

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

An associative learning account for retrieval-induced forgetting

Permalink

<https://escholarship.org/uc/item/9rm3942g>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 42(0)

Authors

Hiatt, Laura M.

Jones, Steven J.

Publication Date

2020

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

An associative learning account for retrieval-induced forgetting

Laura M. Hiatt (laura.hiatt@nrl.navy.mil)

US Naval Research Laboratory
Washington, DC USA

Steven J. Jones (scijones@umich.edu)

University of Michigan
Ann Arbor, MI USA

Abstract

Retrieval-induced forgetting (RIF) is a paradigm where repeated study and cue-based retrieval of words impair retrieval of related, but unstudied, words. We present a process model, situated in the ACT-R/E cognitive architecture, that accounts for the RIF task using the architecture's overarching theory of associative learning. In this theory, studying words strengthens their association with their related cues; this, in turn, weakens the association between those cues and any other words they are related to. We show this account fits a recent dataset that explores cueing in the RIF task (Perfect et al., 2004).

Keywords: associative learning; spreading activation; priming; cognitive architecture; retrieval-induced forgetting

Introduction

The *retrieval-induced forgetting* (RIF) paradigm, first introduced by M. C. Anderson, Bjork, and Bjork (1994), is a cognitive task where repeated study and retrieval of words impairs retrieval of related, but unstudied, words. In a typical RIF experimental task, participants learn category-exemplar pairs (i.e., SPORT-rugby, SPORT-tennis, etc.), where each category contains several members. During a practice phase, participants rehearse only some members of some categories. Then, during testing, participants are cued with each category and are asked to recall all members of that category. Using items of unpracticed categories as a baseline, *practiced* items of practiced categories are more likely to be recalled, whereas *unpracticed* items of practiced categories are less likely to be recalled. This forgetting is the hallmark of the RIF effect.

One explanation for this diminished availability is via inhibition (Levy & Anderson, 2002; M. C. Anderson et al., 1994), where the unstudied words of practiced categories are suppressed during retrieval so that they do not interfere with the practiced items. Because of this suppression, the practiced items are more likely to be recalled than the unstudied words, causing the canonical RIF effect.

A second, competing, account of this effect is an associative model, where associations between the categories and words are strengthened with practice (M. C. Anderson & Spellman, 1995). Here, when cued by the category, the stronger association between the category and the practiced items makes the practiced items more available for retrieval when the category is used as a cue. Thus, even though the unstudied word is not necessarily inhibited, the encouragement of the studied word increases its chance of being recalled relative to unstudied words, again causing the RIF effect.

To test between these, M. C. Anderson and Spellman (1995) argue that it is necessary to use an independent retrieval cue at the final test stage, instead of the category itself. This independent retrieval cue is unique to each exemplar and is not practiced during the rehearsal phase. If inhibition is the correct account, forgetting should occur with both the category cue and the independent cue. If associative learning is the correct account, then forgetting should only occur with the category cue, not the independent cue, since that is what is being affected by the rehearsal phase.

To this end, Perfect et al. (2004) ran a version of the RIF experiment that included a second cue, unique to each exemplar, that participants learned along with the categories and exemplars. The cue was not present during the rehearsal phase. Then during testing, participants were either cued with the category, the unique cue, or both. The results showed that, while the typical RIF patterns were exhibited when participants were cued with the category during testing, that pattern did not hold when cued with the unique cue. This supports the associative explanation of the RIF task.

Here, we provide further evidence that the RIF paradigm can be explained via associative learning. We present a cognitive process model of the task, situated in the ACT-R/E architecture (Trafton et al., 2013), that uses associative learning to capture and explain the experimental results. The underlying associative learning mechanisms are part of our overall architectural account, and have been used to capture a large body of other cognitive effects that involve associative learning (e.g., Hiatt & Trafton, 2017; Hiatt, 2017; Hiatt & Trafton, 2015a). A critical feature of our associative learning account is that associations can both strengthen and weaken, depending on the model's exposure to items. We show our model's ability to account for the main effects of the RIF paradigm by modeling Experiment 1 from Perfect et al. (2004).

Experiment

The specific experiment we model is Experiment 1 from Perfect et al. (2004). In it, participants learn a set of exemplars (the *training phase*), rehearse a subset of them (the *rehearsal phase*), perform a distractor task, and then are asked to retrieve the exemplars under several conditions (the *testing phase*). There were 24 exemplars in the study, each assigned on one of six categories, with four exemplars per category. The exemplars and categories are shown in Table 1. Ninety

participants were involved in the study.

SPORT	rugby, tennis, swimming, hockey
COUNTRY	France, Italy, Spain, Greece
PET	hamster, parakeet, gerbil, rabbit
FOOD	cheese, fish, pasta, burger
DRINK	tequila, vodka, wine, lager
HOBBY	cooking, reading, painting, drawing

Table 1: The categories and exemplars (Perfect et al., 2004).

During the training phase, participants viewed exemplars one by one, with two types of cues: their category, and a photograph of a human *face*, which is arbitrarily paired with, and unique to, each exemplar. Participants were instructed to try to relate each item to its category and to the face. Exemplars were presented at a rate of 4 seconds per item, with a 1 second inter-stimulus interval.

During the rehearsal phase, participants practiced two items from four different categories. For the rehearsal, participants saw the category, plus the (unique) first two letters of the exemplar as an *exemplar stem*. Participants were instructed to complete the word stem with the appropriate exemplar. Participants had 4 seconds to do so. The 8 total rehearsal items were repeated three times in random order.

After the rehearsal phase, participants engaged in a 7 minute distractor task of performing visual puzzles. Participants then underwent one of three testing conditions. In the *category condition*, participants were cued with each of the six categories in random order and, for each, had 20 seconds to write down all exemplars of that category that they could remember. Between categories there was a 1 second interval.

In the *face condition*, participants were shown each of the 24 faces in random order and were asked to respond with the associated exemplar. They had 5 seconds to write down their response, with a 1 second interval between faces.

Finally, in the *joint condition*, participants saw both the category and face of each of the exemplars as cues, in random order, and had 5 seconds to write down their response of the exemplar associated with the face. As before, there was a 1 second inter-stimulus interval.

Experimental Results

The data reveal two interesting trends. First, in all three conditions, the practiced items (*RP+* items) had the best recall performance, on average; in contrast, unpracticed items from the practiced categories (*RP-* items) had the lowest. Items from unpracticed categories (*U* items) fell generally in the middle. This first trend represents the canonical RIF effect, where unpracticed items from practiced categories seem to be forgotten (*RP-*), relative to both practiced items in practiced categories (*RP+*) and items in unpracticed categories (*U*).

Second, on average, this effect was very obvious in the category condition, slightly present in the joint condition, and only marginally present in the face condition. Recall from the

discussion in the introduction that participants' responses in the presence of an independent cue are viewed as providing evidence for either the inhibition account or the associative learning account of the RIF task. Here, the face cue provides that independent cue, since it is unique to each exemplar and is not present in the rehearsal condition.

In an inhibitory account, then, one would expect the patterns of responses in the face cue to be similar to those of the category condition, with *RP-* inhibited relative to both *RP+* and *U*. In an associative learning account, however, one would expect it to have a greatly modulated RIF effect, with little difference between the conditions, since the independent cue is not affected by the rehearsal and so remains constant across conditions. These results clearly support the pattern one would expect for the associative learning account. Accordingly, the account provided by Perfect et al. (2004) for these results supports the associative learning view of RIF.

Next, we discuss our models of these results, including how we interpret and capture the two main trends.

Model

We developed a process model of retrieval-induced forgetting situated within the computational cognitive architecture ACT-R/E. As part of associative learning in this architecture, concepts that are thought about in working memory at the same time become associated, and then can prime one another if one is thought about at a later time, facilitating retrieval into working memory. The more that two concepts are thought about together, the stronger their association becomes; however, if the two concepts are thought about separately, their connection weakens. It is this combination of strengthening and weakening that, in large part, we will later rely upon to explain retrieval-induced forgetting.

Model Architecture

ACT-R/E (Trafton et al., 2013) is an embodied version of the ACT-R cognitive architecture (J. R. Anderson, 2007). At a high level, ACT-R/E is an integrated, production-based system. At the core of ACT-R/E is its working memory; the contents of working memory indicate, for example, what the model is looking at, what it is thinking, and its current goal. Working memory is represented as a set of limited-capacity buffers that can contain thoughts or memories.

At any given time, there is a set of *productions* (if-then rules) that may fire because their preconditions are satisfied by the current contents of working memory. From this set, the production with the highest predicted usefulness is selected to fire. The fired production can either change the model's internal state (e.g., by adding something to working memory) or its physical one (e.g., by pressing a key on a keyboard). In our discussion, we abstract over these productions and instead describe processes at a higher level (i.e., we say that we look at an object, instead of discussing the 3-4 productions that must fire to achieve that).

In addition to containing symbolic information (i.e., factual information), memories have activation values that determine

how easy they are to remember at any given time. When a request is made for a memory to be retrieved into working memory, the memory that matches the request and has the highest activation is the one that is added into working memory. Memories with higher activations are added to working memory more quickly; memories with very low activations may not be able to be retrieved at all. Activation has three components, activation strengthening, spreading activation, and activation noise, that together have shown to be an excellent predictor of human declarative memory (J. R. Anderson, Bothell, Lebiere, & Matessa, 1998; Schneider & Anderson, 2011; Thomson, Harrison, Trafton, & Hiatt, 2017). Activation noise is an instantaneous random component that models the noise of the human brain. Activation strengthening¹ is learned over time and is based on the frequency and recency with which a memory has been in working memory in the past. It is designed to represent the activation of a memory over long periods of time.

Spreading activation, in turn, is a short-term activation boost that is meant to capture a memory’s relevance to the current situation. Spreading activation is based on *associations* between different concepts in memory. Memories become associated when they are in working memory at the same time; associations are not created ahead of time. Associations are also directed, but may exist between concepts in both directions. Once established, an association from memory *j* to memory *i* has a strength value that affects the degree to which *j* spreads activation to *i*, and intuitively reflects the probability that memory *i* is expected to be relevant while thinking of memory *j*. This allows associative learning to capture correspondences between memories that are expected to be relevant at the same time, as well as memories that are semantically related (such as an object and its corresponding color and shape).

Association strengths are calculated in a Bayesian-like way, and are a non-standard adaptation of ACT-R’s Bayesian-based associative mechanisms. Because of their Bayesian-like calculation, association strengths increase when concepts are thought about together, but can also weaken between two concepts if they are thought about separately instead of together. We use this adaptation to account for the large numbers of associations and objects needed by the experiment we consider here, which ACT-R’s original formulation cannot do. These equations have been successful in modeling associations and priming in other work across a variety of domains (e.g., Harrison & Trafton, 2010; Trafton et al., 2013; Hiatt & Trafton, 2015a, among others). The equations are included and described in (Hiatt & Trafton, 2017).

Spreading activation sources from the contents of working memory, and distributes activation along associations leading from those contents to concepts in long-term memory; in other words, the contents of working memory serve as cues for retrieval and *prime* their connected concepts. Generally,

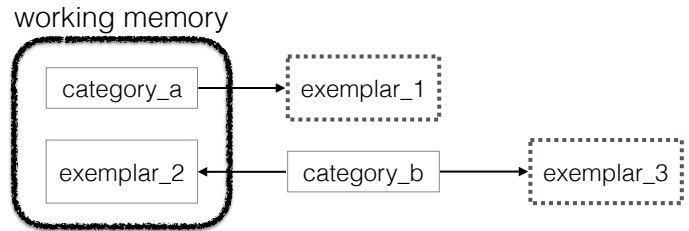


Figure 1: Spreading activation in ACT-R/E. Working memory contains two items; items with spreading activation are designated with dotted outlines. Typically, activation spreads one “hop” out from working memory (*concept_a* to *exemplar_1*); if certain associative patterns appear, however, association can also spread two “hops” from working memory (*exemplar_2* to *exemplar_3*).

activation spreads to a depth of one, and primes only concepts with a direct connection to those in working memory.

In addition, ACT-R/E’s theory of associative learning also allows spreading activation to travel along two associations when a specific associative pattern is present. If there is a concept *k* that is in working memory, and there exist two other concepts *i* and *j* such that *j* primes *i* and *j* primes *k*, spreading activation flows through *j* and primes *i*. We refer to this special case of spreading activation as “two-hop” priming. The two cases of how activation spreads in ACT-R/E are shown in Figure 1. See Hiatt and Trafton (2017) for more information on ACT-R/E’s learning mechanisms.²

ACT-R/E models interact with the world using ACT-R/E’s built-in functionality for doing so. Models can view visual items on a virtual monitor, and can act on the world by pushing keys on a virtual keyboard and clicking a virtual mouse. ACT-R/E models are also inherently tied to physical embodiment (i.e., executing models on a robot), but we do not use that functionality in this paper.

Model Details

The ACT-R/E model for this experiment includes four basic steps consistent with prior models used to capture human performance on a task that requires the use of associative memory: look at experimental stimuli, add experimental stimuli to memory, retrieve other knowledge necessary to perform the task, and respond as appropriate (Hiatt & Trafton, 2017; Hiatt, 2017; Hiatt & Trafton, 2015b, 2015a). Each of these steps is kept as simple as possible throughout the model. As with these other models, properties of associative memory are not modeled in isolation. The model starts out with the task knowledge and productions necessary to complete the tasks. It also assumes moderate prior exposure to the category and exemplar names, since the participants would have encountered them frequently in their daily lives. It does not

¹Activation strengthening is also referred to as base-level activation in other literature describing ACT-R and ACT-R/E.

²Note that in this manuscript, the phrase “two-hop” isn’t used; instead, the two steps are referred to separately as source activation and spreading activation.

include, however, any prior exposure to the face photographs, since it is unlikely participants would have seen those specific photographs before. There are no initial associations; all are learned during the experiment.

The model is shown the stimuli as the participants are, and preserves correspondence with the human timing information. The model “looks at” the stimuli as participants do via its virtual monitor; it digitally logs its responses in lieu of participants’ handwritten responses. As part of the architecture, ACT-R/E encodes each face as a unique token concept.

During the training phase, when the category, exemplar and face are shown at the same time, the model looks at the items randomly and repeatedly for the entire 4 seconds that they are on the screen. Each time the model looks at an item, it encodes the item and adds it to working memory. While this encoding takes place, the model proactively moves on to look at another item. Consequently, items viewed sequentially become associated with one another since they appear in working memory at the same time: one at the end of the encoding process, and the next as it is being looked at and queued up for encoding. Note that, although on average the associations are equal between the three items, the model is looking around randomly and so they are not always equal. At the end of training, all of the RP+, RP- and U exemplars have roughly identical associations, with moderate associative strengths to and from from their categories and faces. Their activation strengthening is also affected by ordering effects: earlier exemplars have lower activation strengthening than those seen near the end of training.

During the rehearsal phase, the model first looks for, and retrieves/encodes, both the category and the prompt. Then, while thinking of both the category and the prompt, the model retrieves the corresponding item. Because our overarching architecture does not have a theoretical way of modeling lexical-based prompts, we simply assume that, given the two-letter prompt, participants always retrieve the correct item. Since the category is in working memory when the exemplar is retrieved, the association from the category to the exemplar is strengthened each time it is rehearsed. The activation strengthening of the exemplar is increased, as well.

During the experimental puzzle task, the model performs a task with productions and concepts unrelated to those of the experiment’s categories, faces and exemplars; it performs no additional rehearsal of the cues/exemplars during this time.

After the distractor task, the model has clear patterns of associations and activation for the RP+, RP- and U groups of exemplars. The RP+ exemplars, which were rehearsed with the corresponding category and prompt, have both higher activation strengthening and stronger associations leading to them from their parent category. The associations from the faces to the RP+ exemplars do not change.

The RP- exemplars have activation strengthening that is slightly lowered, since time has passed without them appearing in working memory. In addition, while the associations from the faces to the RP- exemplars have also not changed,

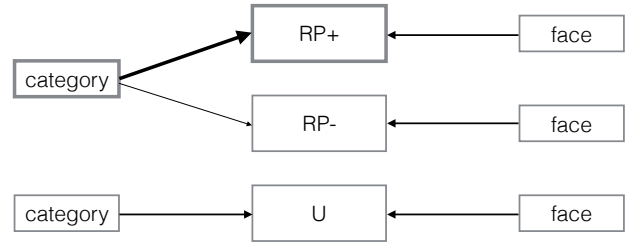


Figure 2: An abstract representation of the model’s memory and associations after the distractor task. Concepts’ outline thickness indicates their activation strengthening; the thickness of the arrows between concepts indicates the strength of their association. For example, RP+ exemplars have the highest activation strengthening, as well as the strongest association from their category; RP- exemplars have the weakest association from their category. Note that, for clarity, not all associations are depicted here (such as links between faces and categories).

the association from the categories to the RP- exemplars have weakened. This is because the RP- categories have been in working memory during the rehearsal, but the RP- exemplars did not appear with them.

The U exemplars also have a slightly lower activation strengthening, since they have not been in working memory for a while. Their associations, however, are unchanged, since the categories, faces and exemplars have not been in working memory. The associations at this point, as well as activation strengthening, are abstractly shown in Figure 2.

During testing, in the *category condition*, the model first looks at and adds the category to working memory. Using the category as a cue, it then tries to sequentially retrieve the four distinct exemplars in memory with the highest total activation. It reports all exemplars it retrieves without further checking their membership in the provided category.

For the *face condition*, the model looks at and adds the face to working memory. While thinking of the face and using it as a cue, it then retrieves an exemplar from working memory and responds with it. The same process occurs during the *joint condition*, except with both the face and category in working memory during the retrieval.

Model Results

We used our model to simulate data from 100 participants per condition (300 participants total) in order to allow our results to better converge on the model’s true predictions. Each model “participant”, regardless of condition, saw the stimuli in a different random order, as well as had different, randomly-chosen RP+ exemplars. The model had the same parameters for each condition. The activation strengthening decay parameter (also called the base level learning parameter) was set to 0.3 instead of its default of 0.5. The activation noise parameter was set to 0.2 instead of its default of 0. The associative learning rate was set to 3.5, representing a mild

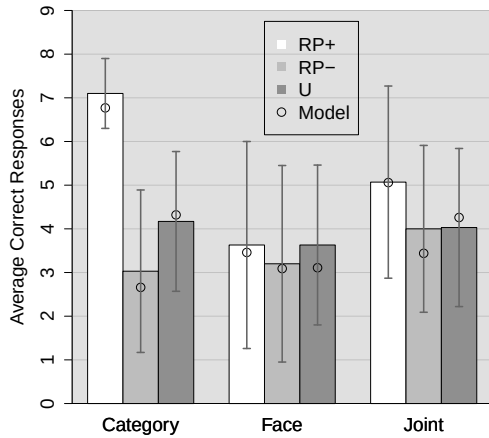


Figure 3: Results of ACT-R/E model. Bars show the experimental data, and error bars show one standard deviation above and below the mean for the experimental data. Dots indicate the model data.

rate of learning. There is no default value for this parameter. All other parameters were at their default values. For more information about how these parameters mathematically relate to the model, please refer to (Hiatt & Trafton, 2017).

The quantitative model results are shown in Figure 3. The y axis indicates the average number of exemplars correctly recalled, out of 8 possible, for each exemplar category. The error bars show one standard deviation above and below the mean for the experimental data. Overall, as the graph shows, the model does an excellent job of capturing the human data, with $R^2 = 0.94$.

During the *category condition*, recall that the model is asked to respond with all exemplars of the category, not just one. Because of the strong association from categories to RP+ exemplars and their higher activation strengthening, the model is almost always able to retrieve a correct RP+ exemplar initially. Then, once the model has retrieved an exemplar, that exemplar is in working memory until the next exemplar is retrieved. That exemplar further cues other exemplars in its category via the two-hop priming described above (see the bottom of Figure 1).

This means that, as the model gets more exemplars correct, it gains *momentum* and becomes slightly more likely to retrieve further correct exemplars. With this initial momentum in place, the model generally retrieves the rest of the RP+ exemplars correctly as well. For the RP- condition exemplars, because those exemplars receive less activation from their associated categories, even with the model’s momentum the model retrieves an exemplar from a different category fairly often, such as one with a higher activation strengthening from ordering effects. The U condition exemplars, which receive a medium amount of priming, lie in the middle, of-

ten being correctly retrieved, but sometimes not. This very closely matches the experimental data.

For the *face condition*, the only difference between the three conditions is in the RP+ exemplar’s higher levels of activation strengthening. All exemplars receive a small amount of activation when cued by their associated face. The RP+ exemplars, therefore, are slightly more likely to have the highest activation overall, and thus be retrieved; the other conditions are equally likely to have correct responses. This is similar to the experimental data, where all three conditions have similar correct responses, on average.

During the *joint condition*, the model behaves similarly to the category condition, with three main differences: (1) it has only one response, not four; (2) it does not have any of the previously discussed “momentum” since the model retrieves only one exemplar for each prompt; and (3) all exemplars receive a small boost in spreading activation from the displayed face concept. Thus, it still performs the best on RP+ exemplars, the worst on RP- exemplars, and in the middle on U exemplars, but with slightly higher correct response rates. This follows the overall pattern of the experimental data, where the RP+ condition had the highest response average, the RP- the lowest, and the U condition is only slightly higher than RP-.

Model Discussion

There are several features of associative learning in ACT-R/E that allow it to capture the experimental data. The first is that associations start out weak and strengthen over time. This explains why, for example, in the face condition, people do not respond perfectly despite the unique cue: they have not been exposed to those two concepts, together, enough times to strengthen the association enough.

A second feature of ACT-R/E that allows it to explain the results is its “two-hop” priming. Recall that this priming occurs during the category condition, where participants recall multiple exemplars for each category cue and, via two-hop priming, build up momentum in their responses. This momentum is what allows ACT-R/E to perform so much better in the category condition than it does in the face and joint conditions: the momentum allows it to correctly retrieve more exemplars than it would otherwise.

Perhaps the most important feature of our associative account with respect to the RIF task is that associative strengths weaken when the underlying concepts are thought about separately. This leads to a weaker association between categories and RP- exemplars after the rehearsal phase, because the categories are thought about in conjunction with the RP+ exemplars. This weakening allows the model to capture the “forgetting” part of retrieval-induced forgetting, where the recall of RP- items is weakened after rehearsal of RP+ items.

General Discussion

We have presented here a process model of Experiment 1 from Perfect et al. (2004) that uses an existing theory of associative learning to explain results of the retrieval-induced forgetting task. This both provides further evidence for our

associative learning theory, as well as supports the thought that RIF is tied to associative learning, as is suggested by the study we model. One conclusion of Perfect et al. (2004), however, that our model does disagree with is whether the two cues (category and face) both refer to the same exemplar representation in memory. The authors of the paper argue that, since RP+ items weren't clearly more available when cued by the face, the rehearsal phase must not have made them more available in memory (via, in our terminology, activation strengthening). Our model disagrees with this conclusion of the experimental data. As our results show, both cues can refer to the same exemplar representation in memory, which does have a higher activation strengthening than the other exemplars. Due to, in part, the large length of time of the distractor task, as well as ordering effects, however, the relative magnitude of this activation strengthening in the face condition is small.

While our results more closely align with the view that the RIF task is driven by associative mechanisms, our model includes some dynamics that can be construed as inhibitory. To clarify, during the task, both the activation strengthening and spreading activation of RP- items lower, with respect to RP+ items, because of the rehearsal phase. Regarding activation strengthening, RP- item's values decay because they are not thought about as recently. However, this does not support the inhibitory account of RIF because the same decay also occurs for U items. For spreading activation, RP- items's association with the category weakens because the category is thought about without the RP- items. And, as we have stated above, a key feature of our model is that if items are seen separately, their association is weakened (Melton & Irwin, 1940). The lessening of activation of RP- items then leaves them more open to interference from other exemplars, leading to their decreased retrieval relative to U items.

So, although the RP- items were not inhibited in memory in our account in the canonical sense, their activation included inhibitory dynamics. We find that the critical distinction between the inhibitory and associative accounts both in the original work and in our model is not simply whether or not there is a decrease in activation for RP- items during competitive retrieval, but also if the mechanism for that decrease differentially lowers the activation of RP- items more than U items. As we have stated, our model captures this. This connects our results with wider thoughts in the cognitive science literature which suggest that, in the debate between decay and interference as the principle driver of memory, *both* decay and interference play a key role (Altmann & Schunn, 2012). By providing an implemented model of how these mechanisms affect behavior and memory as the RIF task unfolds, we are able to relate together inhibition, interference, decay and associations in a unified fashion to better understand retrieval-induced forgetting. Still, in this work we have modeled only one of many studies on RIF; the next step in this work is to apply our associative learning model to a broader range of RIF experiments.

References

- Altmann, E., & Schunn, C. (2012). Decay versus interference: A new look at an old interaction. *Psychological Science*, 23(11), 1435–1437.
- Anderson, J. R. (2007). *How can the human mind occur in the physical universe?* Oxford University Press.
- Anderson, J. R., Bothell, D., Lebiere, C., & Matessa, M. (1998). An integrated theory of list memory. *Journal of Memory and Language*, 38(4), 341–380.
- Anderson, M. C., Bjork, R. A., & Bjork, E. L. (1994). Remembering can cause forgetting: Retrieval dynamics in long-term memory. *Journal of Experimental Psychology: Learning, Memory and Cognition*, 20, 1063–1087.
- Anderson, M. C., & Spellman, B. A. (1995). On the status of inhibitory mechanisms in cognition: Memory retrieval as a model case. *Psychological Review*, 102, 68–100.
- Harrison, A. M., & Trafton, J. G. (2010). Cognition for action: an architectural account for “grounded interaction”. In *Proceedings of the Annual Meeting of the Cognitive Science Society*.
- Hiatt, L. M. (2017). A priming model of category-based feature inference. In *Proceedings of the Annual Meeting of the Cognitive Science Society*.
- Hiatt, L. M., & Trafton, J. G. (2015a). An activation-based model of routine sequence errors. In *Proceedings of the International Conference on Cognitive Modeling*.
- Hiatt, L. M., & Trafton, J. G. (2015b). A computational model of mind wandering. In *Proceedings of the Annual Meeting of the Cognitive Science Society*.
- Hiatt, L. M., & Trafton, J. G. (2017). Familiarity, priming and perception in similarity judgments. *Cognitive Science*, 41(6), 1450–1484.
- Levy, B. J., & Anderson, M. C. (2002). Inhibitory processes and the control of memory retrieval. *Trends in cognitive sciences*, 6(7), 299–305.
- Melton, A. W., & Irwin, J. M. (1940). The influence of degree of interpolated learning on retroactive inhibition and the overt transfer of specific responses. *American Journal of Psychology*, 53, 173–203.
- Perfect, T. J., Stark, L.-J., Tree, J. J., Moulin, C. J., Ahmed, L., & Hutter, R. (2004). Transfer appropriate forgetting: The cue-dependent nature of retrieval-induced forgetting. *Journal of Memory and Language*, 51(3), 399 – 417.
- Schneider, D. W., & Anderson, J. R. (2011). A memory-based model of hick's law. *Cognitive Psychology*, 62(3), 193–222.
- Thomson, R., Harrison, A. M., Trafton, J. G., & Hiatt, L. M. (2017). An account of interference in associative memory: Learning the fan effect. *Topics in Cognitive Science*, 9(1), 69–82.
- Trafton, J. G., Hiatt, L. M., Harrison, A. M., Tamborello, II, F., Khemlani, S. S., & Schultz, A. C. (2013). ACT-R/E: An embodied cognitive architecture for human-robot interaction. *Journal of Human-Robot Interaction*, 2(1), 30–55.