each vector is the same on average as the magnitude of the right half of each vector. Effectively this means that the magnitudes of the four vectors $||V_{vl}||$, $||V_{vr}||$, $||V_{nl}||$, and $||V_{nr}||$ all have the same value M . Combining these assumptions with the finding above that $\langle ||V|| \rangle_R = \sqrt{R} ||V||$, replacing R with R_l for the effect of the lesion on the magnitudes of the left side vectors and replacing R with R_r for the effect of the lesion on the magnitudes of the right side vectors, substituting into the magnitude-weighted cosine equation and simplifying, we find that

$$\langle \cos(V_v, V_n) \rangle_{R_l, R_r} = \frac{R_l}{R_l + R_r} \langle \cos(V_{vl}, V_{nl}) \rangle_{R_l} + \frac{R_r}{R_l + R_r} \langle \cos(V_{vr}, V_{nr}) \rangle_{R_r}$$

That is, after asymmetric damage, the relative weight of the left and right sides, w_l and w_r are $w_l = R_l/(R_r + R_l)$ and $w_r = R_r/(R_r + R_l)$ respectively, where R_r and R_l are the proportions of connections remaining on the right and left sides.

In this equation, consider the effect of removal of some proportion of the connections coming into the right half of the hidden units. This will reduce the magnitudes of the net inputs to the right hidden units, and correspondingly it will therefore reduce the relative weight of the right-side cosine – in the extreme case of a total right-sided lesion, sparing all of the weights on the left, ($R_l = 1$, $R_r = 0$), the right-side contribution goes to 0 and we are left only with the intact left side contribution, unaffected by the lesion. Crucially,

in a bilaterally damaged network, where after the lesion V_{vl} and V_{nl} would have the same magnitude as V_{vr} and V_{nr} , the magnitude normalization would not result in domination by either side, and the distortion to the overall cosine would be the same as the equivalent distortion to each of the two cosines. These effects are depicted in Figure 8C, where we plot the overall cosine resulting from a pure one-sided lesion and the overall cosine resulting from an equal bilateral lesion, under the further simplifying assumption that the right and left sided cosines are both equal to 1 before the lesion. In the latter case the cosine falls linearly with damage; in the former case it actually follows a U-shaped pattern, decreasing then increasing again as damage weakens the relative contribution of units on the right hand side of the network. Once again the equation closely matches the result obtained by simulation.

Figure 8 - about here

The results of the analysis just described are closely approximated by the calculated values of the cosines after simulations of damage to the network, as shown by the agreement of the theoretical and simulation-based curves shown in Figure 8. (For the simulation results, the original values of the separate left, right, and overall cosines are not exactly equal to 1, since the net inputs from the name and visual units are not exactly identical. The results shown are normalized by the pre-lesion value of the relevant cosine term – left, right, or overall.) As with fidelity, the expected (average) value of the cosine is not a direct predictor of performance, since performance is based on the relative values of particular cosines whose values vary from case to case due to the random number of affected connections for each input pattern. Nevertheless, they provide a clear demonstration of the reason why a unilateral lesion has a relatively benign effect: while the lesioned side may be grossly affected, its contribution to performance is reduced as a consequence of the lesion.

3.4 Fidelity and the unilateral vs. bilateral advantage in the recurrent network

Returning now to the more realistic recurrent architecture (from section 3.2), we can apply the fidelity analysis as before. Figure 7C shows the decreasing fidelity for unilateral and bilaterally damaged networks at increasing levels of damage. These curves resemble those for performance on word-to-picture matching and naming (Figure 4), indicating that these analyses provide a good account of the basis for the difference between unilateral and bilateral performance in those tasks.

It is interesting to consider why the fidelity falls more rapidly overall with damage in the recurrent network than in the feedforward network. This can be understood by noticing that in the recurrent network, a lesion to connections coming into hidden units on one side has a compounding effect. The first update to the activations of the hidden units in both demi-hubs is affected only by the damage to the connections from the pool that has inputs clamped, just as in the feedforward network. The distortion to unit activations is propagated via these recurrent connection weights which themselves have been damaged, thereby increasing the distortion. Note that the propagated distortion affects both hubs to a degree, but because the extent of connectivity is greater within a hub than between hubs, the damaged hub has fairly little effect on the hub that was not damaged, contributing to the unilateral advantage.

3.5 Effect of damage to units in a recurrent network

Our simulations thus far have considered the effects of connection damage, but it might be argued that damage to whole units would be a more realistic model of damage in the brain, since degenerative illness, stroke, or resection might affect whole neurons or groups of neurons. To extend our findings to address this possibility, we ran damage simulations removing units instead of connections, where removal of a unit effectively removed all the connections into and out of that unit at once. We tested 10 networks using the recurrent network architecture, using a larger total number of hidden units (400 instead of 64) and longer training time (400,000 epochs). For simplicity, this simulation employed full connectivity within the left and the right side semantic

layers, no connections between the left and right sides, and no asymmetry in the output from hidden to verbal descriptors. We chose the larger number of hidden units since information is often not distributed very evenly across units in smaller networks and, consequently, removal of a unit can cause large network-specific changes that would not be expected in networks – like the brain – with much larger numbers of units.³ The fidelity measure showed that unilateral damage was again less detrimental than bilateral damage and that the difference between the unilateral and bilateral damage increased with increasing amounts of unit damage (Figure 7D). These effects on fidelity carry over, as in other cases, to word-to-picture matching and naming.

The difference between unilateral and bilateral damage here arises entirely from the difference in recurrent connections within a hub vs. between hubs. Recall that the net input to a unit is equal to a sum of terms, each of which is the activation of a sending unit multiplied by its connection weight to the receiving unit. Removal of units within a hub removes some of the terms from the net input to other units within the same hub, just as removal of connection weights among the units within a hub would. Because there is greater connectivity among hidden units within a hub than between hubs, the effect of unit removal is largely confined to units in the same hub. There can be indirect effects via recurrent influences propagating throughout the network, but as before, in the limit of removing all of the units on one side of the network, neither signal nor noise are propagated by this side of the network. In summary, a lesioned unit causes relatively focal severing of connections to other units in the same hub. As with connection damage, when lesioning is restricted to one side, that side's contribution is both weakened and distorted, allowing the intact side to dominate, and thereby producing less of an overall deficit than a comparable amount of damage spread evenly across both sides.

3.6 Within versus between hemisphere redundancy

³ Because there are far more connections than units, a lesion to connections tends to produce relatively more uniform effects, making smaller networks (which train far more quickly) sufficient for the connection damage simulations.

We now consider a related but distinct possible explanation for the different consequences of unilateral and bilateral damage in humans and other animals. According to this explanation, the information contained in two homotopic areas is more redundant between than within each of the two areas, so if one is damaged, the information lost may still be available on the other side, but if both are damaged, some information lost from one side will also be lost from the other. While we cannot rule out the possibility that such an account applies to lesions in the brain, it does not seem to be the cause of the differential effect of unilateral vs. bilateral damage in our networks.

In our networks, there does not appear to be a force that produces greater redundancy of representation between vs. within a hemisphere. Such a force certainly does not exist in our feedforward network (Figure 6). Because there are no recurrent connections and all input units are connected to all hidden units and all hidden units are connected to all output units, the division of the hidden layer into two pools has no effect whatsoever during training – the architecture is indistinguishable from that of a network with a single hidden layer. Thus, two units that might represent a particular feature are equally likely to be located on the same side of the network as on two different sides of the network. A consequence of this is that removing x units from each side has the same average effect on the fidelity of the representations in the network as removing $2x$ units from just one side.

To demonstrate that this is indeed the case for the feedforward network, we presented each item's name and visual input pattern, then computed the fidelity measure after removing $2x$ units from one side or x units from both sides. We found no reliable difference across ten networks of unilateral vs. bilateral removal across several values of the parameter x ($x=0.1, 0.3$, and 0.5 ; smallest $p=0.689$). We also carried out a similar analysis for two versions of the recurrent network: the network from section 3.2 and the network with 400 hidden units. In both cases, we examined the fidelity of combined patterns across the two hidden layers after the network was tested without damage, simply excluding units either unilaterally ($2x$ units on one

side) or bilaterally (x units on both sides) in the analysis. As with the feedforward networks, for the same three values of x , we found no reliable effect of unilateral vs. bilateral exclusion of units in either case (smaller network: smallest $p=0.512$; larger network: smallest $p=0.558$). These simulations are consistent with the view that there is no difference in redundancy of units within vs. between hemispheres in any of our networks.

Thus, although the greater redundancy between than within hemispheres account may seem in some ways similar to the account offered by our models, the accounts are in fact different. In the models, the removal of connection weights on just one side distorts the overall internal representation less than the removal of connection weights on both sides because the lesion to both sides distorts the representation on both sides, while the lesion on one side distorts the representation on only one side. While signal distortion is an important part of the story, by itself it is incomplete. If damage results in signal distortion, and downstream brain areas collect information from both hemispheres, the intact hemisphere would not necessarily override the noise arising from the damaged hemisphere. The insight from the current modelling work is that unilateral damage distorts the signal on that side, but also results in a reduction in signal magnitude on that side. This magnitude reduction is what allows the intact side to dominate performance.

4. Discussion

Across different species and multiple processing domains, unilateral damage can produce minimal effects while profound impairment is only observed after bilateral damage. Despite the fact that evidence for this difference in humans can be traced back for over half a century (to the famous case of patient HM) and longer in non-human primates (Brown & Schafer, 1888; Klüver & Bucy, 1939), there has been no computational investigation of the factors that underlie this

difference until now (though there has been computational work on hemispheric specialization; e.g., Lambon Ralph et al., 2001; Weems & Reggia, 2004). We implemented a neural network model that simulates this difference and applied it to the domain of semantic memory, a domain in which there is a large body of relevant neuropsychological data. In line with the more general literature noted above, the patient data indicated that unilateral damage to the anterior temporal region generates only mild semantic impairment, whereas substantial semantic deficits follow bilateral ATL damage (as observed in SD and other bilateral diseases). While matching for volume is difficult, the data suggest that unilateral damage is somehow inherently less deleterious than bilateral.

Our computational model exhibited just such a unilateral advantage. For symmetrically represented functions such as receptive comprehension (as measured by word-to-picture matching), unilateral damage only produced mild impairments, whereas the same quantity of simulated bilateral damage generated greater deficits at all levels of damage. For asymmetrically-represented tasks such as picture naming, the simulations exhibited the expected laterality effect (lesions to the left demi-hub generated greater naming impairment than equivalent right demi-hub lesions; see also Lambon Ralph et al., 2001), while still showing an advantage, even for a unilateral lesion affecting the somewhat more dominant hemisphere, compared to equivalent damage spread across both hemispheres.

Part of the reason why performance after bilateral damage can be very poor is that it can affect *all* the units and connections whose function together determines network performance. As a result, as shown in Figure 4, the bilateral damage performance curves continue past the range possible for unilateral damage, to yield very low levels of performance. It is probable that some of the patients with bilateral damage (e.g., those with severe SD) have similarly high amounts of atrophy.

Our work indicates that there are also other important parts of the story. Within ranges

where the model received equal amounts of damage unilaterally and bilaterally, performance after bilateral damage was much worse than after unilateral damage. The restriction of damage to one side of the model provides a considerable advantage over the same amount of damage distributed across both sides. When damage is restricted to one side, the other side can continue to perform much as it did when the network was fully intact, with relatively little disruption from the damaged side. Though it may be profoundly impaired, the damaged side contributes far less than the intact side due to the reduction in the variation in its activation from item to item, thereby allowing the intact side to dominate performance. When damage is applied to both halves of the network, processing is disrupted throughout, leading therefore to significantly greater behavioral impairment.

A mathematical analysis of the effects of damage demonstrated how distortion and signal magnitude interact in each hemisphere to produce different amounts of overall disruption in the unilateral and bilateral damage cases. This analysis combined with an understanding of the distributed representations and connectivity patterns of the network provides a clear picture of the causes underlying the differential effects of focal and diffuse damage.

Our results apply very generally to information processing in a highly interconnected system that is distributed bilaterally across the hemispheres of the brain, and may also be relevant to understanding effects of focal vs. diffuse damage within a hemisphere as well, in cases where a function is unilaterally represented over a brain region of moderate extent. If in that case local connectivity among participating neurons is greater than long-range connectivity, then focal damage within the larger area would be expected to behave similarly to unilateral damage within our networks. More generally, although our model simulated the effects of unilateral vs. bilateral ATL damage on semantic memory, our observations about the differential consequences of focal and diffuse damage are likely to be relevant to similar effects observed in many other domains.

Acknowledgements

We thank Ken Norman for helpful discussion and ideas about the simulation analyses. This work was supported by the National Science Foundation Graduate Research Fellowship under Grant No. DGE-0646086 to A.C.S and by Medical Research Council programme grants to M.A.L.R. (MR/J004146/1).

References

- Adolphs, R., & Tranel, D. (2004). Impaired judgments of sadness but not happiness following bilateral amygdala damage. *Journal of Cognitive Neuroscience*, 16(3), 453-462.
- Binney, R. J., Embleton, K. V., Jefferies, E., Parker, G. J., & Lambon Ralph, M. A. (2010). The ventral and inferolateral aspects of the anterior temporal lobe are crucial in semantic memory: evidence from a novel direct comparison of distortion-corrected fMRI, rTMS, and semantic dementia. *Cereb Cortex*, 20(11), 2728-2738. doi: bhq019 [pii]
10.1093/cercor/bhq019
- Bozeat, S., Lambon Ralph, M. A., Patterson, K., Garrard, P., & Hodges, J. R. (2000). Non-verbal semantic impairment in semantic dementia. *Neuropsychologia*, 38(9), 1207-1215. doi: S0028-3932(00)00034-8 [pii]
- Braitenberg, V., & Schüz, A. (1998). *Cortex: statistics and geometry of neuronal connectivity* NY: Springer.
- Brown, S., & Schafer, E. (1888). An investigation into the functions of the occipital and temporal lobes of the monkey's brain. *Philosophical Transactions of the Royal Society of London*, 179, 303-327.
- D'Esposito, M., Cooney, J. W., Gazzaley, A., Gibbs, S. E., & Postle, B. R. (2006). Is the prefrontal cortex necessary for delay task performance? Evidence from lesion and FMRI data. *J Int Neuropsychol Soc*, 12(2), 248-260. doi: S1355617706060322 [pii]
10.1017/S1355617706060322
- Devlin, J. T., Russell, R. P., Davis, M. H., Price, C. J., Wilson, J., Moss, H. E., . . . Tyler, L. K. (2000). Susceptibility-induced loss of signal: comparing PET and fMRI on a semantic task. *Neuroimage*, 11(6 Pt 1), 589-600. doi: 10.1006/nimg.2000.0595
S1053-8119(00)90595-0 [pii]
- Farah, M. J., & McClelland, J. L. (1991). A computational model of semantic memory impairment: modality specificity and emergent category specificity. *J Exp Psychol Gen*, 120(4), 339-357.
- Gee, D. G., Biswal, B. B., Kelly, C., Stark, D. E., Margulies, D. S., Shehzad, Z., . . . Milham, M. P. (2011). Low frequency fluctuations reveal integrated and segregated processing among the cerebral hemispheres. *Neuroimage*, 54(1), 517-527. doi: S1053-8119(10)00825-6 [pii]
10.1016/j.neuroimage.2010.05.073
- Giovagnoli, A. R., Erbetta, A., Villani, F., & Avanzini, G. (2005). Semantic memory in partial epilepsy: verbal and non-verbal deficits and neuroanatomical relationships. *Neuropsychologia*, 43(10), 1482-1492. doi: S0028-3932(05)00016-3 [pii]
10.1016/j.neuropsychologia.2004.12.010
- Glosser, G., Salvucci, A. E., & Chiaravalloti, N. D. (2003). Naming and recognizing famous faces in temporal lobe epilepsy. *Neurology*, 61(1), 81-86.
- Heffner, H.E., & Heffner, R.S. . (1986). Effect of unilateral and bilateral auditory cortex lesions on the discrimination of vocalizations by Japanese macaques. *J. Neurophysiol.*, 56, 683-701.
- Hermann, B. P., Seidenberg, M., Haltiner, A., & Wyler, A. R. (1995). Relationship of age at onset, chronologic age, and adequacy of preoperative performance to verbal memory change after anterior temporal lobectomy. *Epilepsia*, 36(2), 137-145.

- Hermann, B. P., Wyler, A. R., Somes, G., Dohan, F. C., Jr., Berry, A. D., 3rd, & Clement, L. (1994). Declarative memory following anterior temporal lobectomy in humans. *Behav Neurosci*, 108(1), 3-10.
- Hinton, G. E., & Sejnowski, T. J. (1983). *Learning and relearning in Boltzmann machines*. Cambridge, MA: MIT Press.
- Keidel, J. L., Welbourne, S. R., & Lambon Ralph, M. A. (2010). Solving the paradox of the equipotential and modular brain: A neurocomputational model of stroke vs. slow-growing glioma. *Neuropsychologia*, 48(6), 1716-1724. doi: Doi 10.1016/J.Neuropsychologia.2010.02.019
- Klüver, H., & Bucy, P.C. (1939). Preliminary analysis of functions of the temporal lobes in monkeys. *Archives of Neurology and Psychiatry*, 42(6), 979-1000.
- Lambon Ralph, M. A., Cipolotti, L., Manes, F., & Patterson, K. (2010). Taking both sides: Do unilateral, anterior temporal-lobe lesions disrupt semantic memory? . *Brain*, 133, 3243–3255.
- Lambon Ralph, M. A., Ehsan, S., Baker, G. A., & Rogers, T. T. (2012). Semantic memory is impaired in patients with unilateral anterior temporal lobe resection for temporal lobe epilepsy. *Brain*, 135(Pt 1), 242-258. doi: awr325 [pii] 10.1093/brain/awr325
- Lambon Ralph, M. A., McClelland, J. L., Patterson, K., Galton, C. J., & Hodges, J. R. (2001). No right to speak? The relationship between object naming and semantic impairment: Neuropsychological abstract evidence and a computational model. *Journal of Cognitive Neuroscience*, 13(3), 341-356.
- Lambon Ralph, M. A., Pobric, G., & Jefferies, E. (2009). Conceptual knowledge is underpinned by the temporal pole bilaterally: convergent evidence from rTMS. *Cereb Cortex*, 19(4), 832-838. doi: bhn131 [pii] 10.1093/cercor/bhn131
- Lewis, J. D., Theilmann, R. J., Sereno, M. I., & Townsend, J. (2009). The relation between connection length and degree of connectivity in young adults: a DTI analysis. *Cereb Cortex*, 19(3), 554-562. doi: bhn105 [pii] 10.1093/cercor/bhn105
- Li, H., Matsumoto, K., & Watanabe, H. (1999). Different effects of unilateral and bilateral hippocampal lesions in rats on the performance of radial maze and odor-paired associate tasks. *Brain Res Bull*, 48(1), 113-119. doi: S0361-9230(98)00157-9 [pii]
- Marinkovic, K., Dhond, R. P., Dale, A. M., Glessner, M., Carr, V., & Halgren, E. (2003). Spatiotemporal dynamics of modality-specific and supramodal word processing. *Neuron*, 38(3), 487-497. doi: S0896627303001971 [pii]
- Martin, R. C., Sawrie, S. M., Roth, D. L., Gilliam, F. G., Faught, E., Morawetz, R. B., & Kuzniecky, R. (1998). Individual memory change after anterior temporal lobectomy: a base rate analysis using regression-based outcome methodology. *Epilepsia*, 39(10), 1075-1082.
- Milner, B., & Penfield, W. (1955). The effect of hippocampal lesions on recent memory. *Trans Am Neurol Assoc*(80th Meeting), 42-48.
- Mion, M., Patterson, K., Acosta-Cabronero, J., Pengas, G., Izquierdo-Garcia, D., Hong, Y. T., . . . Nestor, P. J. (2010). What the left and right anterior fusiform gyri tell us about semantic memory. *Brain*, 133(11), 3256-3268. doi: awq272 [pii] 10.1093/brain/awq272
- Nestor, P. J., Fryer, T. D., & Hodges, J. R. (2006). Declarative memory impairments in Alzheimer's disease and semantic dementia. *Neuroimage*, 30(3), 1010-1020. doi: S1053-8119(05)00789-5 [pii] 10.1016/j.neuroimage.2005.10.008

- Noppeney, U., Patterson, K., Tyler, L. K., Moss, H., Stamatakis, E. A., Bright, P., . . . Price, C. J. (2007). Temporal lobe lesions and semantic impairment: a comparison of herpes simplex virus encephalitis and semantic dementia. *Brain*, 130(Pt 4), 1138-1147. doi: awl344 [pii]
10.1093/brain/awl344
- Patterson, K., & Hodges, J. R. (1992). Deterioration of word meaning: implications for reading. *Neuropsychologia*, 30(12), 1025-1040.
- Patterson, K., Nestor, P. J., & Rogers, T. T. (2007). Where do you know what you know? The representation of semantic knowledge in the human brain. *Nat Rev Neurosci*, 8(12), 976-987. doi: nrn2277 [pii]
10.1038/nrn2277
- Pobric, G., Jefferies, E., & Lambon Ralph, M. A. (2007). Anterior temporal lobes mediate semantic representation: mimicking semantic dementia by using rTMS in normal participants. *Proc Natl Acad Sci U S A*, 104(50), 20137-20141. doi: 0707383104 [pii]
10.1073/pnas.0707383104
- Pobric, G., Jefferies, E., & Lambon Ralph, M. A. (2010). Amodal semantic representations depend on both anterior temporal lobes: evidence from repetitive transcranial magnetic stimulation. *Neuropsychologia*, 48(5), 1336-1342. doi: S0028-3932(09)00522-3 [pii]
10.1016/j.neuropsychologia.2009.12.036
- Rogers, T. T., Lambon Ralph, M. A., Garrard, P., Bozeat, S., McClelland, J. L., Hodges, J. R., & Patterson, K. (2004). Structure and deterioration of semantic memory: a neuropsychological and computational investigation. *Psychol Rev*, 111(1), 205-235. doi: 10.1037/0033-295X.111.1.205
2004-10332-012 [pii]
- Rohrer, J. D., Warren, J. D., Modat, M., Ridgway, G. R., Douiri, A., Rossor, M. N., . . . Fox, N. C. (2009). Patterns of cortical thinning in the language variants of frontotemporal lobar degeneration. *Neurology*, 72(18), 1562-1569. doi: 72/18/1562 [pii]
10.1212/WNL.0b013e3181a4124e
- Scoville, W. B., & Milner, B. (2000). Loss of recent memory after bilateral hippocampal lesions. 1957. *J Neuropsychiatry Clin Neurosci*, 12(1), 103-113.
- Seidenberg, M., Hermann, B., Wyler, A. R., Davies, K., Dohan, F. C., Jr., & Leveroni, C. (1998). Neuropsychological outcome following anterior temporal lobectomy in patients with and without the syndrome of mesial temporal lobe epilepsy. *Neuropsychology*, 12(2), 303-316.
- Sharp, D. J., Scott, S. K., & Wise, R. J. (2004). Retrieving meaning after temporal lobe infarction: the role of the basal language area. *Ann Neurol*, 56(6), 836-846. doi: 10.1002/ana.20294
- Snowden, J. S., Goulding, P. J., & Neary, D. . (1989). Semantic dementia: A form of circumscribed cerebral atrophy. *Behav. Neurobio*, 2, 167-182.
- Vandenberghe, R., Price, C., Wise, R., Josephs, O., & Frackowiak, R. S. (1996). Functional anatomy of a common semantic system for words and pictures. *Nature*, 383(6597), 254-256. doi: 10.1038/383254a0
- Visser, M., Embleton, K. V., Jefferies, E., Parker, G. J., & Lambon Ralph, M. A. (2010). The inferior, anterior temporal lobes and semantic memory clarified: novel evidence from distortion-corrected fMRI. *Neuropsychologia*, 48(6), 1689-1696. doi: S0028-3932(10)00070-9 [pii]
10.1016/j.neuropsychologia.2010.02.016

- Visser, M., Jefferies, E., & Lambon Ralph, M. A. (2010). Semantic processing in the anterior temporal lobes: a meta-analysis of the functional neuroimaging literature. *J Cogn Neurosci*, 22(6), 1083-1094. doi: 10.1162/jocn.2009.21309
- Visser, M., & Lambon Ralph, M. A. (2011). Differential contributions of bilateral ventral anterior temporal lobe and left anterior superior temporal gyrus to semantic processes. *J Cogn Neurosci*, 23(10), 3121-3131. doi: 10.1162/jocn_a_00007
- Warden, C.J., Barrera, S.E., & Galt, W. (1942). The effect of unilateral and bilateral frontal lobe extirpation on the behavior of monkeys. *Journal of Comparative Psychology*, 34, 139-171.
- Warrington, E. K. (1975). The selective impairment of semantic memory. *Q J Exp Psychol*, 27(4), 635-657. doi: 10.1080/14640747508400525
- Weems, S. A., & Reggia, J. A. (2004). Hemispheric specialization and independence for word recognition: a comparison of three computational models. *Brain Lang*, 89(3), 554-568. doi: 10.1016/j.bandl.2004.02.001
- Welbourne, S. R., & Lambon Ralph, M. A. (2005). Exploring the impact of plasticity-related recovery after brain damage in a connectionist model of single-word reading. *Cognitive Affective & Behavioral Neuroscience*, 5(1), 77-92.
- Wilkins, A., & Moscovitch, M. (1978). Selective Impairment of Semantic Memory after Temporal Lobectomy. *Neuropsychologia*, 16(1), 73-79.

Figure 1. Example scans of patients with unilateral and bilateral anterior temporal lobe damage. **(a)** An example axial MRI for a patient with semantic impairment resulting from bilateral ATL atrophy (orange arrows). **(b)** A comparable axial slice from a patient following ATL unilateral resection for temporal lobe epilepsy (red arrow). The red region on the lateral view shows the resected area. Adapted from Fig. 1 of Lambon Ralph, Ehsan, Baker, & Rogers (2012).

Figure 2. Patient performance. The proportion of items correct (out of 64) on naming and word-to-picture matching tasks of patients with unilateral left, unilateral right, and bilateral ATL damage. Averages for each patient group are included on the far right of each graph, with error bars indicating ± 1 SEM. The full horizontal line denotes control-participant mean performance, and the dashed line shows 2 SD below the control mean. Adapted from Fig. 2 of Lambon Ralph, Cipolotti, Manes, & Patterson (2010).

Figure 3. Model Architecture. The model consists of two layers, *verbal descriptors* and *visual features*, which serve as input and output, and two hidden layers, *left demi-hub* and *right demi-hub*. Arrows indicate full connectivity from all the units in the projecting to the receiving layer, except where there are percentages written which specify a different degree of connectivity.

Figure 4. Performance of the network on word-to-picture matching and naming with unilateral and bilateral damage. Proportion correct on these tasks is shown for 20 levels of damage in each graph, labelled by the proportion of connections removed from units in both hidden layers or one hidden layer. Deterioration of performance is shown with no recovery throughout the course of damage, one epoch of recovery between each episode of damage, and ten epochs between each episode of damage.

Figure 5. Degradation of cosine structure in unilateral and bilateral networks. The colors in the heatmaps represent the value of the cosines between hidden layer activation vectors after each set of visual features (indexed by rows) and each name (indexed by columns) were presented, for the intact model, and for the model with two levels of unilateral and bilateral damage. Names and visual features are clustered by category. Cosine values range from 0 (blue) to 1 (red).

Figure 6. Feedforward model architecture.

Figure 7. Fidelity of semantic representations as a function of damage. Fidelity was defined as the difference between mean diagonal and mean off-diagonal cosine values. The value of this difference was normalized by the value for the intact network, which therefore had a fidelity value of 1. **(a)** Fidelity in the damaged layer alone (damaged unilateral), the remaining, undamaged layer (intact unilateral); the average of these two (average unilateral); and the average of the fidelity across two bilaterally damaged semantic layers matched for damage to the unilateral cases (average bilateral); **(b)** Fidelity measure applied directly over both hidden layers in the cases of unilateral and bilateral damage in the feedforward network; **(c)** Same as B but in the recurrent architecture; **(d)** Same as C, but after unit instead of connection damage. This network had 400 instead of 64 hidden units, as discussed in the text.

Figure 8. Effect of unilateral and bilateral connection weight removal on the cosine of the net input vectors at the hidden layer in the feedforward network. Analytic approximation (labelled ‘Equations’; grey) and actual values obtained in simulations (black). The cosine is calculated

between the net input vector produced by inputting the name of an item with the cosine of the vector produced by the visual input pattern for the same item. **(a)** Effect of a lesion to weights coming into hidden units on one side on the cosine across the units on that side. **(b)** Effect of a lesion to weights on one side on the product of the magnitudes of the two vectors on that side. **(c)** Effect of a lesion on the cosine of the net input vectors across both hidden layers, either for a unilateral lesion or for a bilateral lesion of a given overall proportion of connection weights. The curves labelled 'equations' are based on the formulas given in the text. The simulation results were obtained by applying the following procedure to each of the 10 feedforward networks and averaging the results obtained for each network: At each lesion level, connections were removed from one or both hidden layers. Then all 48 names were presented and the resulting hidden unit activation pattern was recorded for each, and similarly for all 48 visual patterns. There was no noise added to the net inputs in these simulations, and the bias term was not included in the net input value. The cosines (panels **a** and **c**) between each corresponding name and each corresponding visual pattern, as well as the product of the magnitudes of the members of each corresponding pair (panel **b**) were calculated. For each network the 48 resulting cosines and the 48 resulting magnitudes were averaged.

Figure 1

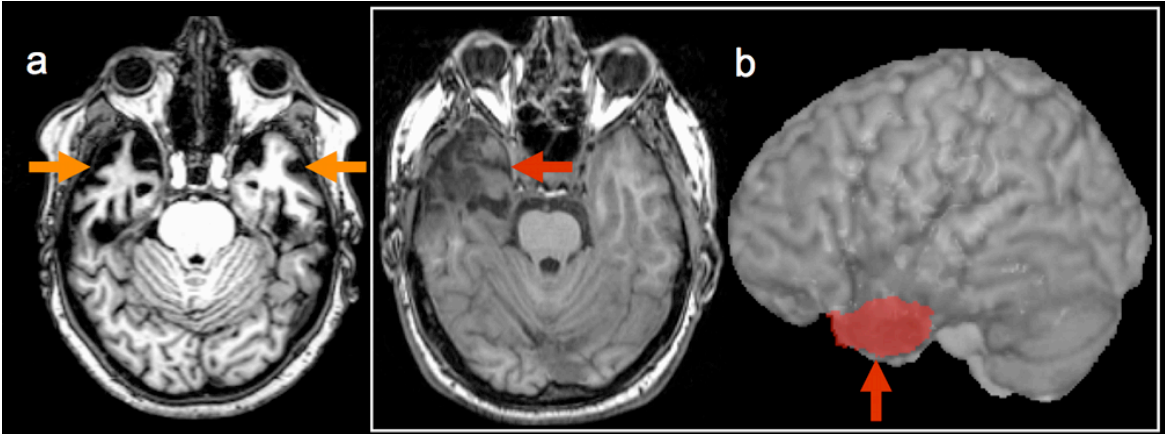


Figure 2

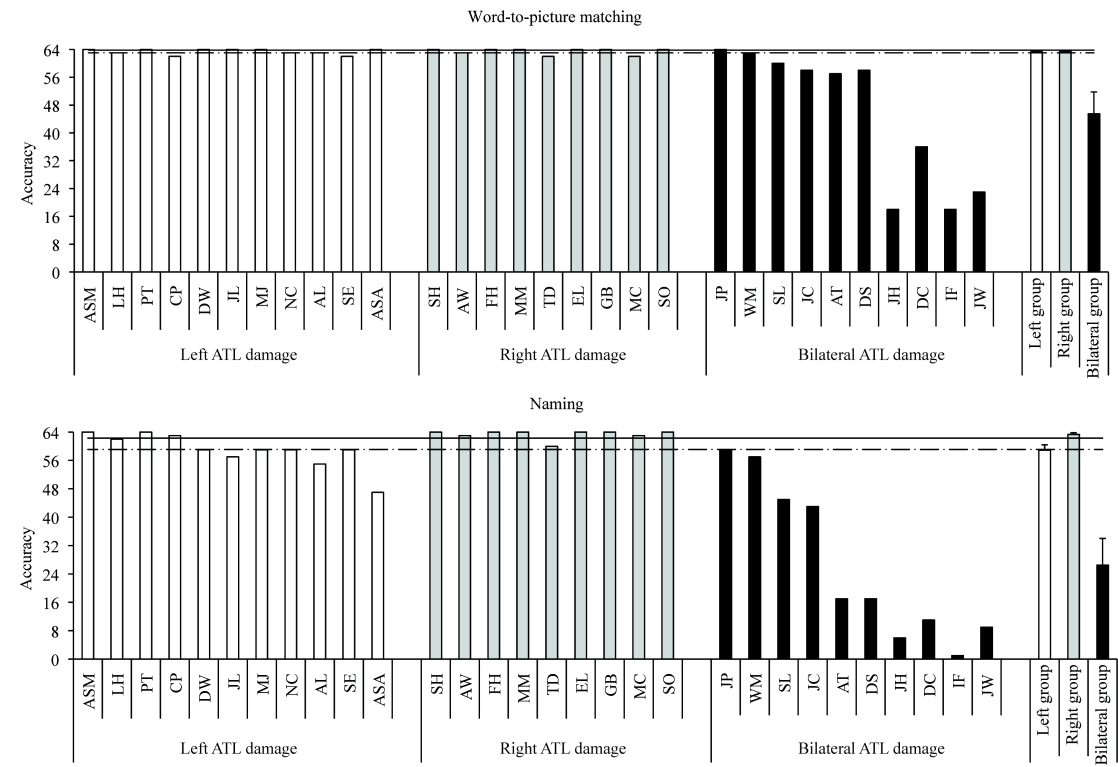


Figure 3

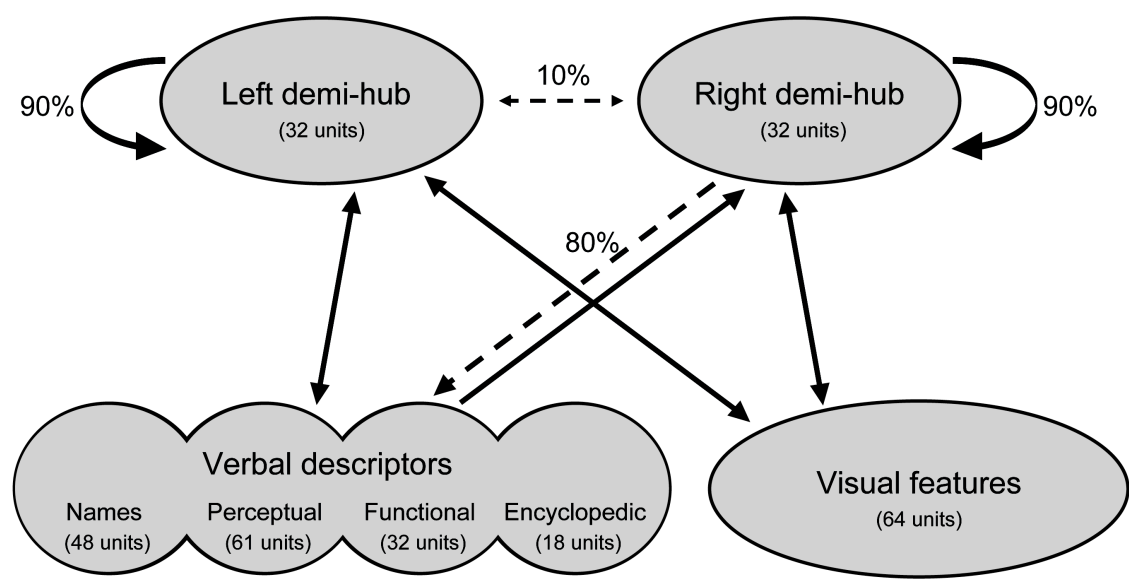


Figure 4

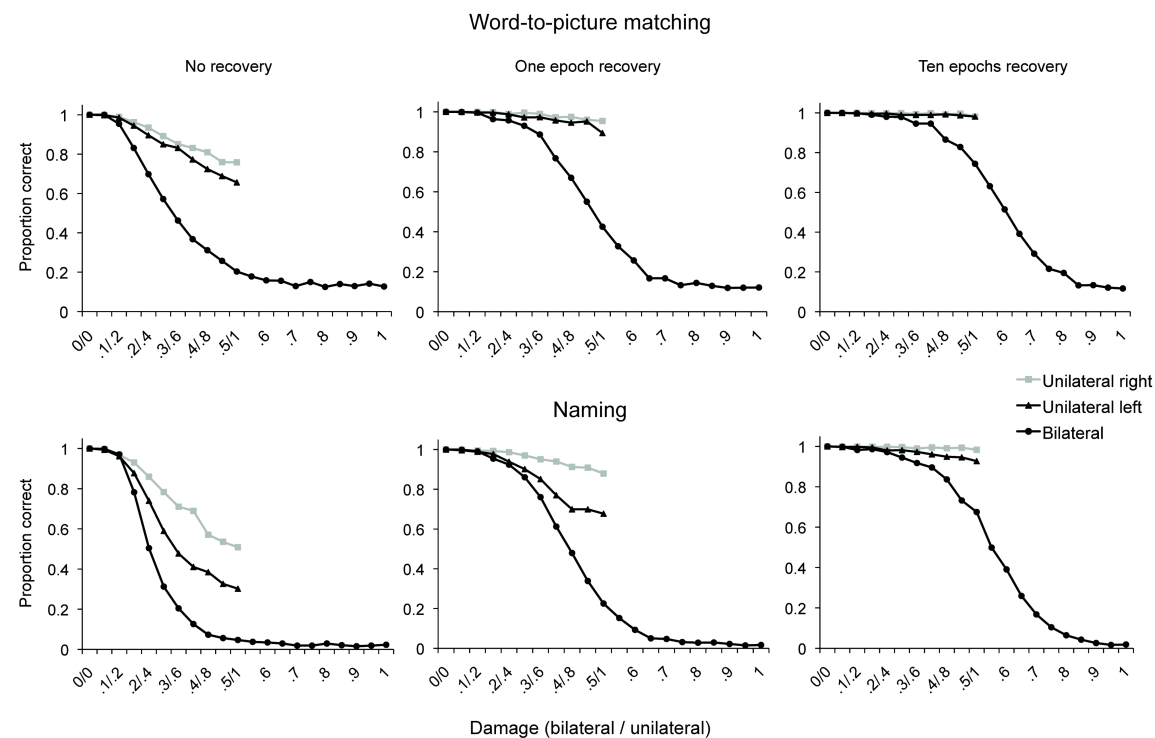


Figure 5

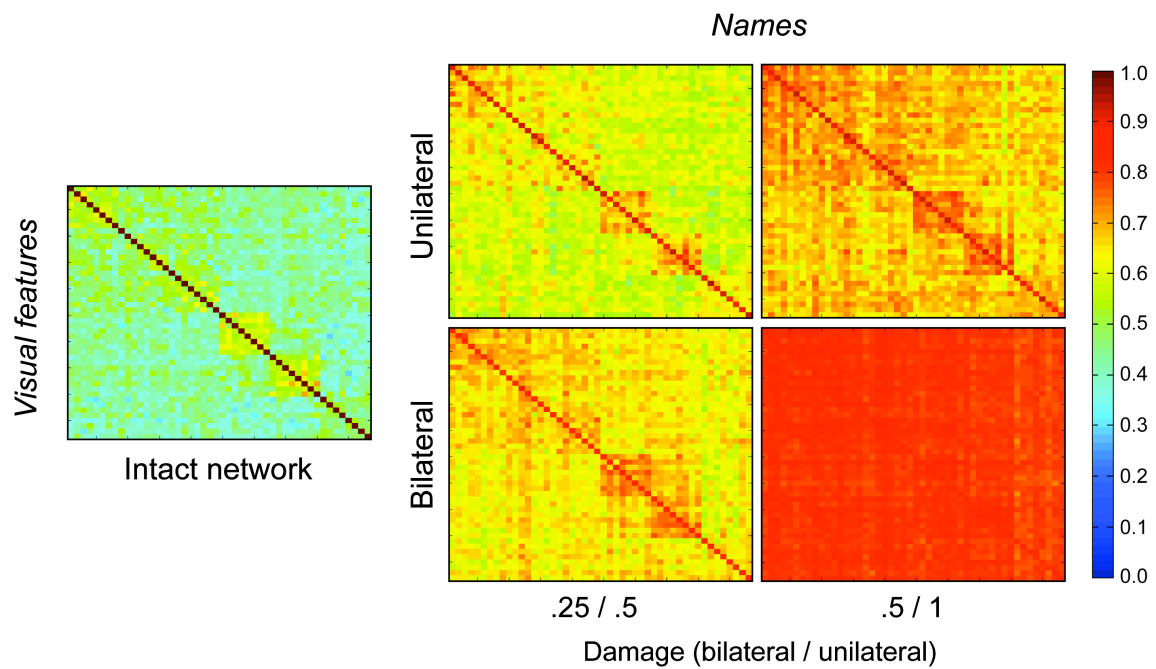


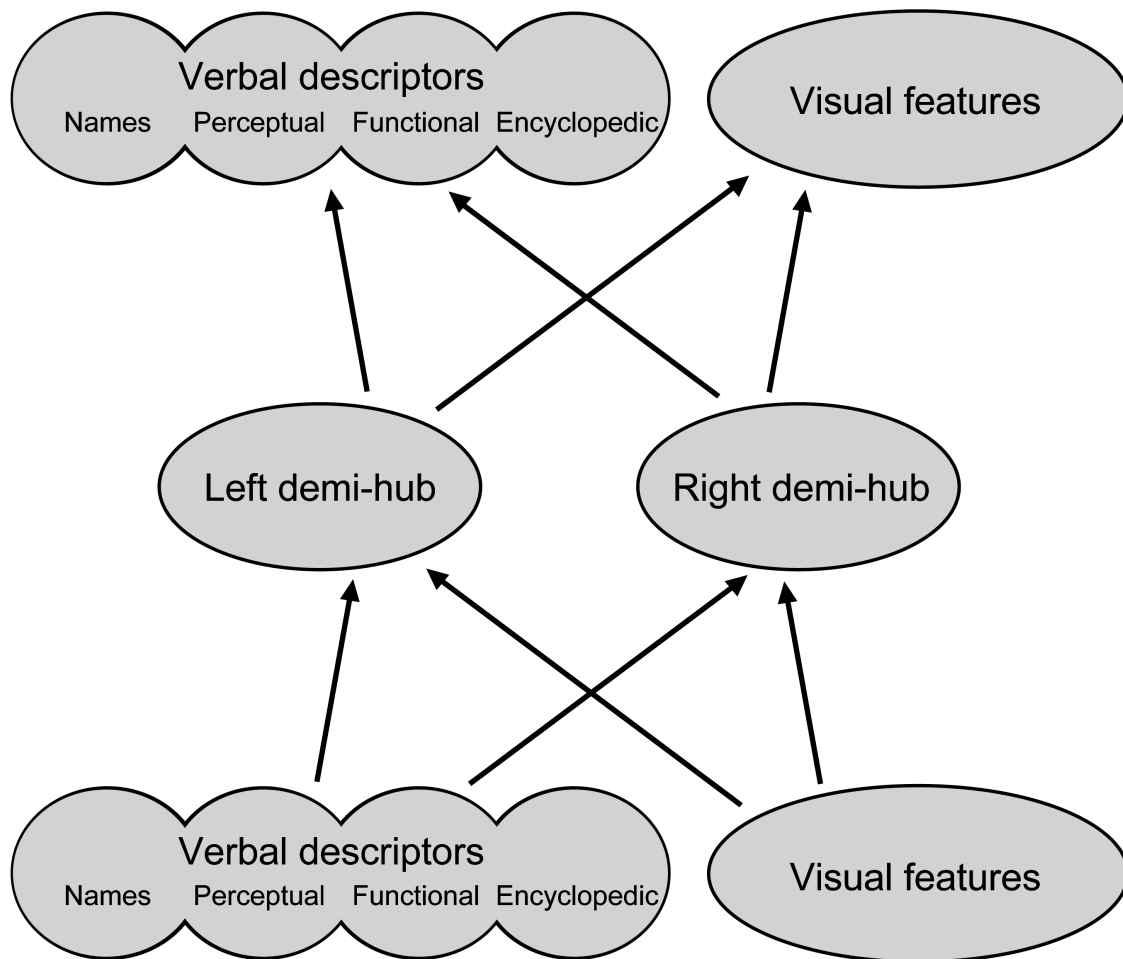
Figure 6

Figure 7

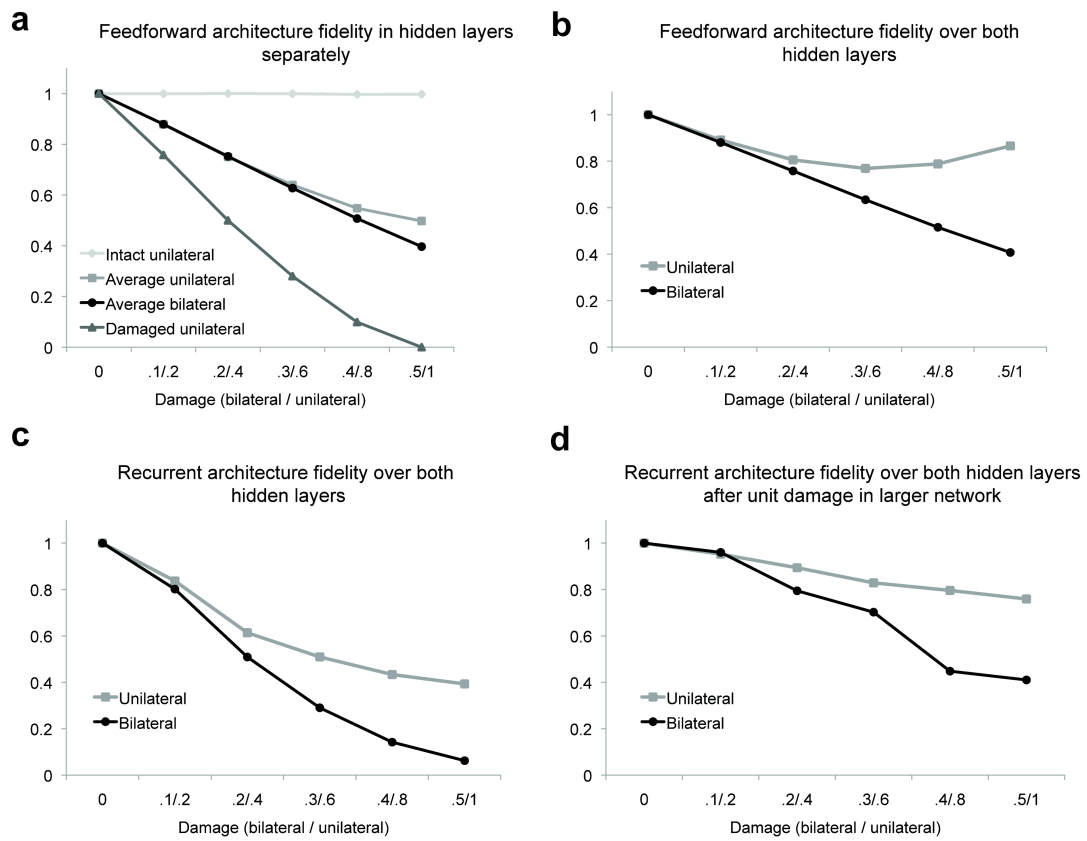


Figure 8

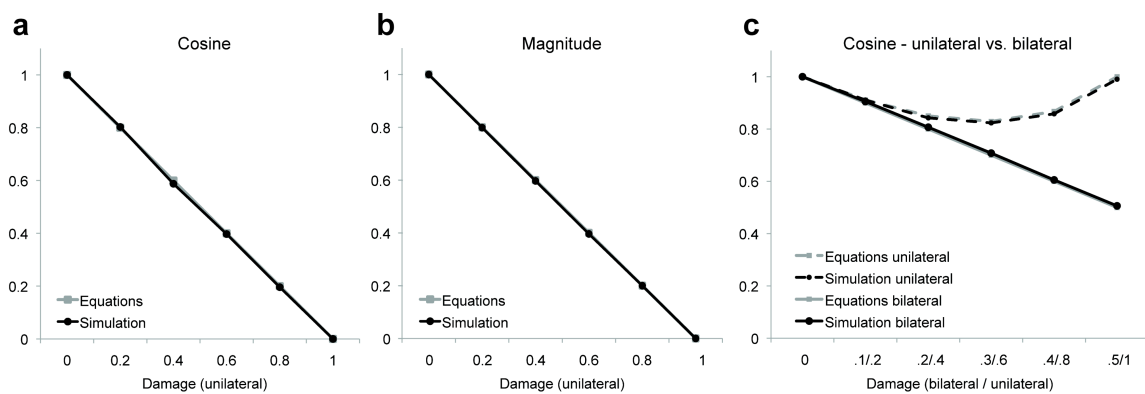


Table 1. Analysis of variance for Figure 4 results

		<i>df</i>	<i>F</i>	<i>p</i>
No recovery				
Word-to-picture matching				
	Damage	10, 90	301.2	<0.0001
	Bilateral vs. mean unilateral	1, 9	264.0	<0.0001
	Damage * unilateral vs. bilateral	10, 90	74.57	<0.0001
	Right vs. left unilateral	1, 9	4.53	0.0622
Naming				
	Damage	10, 90	888.7	<0.0001
	Bilateral vs. mean unilateral	1, 9	277.5	<0.0001
	Damage * unilateral vs. bilateral	10, 90	68.07	<0.0001
	Right vs. left unilateral	1, 9	62.52	<0.0001
One epoch recovery				
Word-to-picture matching				
	Damage	10, 90	405.9	<0.0001
	Bilateral vs. mean unilateral	1, 9	849.1	<0.0001
	Damage * unilateral vs. bilateral	10, 90	190.6	<0.0001
	Right vs. left unilateral	1, 9	16.49	0.003
Naming				
	Damage	10, 90	372.1	<0.0001
	Bilateral vs. mean unilateral	1, 9	692.2	<0.0001
	Damage * unilateral vs. bilateral	10, 90	196.0	<0.0001
	Right vs. left unilateral	1, 9	119.40	<0.0001
Ten epochs recovery				
Word-to-picture matching				
	Damage	10, 90	81.85	<0.0001
	Bilateral vs. mean unilateral	1, 9	257.8	<0.0001
	Damage * unilateral vs. bilateral	10, 90	52.81	<0.0001
	Right vs. left unilateral	1, 9	5.39	0.045
Naming				
	Damage	10, 90	172.3	<0.0001
	Bilateral vs. mean unilateral	1, 9	305.47	<0.0001
	Damage * unilateral vs. bilateral	10, 90	95.56	<0.0001
	Right vs. left unilateral	1, 9	73.45	<0.0001

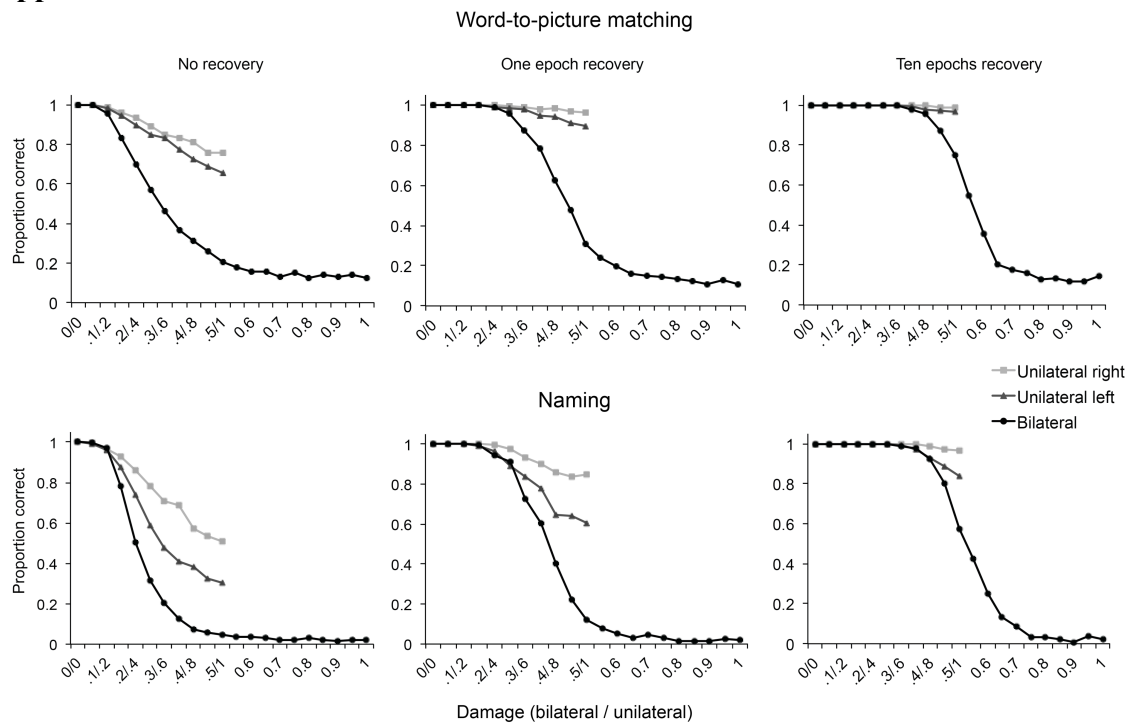
Note. Two factor repeated measures ANOVAs treat network initializations as the random effects factor, assessing overall effect of damage extent, effect of unilateral (averaged across right and left damage) compared to bilateral damage, the interaction between damage extent and the difference between unilateral and bilateral damage, and the difference between right and left unilateral damage. Damage levels include zero to full unilateral damage over 11 increments (as in Figure 4 unilateral damage).

Table 2. Analysis of variance for Figure 7 results

		<i>df</i>	<i>F</i>	<i>p</i>
Feedforward architecture fidelity over both hidden layers				
	Damage	5, 45	1370.5	<0.0001
	Bilateral vs. mean unilateral	1, 9	683.52	<0.0001
	Damage * unilateral vs. bilateral	5, 45	411.84	<0.0001
Recurrent architecture fidelity over both hidden layers				
	Damage	5, 45	861.61	<0.0001
	Bilateral vs. mean unilateral	1, 9	135.36	<0.0001
	Damage * unilateral vs. bilateral	5, 45	64.68	<0.0001
Recurrent architecture fidelity over both hidden layers after unit damage in larger network				
	Damage	5, 45	71.09	<0.0001
	Bilateral vs. mean unilateral	1, 9	87.20	<0.0001
	Damage * unilateral vs. bilateral	5, 45	25.57	<0.0001

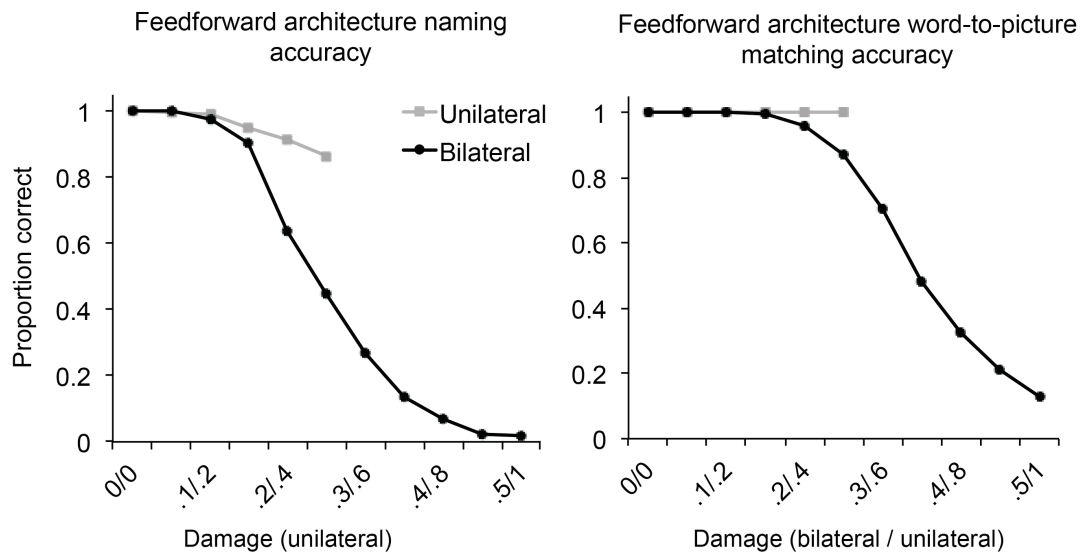
Note. Two factor repeated measures ANOVAs treat network initializations as the random effects factor, assessing overall effect of damage extent, effect of unilateral compared to bilateral damage, and the interaction between damage extent and the difference between unilateral and bilateral damage. Damage levels include zero to full unilateral damage over 6 increments (as in Figure 7).

Appendix 1.



Effects of unilateral and bilateral damage applied all at once (as opposed to progressively). Deterioration of performance is shown with no recovery (equivalent to the Figure 4 results), one epoch of recovery after the specified amount of damage, and ten epochs of recovery after the specified amount of damage.

Appendix 2.



Effects of unilateral and bilateral damage on naming and word-to-picture matching in the feedforward architecture. Performance is shown for the intact network and at ten levels of damage, labeled by the proportion of connections removed from units in both hidden layers or one hidden layer. Deterioration of performance is shown with no recovery allowed.