

**No Position-Specific Interference from Prior Lists in Cued Recognition:
A Challenge for Position Coding (and Other) Theories of Serial Memory**

Gordon D. Logan,¹ Gregory E. Cox,² Simon D. Lilburn,¹ Jana E. Ulrich¹

¹Department of Psychology, Vanderbilt University, Nashville TN, 37240, USA

²Department of Psychology, University at Albany, State University of New York, Albany NY,
12222, USA

Corresponding author:

Gordon D. Logan

Email: gordon.logan@vanderbilt.edu

Cognitive Psychology, in press.

Acknowledgments

This research was supported by National Science Foundation grant BCS 2147017 to GDL. We are grateful to Klaus Oberauer, Jeremy Caplan, and two anonymous reviewers for detailed, constructive criticism that greatly improved the article.

Abstract

Position-specific intrusions of items from prior lists are rare but important phenomena that distinguish broad classes of theory in serial memory. They are uniquely predicted by position coding theories, which assume items on all lists are associated with the same set of codes representing their positions. Activating a position code activates items associated with it in current and prior lists in proportion to their distance from the activated position. Thus, prior list intrusions are most likely to come from the coded position. Alternative “item dependent” theories based on associations between items and contexts built from items have difficulty accounting for the position specificity of prior list intrusions. We tested the position coding account with a position-cued recognition task designed to produce prior list interference. Cuing a position should activate a position code, which should activate items in nearby positions in the current and prior lists. We presented lures from the prior list to test for position-specific activation in response time and error rate; lures from nearby positions should interfere more. We found no evidence for such interference in 10 experiments, falsifying the position coding prediction. We ran two serial recall experiments with the same materials and found position-specific prior list intrusions. These results challenge all theories of serial memory: Position coding theories can explain the prior list intrusions in serial recall and but not the absence of prior list interference in cued recognition. Item dependent theories can explain the absence of prior list interference in cued recognition but cannot explain the occurrence of prior list intrusions in serial recall.

Word count: 257

Keywords: serial memory, position coding, prior list intrusions, prior list interference

Introduction

The problem of serial order has been a central topic in psychology and neuroscience for nearly 150 years (Ebbinghaus, 1885; Ladd & Woodworth, 1911; Lashley, 1951). It is important practically because it is ubiquitous in daily life, addressing how we perceive structure in the world, how we structure our actions in time and space, and how we structure our memories of those percepts and actions. It is challenging theoretically. The 150 years were filled with controversy, pitting *item-dependent* theories that explain serial order in terms of associations between the elements of the structure (Ebbinghaus, 1885; Ebenholtz, 1963; Hull, 1932, 1934) against *item-independent* theories that explain order in terms of associations between the elements and a separate set of codes that represent temporal or spatial positions. (Ladd & Woodworth, 1911; Tolman, 1948; Young, 1961). For the last 25 years, item-independent *position coding* theories have dominated research on serial memory, following an influential paper by Henson et al. (1996), who showed that item-dependent theories based on simple chains of associations between adjacent elements could not explain how people recover from errors, respond to manipulations of phonological similarity, produce transpositions to earlier list positions, or produce position-specific intrusions from previous lists. Their findings inspired many researchers to develop theories that implement position coding in various ways (Anderson & Matessa, 1997; Brown et al., 2000, 2007; Burgess & Hitch, 1999; Farrell, 2012; Hartley et al., 2016; Henson, 1998; Lewandowsky & Farrell, 2008; Oberauer et al., 2012). Only a few developed item-dependent theories (Botvinick & Plaut, 2006; Dennis, 2009; Logan, 2021; Solway et al., 2012; also see Lewandowsky & Murdock, 1989; Murdock, 1995).

Recent investigations have shown that item-dependent theories can account for three of the four phenomena that are incompatible with simple chaining theories, by assuming compound retrieval cues and remote associations (Lewandowsky & Li, 1994; Murdock, 1995; Solway et al., 2012) or associations between items and contexts made of fading traces of past items (Logan, 2021). These more elaborate theories can explain recovery from errors (Lewandowsky & Li, 1994; Logan, 2018, 2021), phonological confusability effects (Osth & Hurlstone, 2023; also see Logan, 2018), and transitions to earlier list positions (Logan, 2021; Logan & Cox, 2023; Solway et al., 2012), but cannot explain position-specific intrusions from prior lists (Osth & Hurlstone, 2023; but see Caplan et al., 2022; Dennis, 2009). Thus, position coding theories uniquely

explain position-specific prior list intrusions (Conrad, 1959; Henson, 1998; Melton & Von Lackum, 1941; Osth & Dennis, 2015).

This article reports a critical test of the position coding explanation of position-specific prior list intrusions, using a *cued recognition* task to elicit *position-specific prior list interference*. Subjects were given lists of six random letters to remember followed by a probe display containing a letter and a position cue. They were asked to decide whether the probe letter occurred in the cued position in the memory list (Logan et al., 2021), and *lures* (probe letters that required a “no” response) were sampled from the prior list and from uncued positions within the current list. For example, given list ABCDEF and prior list QRSTUV, ##C### is a *matching* probe that requires a “yes” response, ##S### is a *prior-list lure* that requires a “no” response, and ##B### is a *within-list lure* that requires a “no” response. We show that position coding theories predict longer response time (RT) and higher error rates for prior list lures the closer they are to the cued position—position-specific prior list interference.

This prediction follows directly from the fundamental assumptions of the position coding account of position-specific prior list intrusions: Items in the current list and the prior list are associated with the same position codes. The associations with items in the prior list are weaker. Items are retrieved by activating position codes and reporting what is associated with them. A position code activates the items on both lists in proportion to their strength of association. Items from the current list are activated more than items from the prior list. Under these conditions, retrieving and reporting an item from the prior list is a prior list intrusion. If it is in the right position in the wrong list, it is position-specific (e.g., Henson, 1998).

The cued recognition task establishes the conditions necessary to produce position-specific prior list intrusions and tests their ability to produce position-specific prior list interference. Cued recognition requires focusing on the cued position, which should activate a position code. The position code should activate items associated with it on the current and prior lists in proportion to their distance from the cued position. (in the example above, C and S would be activated more than B and Q). Under these conditions, prior-list lures should match the activated memory items, providing evidence for a “yes” response instead of the required “no” response, which should increase RT and error rate in proportion to the proximity of the lure to the cued position (Logan et al., 2021)—position-specific prior-list interference.

The cued recognition task provides more information about prior list activation than recall tasks. In recall tasks, prior list activation is apparent as prior-list intrusion errors, which occur only when a prior list item wins the competition with the correct item and the within-list items. These errors are rare because prior list items have less activation, so they usually lose the competition. Recall tasks provide no information about prior list activation when the correct item or a within-list item wins the competition. On those trials, the prior list items could be activated less than current list items or not activated at all. Like recall tasks, the cued recognition task provides information about prior list activation on error trials, when subjects respond “yes” to prior list lures, analogous to prior list intrusion errors. The cued recognition task also provides information about prior list activation on correct trials, when subjects respond “no.” The prior list lure will match the prior list item and activate the “yes” response on all trials, and this will increase RTs for correct “no” responses, as we show below. Thus, cued recognition provides information about prior list activation in both false alarm rate and correct-rejection RT.

The cued recognition task allows stronger conclusions than recall tasks. Position-specific prior list interference is *elicited* by an experimental manipulation (the presentation of a prior-list lure) that allows us to assess prior list activation in RT and error rate on any trial. Observing such interference would support position coding predictions and failing to observe it would falsify them. Position-specific prior list intrusions are *emitted* occasionally by subjects. Observing such intrusions supports position coding predictions but failing to observe them does not falsify them. The prior list item could be activated, as the theory predicts, but not strongly enough to produce an error. The cued recognition task allows us to measure the activation of prior list items when the activation is not strong enough to produce an error.

Position Coding Model

We used a simple generic position coding model, depicted in Figure 1, to formalize predictions and test hypotheses. It embodies the core assumptions of established position coding theories that predict position-specific prior list intrusions, so its predictions generalize to all those theories. Like all position coding theories, the generic model assumes that items on each list are associated with an ordered set of position codes (Anderson & Matessa, 1997; Brown et al., 2000, 2007; Burgess & Hitch, 1999; Farrell, 2012; Henson, 1998; Lewandowsky & Farrell, 2008; Oberauer et al., 2012). Like all position coding theories of prior-list intrusions, the strength of

associations between position codes and items is weaker for the prior list than for the current list because of decay or reduced contextual similarity (Brown et al., 2007; Burgess & Hitch, 1999; Henson, 1998). We assume association strength s equals 1 for the current list and $0 \leq s_{prior} \leq 1$ for the prior list. This is illustrated by the lighter dashed lines in the top left panel of Figure 1. Like all position coding theories, the generic model assumes that items are retrieved by activating position codes. Activation spreads from the position codes to the associated items in proportion to their associative strength. Current list items have stronger associations than prior list items, and so are more likely to be retrieved. Prior list intrusions occur when an item is retrieved from the prior list instead of the current one.

Like all position coding theories, the model assumes that cuing a list position activates position codes in proportion to their distance from the cued position. The activation of the item in position i given a cue in position j is:

$$a(i|j) = s\rho^{|i-j|} \quad (1)$$

where $0 \leq \rho \leq 1$ is the rate at which activation decreases with distance. For the current list, $s = 1$; for the prior list, $s = s_{prior}$. Equation 1 is a common expression for contextual drift (Estes, 1955; Murdock, 1997) that is used explicitly to model within-list distance effects in models of serial recall (Farrell, 2012; Lewandowsky & Farrell, 2008; Logan, 2021; Logan & Cox, 2021). Equation 1 is responsible for order errors (transpositions) that dominate serial recall. It is also responsible for the position specificity of prior list intrusions (and interference). Activation is higher for the cued position than for its neighbors on both the current and prior lists, so items retrieved from both lists are more likely to come from the cued position than its neighbors. The activation across positions in both lists is illustrated in the top right panel of Figure 1. In the generic model, the activation produced by a cue is represented as a vector $\mathbf{m}$ whose elements correspond to the set of possible items, which is shown in the top row of Table 1. The values for items on the current and prior list are given by Equation 1. The values for items that were not on either list are set to 0. Importantly, we assume that the activation values in $\mathbf{m}$ – the results of cuing a position -- are the same whether the retrieval task is recall or cued recognition.

These assumptions are common to all position coding theories of serial recall (Anderson & Matessa, 1997; Brown et al., 2000, 2007; Burgess & Hitch, 1999; Farrell, 2012; Henson, 1998; Lewandowsky & Farrell, 2008; Oberauer et al., 2012) and all position coding accounts of prior list intrusions (Brown et al., 2007; Burgess & Hitch, 1999; Henson, 1998). The theories

share these assumptions but differ in ancillary assumptions like response suppression, primacy gradients, etc. that are designed to address specific effects in serial recall (Lewandowsky & Farrell, 2008). The core assumptions are at issue here. We believe that the predictions of the generic model represent the predictions of the general class of position coding theories and the subclass of position coding theories that address position specific prior list intrusions.

Confirmation of the predictions would support position coding theories of position specific prior list intrusions. Failure to confirm the predictions would falsify some of the assumptions (depending on the nature of the failure), challenge position coding accounts of position specific prior list intrusions, and more generally, challenge the dominance of position coding theories of serial memory.

We apply the generic model to recall and cued recognition tasks by assuming that they access the same memory representations in different ways (i.e., m is the same but the decision process applied to it is different). This assumption has a long history in computational models of memory. Models that relate recognition and recall generally assume that the representations are the same in the two tasks but the decision processes are different (Anderson et al., 1998; Gillund & Shiffrin, 1981; Hintzman, 1984, 1988; Humphreys et al., 1989; Murdock, 1982, 1983; Raijmakers & Shiffrin, 1984). We view memory retrieval as attention turned inward (Logan et al., 2021) and decision processes as mechanisms of attention (Logan et al., 2023a), so we think of recognition and recall as requiring attention to different aspects of memory representations. It is possible that recognition and recall rely on different representations as well as decision processes. Our assumption of a common representation is simpler and consistent with existing computational models.

Serial Recall. In serial recall, m represents the strengths with which the items on the current and prior lists compete with each other for retrieval. We model the competition as a limited-capacity racing diffusion decision process, which accounts for response time (RT) and response probability (accuracy; Logan et al., 2021; Tillman et al., 2020). There is one runner for each possible response, and the first runner to finish is retrieved. The finishing time for each runner depends on its drift rate (v) and its threshold (θ). The drift rate is Equation 1 normalized by 1 plus the length of m , which represents the activity produced by the retrieval cue (Carandini & Heeger, 2012; Lo & Wang, 2006), multiplied by a constant κ , which represents capacity limitations:

$$v_{i,recall} = \frac{a(i|j)}{1+\kappa\|\mathbf{m}\|} \quad (2)$$

If $\kappa = 0$, capacity is unlimited; if $\kappa > 0$, capacity is limited.

The finishing time distribution for each runner is Wald (Inverse Gaussian) with a drift given by Equation 1 and a common threshold. The density and distribution functions are:

$$f(t|v, \theta) = \frac{\theta}{\sqrt{2\pi t^3}} \exp\left[-\frac{(vt-\theta)^2}{2t}\right] \quad (3)$$

and

$$F(t|v, \theta) = \Phi\left(\frac{vt-\theta}{\sqrt{t}}\right) + \exp(2\theta v)\Phi\left(-\frac{vt+\theta}{\sqrt{t}}\right) \quad (4)$$

where $\Phi(\cdot)$ is the standard normal cumulative distribution function. The finishing time distribution for item i in a race between N items is:

$$f(t, i) = f_i(t) \prod_{j \neq i}^N [1 - F_j(t)] \quad (5).$$

The probability that i finishes first is given by the integral of Equation 5. The decision process is illustrated in the second row of Figure 1.

Cued Recognition. Our model of the cued recognition task makes the same assumptions about representation and activation (Figure 1, top) and uses the same vector $\mathbf{m}$ to represent the activation from the position cue, but it makes different assumptions about the decision process applied to $\mathbf{m}$. In serial recall, the decision is based on the activation of individual items, each of which requires a separate response. In cued recognition, we adopted the decision model Logan et al. (2021) applied to the task. In this model, the decision is based only on the activation of the item in the probed position. High activation is evidence for a “yes” response; low activation is evidence for a “no” response. Lures from nearby positions in either list will have greater activation than lures from more distant positions, and so provide evidence for a “yes” response, which increases RT and error rate for the required “no” response. This is illustrated in the bottom panels of Figure 1.

We assume that the activated items on both lists are represented in vector $\mathbf{m}$ with one element for each possible item, whose value is specified by Equation 1, as in serial recall. The probe item is represented as a vector $\mathbf{q}$ with the same dimensionality as $\mathbf{m}$, with 1 in the element representing the probe item and 0 in all other elements. Table 1 presents $\mathbf{q}$ vectors for matching probes, within-list probes, and prior-list probes. The probe is matched to the activated items by taking the dot product of the vectors ($\mathbf{m} \cdot \mathbf{q}$). As illustrated in Table 1, this amounts to multiplying the memory list item corresponding to the probe by 1 and multiplying all other items

by 0, so the match value depends only on the activation of the probe item in the probed position whether the activation comes from the current or prior memory list. Consequently, the dot product $\mathbf{m} \cdot \mathbf{q}$ is given by Equation 1 times 1. The process is illustrated in the bottom panel of Figure 1. The lines represent the activation of $\mathbf{m}$ and the red box represents the nonzero element in $\mathbf{q}$ and the contribution of $\mathbf{m}$ to the dot product. Table 1 contains numerical examples.

The decision process uses the limited-capacity racing diffusion model as serial recall but configures it differently. There are only two runners, one for a “yes” response and one for a “no” response. Equation 1 provides positive evidence for a “yes” response. The larger the value of $a(i,j)$, the more likely the response should be “yes.” The drift rate for the “yes” response is simply Equation 1 normalized by 1 plus the length of $\mathbf{m}$ multiplied by a constant κ to implement capacity limitations and an additional scaling constant λ to balance “yes” and “no” evidence:

$$v_{yes} = \frac{a(i,j)}{1 + \kappa \lambda \|\mathbf{m}\|} \quad (6)$$

Equation 1 provides negative evidence for a “no” response. The higher the value of $a(i,j)$, the less likely the response should be “no.” The racing diffusion model (and neurons) require positive evidence (because the diffusion has a single upper bound and neurons can only have positive firing rates). We create positive evidence by defining the drift rate for the “no” response is the length of the vector $\mathbf{m}$, which represents the largest possible dot product of the probe and the activated memory items (Logan et al., 2021), multiplied by λ to balance “yes” and “no” evidence, and divided by 1 plus the evidence for a “yes” response multiplied by κ to implement capacity limitations:

$$v_{no} = \frac{\lambda \|\mathbf{m}\|}{1 + \kappa a(i,j)} \quad (7)$$

In Equation 7, “no” drift rate decreases as the evidence for a “yes” response increases. “No” drift rate is highest when there is no evidence for a “yes” response (i.e., $a(i,j) = 0$) and lowest on match trials when the evidence for a “yes” response is strongest (i.e., $a(i,j) = 1$).

The denominators that normalize the drift rates are different in recall (Equation 2) and cued recognition (Equations 6-7). In recall, each response is normalized by the total activity produced by the retrieval cue (i.e., the length of $\mathbf{m}$), while in cued recognition, each response is normalized by the activity supporting the other response. Normalization can be viewed as inhibition (Caradini & Heeger, 2012; Lo & Wang, 2006). In recall, each possible response

inhibits every other possible response. In recognition, the two responses inhibit each other, as in lateral inhibition.

The finishing time distributions for “yes” and “no” runners are Wald with drift rates v_{yes} , v_{no} , and thresholds θ_{yes} and θ_{no} . The finishing time distributions for “yes” and “no” responses are

$$f(t, \text{"yes"} | v_{yes}, v_{no}, \theta_{yes}, \theta_{no}) = f(t | v_{yes}, \theta_{yes}) [1 - F(t | v_{no}, \theta_{no})] \quad (8)$$

and

$$f(t, \text{"no"} | v_{yes}, v_{no}, \theta_{yes}, \theta_{no}) = f(t | v_{no}, \theta_{no}) [1 - F(t | v_{yes}, \theta_{yes})] \quad (9).$$

The accuracy of “yes” and “no” responses is given by the integrals of Equations 8 and 9, respectively.

Again, it is important to emphasize that the cued recognition model makes the same assumptions about representation and activation as the serial and recall model. It differs only in the configuration of the decision process, as if subjects are attending to the same information in different ways (Logan et al., 2021, 2023a).

Four Core Predictions

The generic position coding model assumes that memory performance is the result of the activation of position codes, which depends on the distance from the cued position (ρ^{i-j}), and the strength of association (s) between the position codes and the items (Equation 1). We derived four core predictions from the model about performance in memory tasks that require serial retrieval (serial recall, cued recognition).

Prediction 1: Within-list transposition errors should decrease with distance from the intended (cued) position (distances -2 -1 1 2). Performance should be worse for positions ± 1 away from the cued position than for positions ± 2 away. This follows from the distance component of Equation 1. This is a core prediction of position coding theories but it is not unique to them. Alternatives to position coding make the same prediction (Logan, 2021; Solway et al., 2012). Nevertheless, it is important to test. Failing to confirm it would challenge position-coding and non-position-coding theories alike.

Prediction 2: Prior-list intrusion errors should show the same distance effect (-2 -1 1 2). This follows from the distance component of Equation 1 and from the assumption that prior-list and current-list items are associated with the same position codes. This is a core prediction

that is unique to the position coding account of position-specific prior list intrusions in recall and interference in cued recognition. It is not predicted by alternatives to position coding theories.

Prediction 3: The prior-list distance effect should be smaller than the within-list distance effect at corresponding distances (-2 -1 1 2). This follows from the multiplication of s and $\rho^{|i-j|}$ in Equation 1. For the current list, $s = 1$, so the distance effect is simply $\rho^{|i-j|}$. For the prior list, $s = sprior < 1$ so the distance effect is $sprior \times \rho^{|i-j|}$, which is smaller. This is a core prediction of the position coding account of position-specific prior list intrusions and interference but it is not unique. Theories that assume no such intrusions or interference also predict a smaller (i.e., null) effect of prior list distance.

Prediction 4: Prior list errors should peak at distance = 0 (distances -1 0 1). This follows from the distance component of Equation 1. This is the strongest prediction of the position coding model. It predicts position-specific prior list intrusions in serial recall, and it predicts position-specific prior list interference in cued recognition. It is unique to the position coding account. Failure to confirm this prediction would seriously challenge the position coding account of position-specific prior list intrusions.

Simulations. We ran simulations of the position coding model to illustrate the four predictions in recall and cued recognition and to assess the effects of varying prior list strength ($sprior$) on the predictions. We assumed five-item lists that were cued in the third (middle) position and used Equation 1 to specify activation for distances of -2, -1, 1, and 2 for within list errors and distances of -2, -1, 0, 1, and 2 for prior list errors. We used Equation 2 to simulate recall and Equations 6-7 to simulate cued recognition. In all simulations, $\rho = .5$ and $\kappa = .2$ for both tasks, $\theta_{recall} = 10.0$ for recall, and $\theta_{yes} = 2.8$, $\theta_{no} = 3.0$, and $\lambda = .8$ for cued recognition. Further details of the simulations are presented in Appendix A. MATLAB code for the simulations is posted on OSF.

Figure 2 shows the effect of prior list strength ($sprior = .1, .2, .3, .5, .7$) on predicted distance effects. The top panel shows predicted within-list transposition errors and prior list intrusions in recall. There are strong within-list distance effects (-2 -1 1 2) at all values of $sprior$, confirming Prediction 1. There are prior-list distance effects (-2 -1 1 2), confirming Prediction 2. Within-list distance effects were stronger than prior-list distance effects at all values of $sprior$, confirming Prediction 3. Prior-list distance effects peaked at the cued position (-1 01), confirming Prediction 4 for values of $sprior \geq .2$. Thus, the position coding model predicts prior

list intrusions in recall. The middle panel shows predicted error rates for within-list and prior-list lures in cued recognition, which also confirm the four predictions. There are strong within-list distance effects and weaker position-specific prior list interference effects with a peak at the cued position at all values of *sprior*. The bottom panel shows predicted RTs for correct responses to matching probes, within-list lures, and prior-list lures in cued recognition (i.e., the additional information that cued recognition provides about prior list activation). The RTs show within-list distance effects and weaker position-specific prior list interference that peaks at the cued position at all values of *sprior*, confirming the four predictions. Prior list interference is greater the stronger the associations of position codes prior list items. Thus, the position coding model predicts position-specific prior list interference in cued recognition over a broad range of prior list association strengths.

The effects of the prior list appear stronger in cued recognition than in recall. This follows from the model. Prior list items may be activated to the same extent in recall and cued recognition, but prior list intrusions only occur if the prior-list item happens to finish first in the decision process, before the correct item or another item from the current list. Cued recognition probes the activation of prior list items directly on every trial, showing prior list interference in both accuracy and RT.

The Experiments

We conducted 12 experiments to test for the prior list intrusions and interference predicted by the position coding model. Experiments 1 and 2 tested serial recall to ensure that position-specific prior list intrusions would occur with our materials (consonants), list length (6 items), exposure duration (1000 ms), and retention interval (1000 ms). The remaining experiments tested cued recognition to determine whether the same study conditions would produce the predicted position-specific prior list interference. Experiments 3-10 manipulated factors intended to increase the likelihood that position codes would be activated. We presented the position component of the probe 500 ms before the probe letter appeared so subjects could begin to focus on the cued position in the list. We cued position with a number or a spatial display depicting its position. Experiments 11-12 tested cued recognition with sequential presentation of the lists instead of simultaneous presentation. Most studies of serial recall,

including those that address position-specific prior list intrusions, use sequential presentation. The goal was to generalize our results and strengthen connections to that literature.

Experiments 1-2: Serial Recall

The first two experiments used serial recall to determine whether it is possible to get position-specific prior list intrusions with the simultaneously presented six-item lists used later in the cued recognition experiments. The purpose was to establish that items on the current list and prior list could be associated with position codes under these conditions. Subjects were given lists of six consonants to remember, presented in a row on the computer screen for 1000 ms. The screen went blank for 1000 ms and then a screen containing “RECALL” appeared, cuing subjects to type the list into their computer keyboards in correct order. Their recall errors were scored as *within-list transpositions* or *prior-list intrusions*, which were analyzed as a function of their distance in the list from the correct letter. In theory, these errors reflect the same activation measured by within-list lures and prior-list lures, respectively, in cued recognition.

The experiments were the same except for the way the lists were constructed. Experiment 1 used lists that were *constrained* so that no letters repeated from one list to the next. Experiment 2 used lists that were *unconstrained*, so letters could repeat from one list to the next. The difference in the lists addresses an alternative interpretation of the cued recognition results and will be addressed in the General Discussion.

Each experiment tested the four predictions for error rate derived from the position coding model: (1) Within-list transposition errors should show a distance effect, with more errors from ± 1 position away from the correct position than from ± 2 positions away. (2) Prior-list intrusion errors should show the same distance effect for positions ± 1 and ± 2 away from the correct position. (3) The prior list distance effect should be smaller than the within-list distance effect at corresponding positions, reflecting the reduced strength of prior-list associations (*sprior*). (4) Prior list intrusion errors should show position-specific interference, manifest as more errors from the correct position in the prior list (distance = 0) than for lures from adjacent positions (distance = ± 1).

Method

Subjects

Each experiment tested 32 subjects recruited online through Prolific (<https://www.prolific.co/>). We included only subjects 18-40 years of age, located in the USA, with English as first language, with an approval rating of at least 95%, who typed at least 40 words per minute (WPM) on the typing test. Subjects who participated in one experiment were excluded from the others. Experiments 1-2 involved a single 1.5-hour session. Subjects were paid USD \$12 per hour. The study was approved by the Vanderbilt University Institutional Review Board.

Subjects reported their age and gender. The mean age (standard deviation in brackets) of the subjects was 30.97 (6.01) and 31.94 (5.60). for Experiments 1-2 respectively. The gender distribution (male:female:prefer-not-to-say) was 15:17:0 and 26:6:0 for Experiments 1-2 respectively. Mean speed on the typing test was 60.80 (17.70) and 64.73 (15.10) for Experiments 1-2, respectively. Mean accuracy was 0.9173 (0.0430) and 0.9272 (0.0427) for Experiments 1-2, respectively.

Apparatus and Stimuli

The experiments were conducted online on subjects' personal computers. Subjects were instructed to use Google Chrome or Mozilla Firefox to complete the experiment. Phone and tablet users were excluded in the Prolific intake, and the experiment would not run on their browsers. The trials for each session were generated individually and sent to subjects' computers using a custom Python backend. The experiment was controlled by Javascript in the web browser using a custom function written to operate in jsPsych (de Leeuw, 2015). When the experiment started, subjects' web browsers were instructed to enter fullscreen mode to reduce distraction.

The memory lists consisted of six uppercase letters selected at random from the set of 20 consonants (excluding vowels and Y), displayed in a row. Experiment 1 used constrained lists, in which no letters were repeated from one trial to the next. Experiment 2 used unconstrained lists, in which letters were allowed to repeat from one trial to the next. Characters were presented in a monospaced typeface (Courier New or Courier, displayed in white, 45 pixels high. The background of the display was set to mid-gray ([127, 127, 127] in 24-bit RGB values).

Procedure

In both experiments, each trial began with a fixation cross presented in the center of the screen for 1000 ms. Then the memory list was presented for 1000 ms, followed by a blank screen for 1000 ms, and then a probe display containing the word RECALL appeared. Subjects were required to type the letters in the list in response to the probe, and the letters they typed were echoed on the screen in left to right order, as in typing text. They were told to type six letters on each trial and hit “return” when they were finished. Then the screen went blank for a 1000 ms intertrial interval. Space and backspace keys were disabled. There were 480 trials in each experiment. Breaks were given every 80 trials.

The instructions were written and presented using a self-paced series of manually controlled slides. Subjects were allowed to review the instructions if they wished. Each subject completed a typing test to ensure they had enough skill to execute keystrokes automatically, without hunting and pecking on the keyboard, which might limit performance. The typing test involved typing a paragraph about the many merits of border collies (Logan & Zbrodoff, 1998). The paragraph was presented on the top of the screen and subjects’ keystrokes were echoed in a panel below the paragraph.

At the end of each block, a screen was presented indicating the overall accuracy for the preceding block, and subjects were allowed to take a self-timed break. Every 5 minutes, the experiment checked whether accuracy was greater than 60%. If subjects fell below this criterion, they were warned to improve performance and given an opportunity to review the instructions. On the third warning, subjects were excluded from the experiment but paid nevertheless.

Data Analysis

Experiments 1 and 2 were designed to measure within-list transposition errors and position-specific prior list intrusions in serial recall. We identified within-list errors as items from the list that were recalled in the wrong position. Distance was defined as the signed difference between the position in the recall sequence and the position in the memory list. We included distances (-2 -1 1 2) to parallel the distance manipulation in the cued recognition experiments. We identified prior-list errors as recalled items that were in the prior list and not in the current list. We defined distance as the signed difference between the position in the prior list and the position that was reported in the current list. For example, if the current list is ABCDEF and the prior list is GHIJKL, then recalling K (in error) after recalling A and B is a

prior list intrusion with distance = 2. We did not normalize within-list transpositions or prior-list intrusions for availability.

We tested the four predictions with contrasts. We tested Predictions 1 and 2 (within- and prior-list distance effects) using contrast weights (-1 1 1 -1) for distances (-2 -1 1 2) to compare distances ± 1 and ± 2 . We tested Prediction 3 by comparing the (-2 -1 1 2) distance contrast for the current list with the (-2 -1 1 2) distance contrast for the prior list, using weights (-1 1 1 -1) for the current list and (1 -1 -1 1) for the prior lists. We tested Prediction 4 (position specific prior list intrusions) using weights (-1 2 -1) for distances (-1 0 1) in the prior list. This is the critical contrast that tests for position-specific prior list intrusions.

For each contrast, we divided the data for each subject into the relevant cells (4 distances for within-list lures; 5 distances for prior list lures) and calculated the proportion of errors. Then we calculated the contrast values for each subject, multiplying the error rates by the contrast weights and summing them. Then, we did a t test asking whether the mean contrast was significantly greater than zero. The error term was the standard error of the mean contrast value. We also counted the number of subjects who showed an effect in the expected direction and reported JZS Bayes Factors (BF) to quantify support for null (BF_{01}) and alternative (BF_{10}) hypotheses.

Our contrasts provide inferential statistical tests of specific hypotheses derived from theory. They evaluate relations between conditions, and the error variability depends on those relations, which cannot be expressed as error bars around individual means. Because of this, we do not present error bars in any of our figures.

Data and programs for presenting the task and analyzing the data for all experiments in this article are available on the Open Science Framework at <https://osf.io/j4z7a/>.

Results

Mean within-list and prior-list error rates for Experiments 1 and 2 are plotted as a function of distance in the left and middle panels of Figure 3, respectively. Table 2 contains contrasts evaluating distance effects. The right panel of Figure 3 contains within-list and prior-list error rates from *position-cued recall* experiments that used the same (unconstrained) lists and probed recall of a single item with a spatial cue (e.g., ###?##, where the underline represents a caret ^ pointing at the cued position; Logan et al., 2023a).

The data from Experiments 1 and 2 confirmed the four core predictions of position coding theory. There were significant distance effects $(-2 -1 1 2)$ in each experiment for both within- and prior-list errors, confirming Predictions 1 and 2. Within-list distance effects were significantly stronger than prior-list distance effects in each experiment, confirming Prediction 3. There were significant position-specific prior list intrusions in each experiment. Intrusions were more frequent at lag 0 than at lags ± 1 in 30 out of 32 subjects in Experiment 1 and in 31 out of 32 subjects in Experiment 2. The contrast assessing position specific prior list intrusions $(-1 0 1)$ was significant in each experiment.

Experiment 2 replicated the results of Experiment 1 very closely. The patterns in Figure 3 are very similar. Table 2 contains t tests comparing prior list contrasts $(-1 0 1)$ and $(-2 -1 1 2)$, within list contrasts $(-2 -1 1 2)$, and contrasts comparing $(-2 -1 1 2)$ in prior versus current lists between experiments. None of the t tests were significant.

The cued recall data in Figure 3 build a bridge between serial recall and cued recognition. Cued recall requires subjects to recall items, like serial recall, while focusing on a single item in the memory list in response to a cue, like cued recognition. The cued recall data were obtained in dual task experiments in which subjects were given 6-item lists to remember and then were given two spatial cues in succession indicating the two items to be reported (e.g., ##### followed by ##### cues the report of the second and the fifth item in the list). The interval between the two cues varied to produce dual-task interference (100, 300, or 900 ms). The data in Figure 3 collapse over four experiments, the interval between cues, and responses to the first and second cue to obtain sufficient observations. The contrast testing position-specific prior list intrusions was significant in each of the four experiments. Figure 3 also shows that within-list distance effects $(-2 -1 1 2)$ were stronger than prior list distance effects $(-2 -1 1 2)$, as in serial recall. In theory, this means that the cue in cued recall activated position codes, the position codes activated items on both lists, and the activation was greater for items on the current list. Thus, the position cue in cued recognition should also activate a position code and the items associated with it on both lists.

Discussion

Experiments 1 and 2 confirmed the four predictions of position coding theory in serial recall and set the stage for the cued recognition experiments to follow. They show that position-

specific prior list intrusions can be observed under our list presentation conditions if the retrieval task is serial recall. In theory, this means that position codes were activated in serial recall, and they activated associated items on the current and prior lists. The data from Logan et al. (2023a) in Figure 4 show that position-specific prior list intrusions can also be observed in cued recall. In theory, this means that the position cues activated position codes, which activated items in the current and prior list. It means that the position cues in the cued recognition experiments should also activate position codes, which should activate items on the current and prior list. Cued recognition allows us to test that activation directly, using lures from uncued positions in the current list and lures from all positions in the prior list.

Experiments 3-10: Cued Recognition

The simulations established that position coding theories predict position-specific prior list interference when cued recognition is tested with prior list lures. Experiments 1 and 2 established that position-specific prior list intrusions occur under our presentation conditions in serial recall. Now, we report the cued recognition experiments that test for position-specific prior list interference under the same conditions. We ran a series of eight experiments with manipulations intended to enhance the activation of position codes. We began with probes that cued position spatially, following our previous experiments on cued recognition (Logan et al., 2021; Logan et al., 2023b). The probe consisted of five # symbols and a letter with a caret (^) underneath it to indicate the cued position (e.g., ####D##, where the underline represents the caret). Then we tried cuing spatial position numerically (e.g., 4D cues the fourth position), thinking that numeric cues might cue position more directly. Then we tried pre-cuing position so subjects could begin to focus on the cued position in the memory list before the letter probe was presented (Logan et al., 2023b). We first ran the series with constrained lists (items could not repeat in consecutive lists) and then replicated it with another four experiments that used unconstrained lists (items could repeat in consecutive lists) to address alternative interpretations (see General Discussion). Altogether, we ran eight experiments in a 2 (probe type) x 2 (pre-cue) x 2 (list type) design.

Each experiment had the same basic design to test the predictions of position coding theory. Half of the probes contained *targets* that matched the item in the cued position in the memory list. The other half of the probes contained lures that did not match the item in the cued

position in the memory list. Half of the lures (*within-list lures*) were sampled from the current list -2, -1, 1, or 2 positions away from the cued position to ensure that subjects focused on the cued position. The other half of the lures (*prior-list lures*) were sampled from the prior list -2, -1, 0, 1, or 2 positions away from the cued position to test for position-specific prior list interference.

Each experiment tested the four predictions for RT and error rate derived from the position coding model: (1) Within-list lures should show a distance effect, with worse performance for lures ± 1 position away from the cued position than for lures ± 2 positions away. (2) Prior-list lures should show the same distance effect for lures ± 1 and ± 2 positions away from the cued position. (3) The prior list distance effect should be smaller than the within-list distance effect at corresponding positions, reflecting the reduced strength of prior-list associations (*sprior*). (4) Prior list lures should show position-specific interference, manifest as worse performance for lures from the cued position in the prior list (distance = 0) than for lures from adjacent positions (distance = ± 1). This is the strongest prediction of the position coding model. The same activation of prior list items predicts position specific interference in cued recognition and position specific intrusions in serial recall. Failure to confirm this prediction would seriously challenge the position coding account of position-specific prior list intrusions.

Method

Subjects

Each experiment recruited 32 subjects from Prolific using the same selection criteria as Experiments 1 and 2. The mean age (standard deviation in brackets) of the subjects was 28.63 (6.96), 30.16 (5.73), 30.53 (5.86), 29.91 (6.79), 29.22 (5.53), 29.06 (5.54), 31.69 (4.43), and 28.66 (6.39) for Experiments 3-10 respectively. The gender distribution (male:female:prefer-not-to-say) was 15:17:0, 18:14:0, 16:15:1, 16:16:0, 18:14:0, 14:17:1, 16:16:0, and 16:16:0 for Experiments 3-10 respectively. No typing test was required because subjects only pressed one of two keys.

Apparatus and Stimuli

The apparatus was the same as in Experiments 1 and 2 (subjects' home computers), and the memory lists were the same: six consonants randomly selected from a set of 20 with the

constraint that no items repeat on consecutive lists (Experiments 3-6) or with no constraint (items could repeat on consecutive lists; Experiments 7-10). The presentation duration of the memory lists (1000 ms), the retention interval between the memory list and the complete probe (1000 ms), and the intertrial interval (1000 ms) were the same as in Experiments 1 and 2, but the probe differed. Experiments 3, 4, 7, and 8 used spatial probes, which displayed an array of five # symbols plus a probe letter with a caret (^) underneath it in the cued position (e.g., ##C###, where the underline represents the caret). Experiments 5, 6, 9, and 10 used numeric probes, which displayed a single number and a probe letter presented in the center of the screen (e.g., 2C). Experiments 3, 5, 7, and 9 had blank 1000 ms retention intervals followed by complete probes (##C### or 2C). Experiments 4, 6, 8, and 10 pre-cued the probed position 500 ms after the memory list. The position component of the probe was presented with a blank instead of the probe letter for 500 ms (e.g., ##_### or 2), followed by the complete probe (##C### or 2C), in which the blank position in the precue was replaced by the probe letter.

Procedure

Each trial began with a fixation cross presented in the center of the screen for 1000 ms. Then the memory list was presented for 1000 ms. In Experiments 3, 5, 7, and 9, the memory list was followed by a blank screen for 1000 ms, and then the complete probe display (containing the position cue and the probe letter) appeared. In Experiments 4, 6, 8, and 10, the position cue appeared 500 ms after the memory list for 500 ms, when the probe letter was added to complete the probe display. In all experiments, the probe display remained onscreen until subjects responded, and then the screen went blank for a 1000 ms intertrial interval.

There were 480 trials per session, constructed by randomly interleaving 240 trials of 120 targets and 120 within-list lures with 240 trials of 120 targets and 120 prior list lures. The targets were no different in the two sets of trials but the lures differed. The design for targets and within-list lures involved 6 probe positions and 4 distances (-2 -1 1 2), creating 24 “no” trials, plus 24 “yes” trials (6 probe positions replicated 4 times), for a total of 48 trials for one replication. The 240 target and within-lure trials replicated this design 5 times. The design for targets and prior-list lures involved 6 probe positions and 5 distances (-2 -1 0 1 2), creating 30 “no” trials and 30 “yes” trials (6 probe positions replicated 5 times), for a total of 60 trials for one replication. The 240 target and prior-list trials replicated this design 4 times. For each

subject, the 240 trials for within- and prior-list lures were randomized separately and then combined randomly to produce the final set of 480 trials.

Subjects were told to indicate whether the cued letter in the probe was presented in the same position in the memory list, pressing the M (or Z) key on the keyboard to indicate a “yes” response and the Z (or M) key to indicate a “no” response. Mapping of response categories to keys was counterbalanced between subjects. The instructions were written and presented using a self-paced series of manually controlled slides. Subjects were allowed to review the instructions if they wished.

Subjects had to respond within 3000 ms of the presentation of the probe or the trial was terminated with the message “TOO SLOW” presented centrally in red font for 3000 ms. These trials were excluded from analysis and treated as errors in calculating feedback during the task. At the end of each block, a screen was presented indicating the overall accuracy for the preceding block, and subjects were allowed to take a self-timed break. Every five minutes, the experiment checked whether accuracy was greater than 60%. If subjects fell below this criterion, they were warned to improve performance and given an opportunity to review the instructions. On the third warning, subjects were excluded from the experiment.

Data Analysis

We tested the four predictions of position coding theory with four contrasts on the mean RTs and error rates. The within-list distance effects in Prediction (1) were tested with a contrast using weights (-1 1 1 -1) for distances (-2 -1 1 2). The corresponding prior-list distance effects in Prediction (2) were tested using the same contrast weights for the same distances in the prior list. The attenuation of distance effects in prior lists relative to current list effects in Prediction (3) was tested with contrast weights (-1 1 1 -1) for within-list distances (-2 -1 1 2) and contrast weights (1 -1 -1 1) for prior-list distances (-2 -1 1 2). The position-specific prior list interference in Prediction (4) was tested with contrast weights (-1 2 -1) for prior-list distances (-1 0 1). The confidence intervals around contrast values cannot be expressed as error bars around the component RTs and error rates. Confidence intervals around mean RTs and error rates cannot support inferences about the significance of the contrasts. Consequently, we present no error bars in the figures.

Results

Mean RT for correct responses (top) and error rate (bottom) for matches (“yes” response), within-list lures (“no” response), and prior-list lures (“no” response) are plotted as a function of distance from the cued position in Figure 4 for Experiments 3-6 and Figure 5 for Experiments 7-10. The pattern of the data was very similar across experiments. It shifted downward but remained the same when probe position was pre-cued (Experiments 4, 6, 8, and 10 vs. Experiments 3, 5, 7, and 9), following previous research (Logan et al., 2023b). RTs were longer with numeric position cues (Experiments 5, 6, 9, 10 vs. Experiments 3, 4, 7, 8) but the pattern of the data was very similar. The pattern was the same whether lists were constrained to exclude letter repetitions in consecutive lists (Experiments 3-6) or unconstrained to allow repetitions (Experiments 7-10). Serial position data are presented in Appendix C in Figure C1.

Position Coding Model Predictions

We assessed the four predictions of the position coding model separately for each experiment. Contrasts evaluating the predictions are presented in Table 3 for Experiments 3-6 and Table 4 for Experiments 7-10.

Prediction 1: Distance Effects for Within-List Lures (-2 -1 1 2). In each experiment, subjects were able to focus on the cued item and ignore the other items in the list: d 's, calculated from hit rates from “yes” trials and false alarm rates from within-list lures, averaged (SEM in brackets) 2.0942 (.1504), 2.5683 (.1527), 1.9732 (.1575), and 2.2042 (.1435) in Experiments 3-6, respectively, and 2.1247 (0.1262), 2.3239 (.1187), 1.9038 (0.1379), and 2.1344 (0.1630) for Experiments 7-10, respectively. In theory, this means position codes for the cued positions were activated. Current list items are associated with position codes with strength ($sprior$) = 1, so within-list lures should be activated in proportion to their distance from the cued position (Equation 1). In each experiment, RT and error rate for within-list lures decreased substantially as the distance between the probed position and the lure's position increased. The within-list distance contrasts were highly significant for both RT and error rate in each experiment. In theory, this means position codes activated neighboring items in proportion to their distance from the cued position.

Prediction 2: Distance Effects for Prior-List Lures (-2 -1 1 2). Having established the conditions necessary to produce prior list intrusions (activation of position codes, activation of

neighboring within-list items), the question is whether lures from prior lists produced the predicted distance effects at the same distances as within-list lures. In each experiment, the answer was clearly negative. RTs and error rates for prior list lures showed no effect of distance in any experiment. The contrast was significant only for RT in Experiment 10, where RTs were *shorter* for distances of ± 1 than for distances of ± 2 (Tables 3 and 4). These results fail to confirm the prediction, but the prediction for distances $(-2 -1 1 2)$ is not strong. The simulations in Figures 2 and 3 show weak effects at these distances.

Prediction 3: Distance Effects are Stronger Within-List than Between-List. Position coding theories assume that items in the current list are more strongly associated with position codes than items in the prior list. This implies that within-list distance effects should be stronger than prior-list distance effects at the same distances $(-2 -1 1 2)$. Contrasts comparing within- and prior-list distance effects supported this prediction. They were highly significant for error rate in every experiment and highly significant for RT in every experiment but Experiment 5 (Tables 3 and 4). On the balance, the data confirm the prediction.

Prediction 4: Distance Effects with Prior List Lures $(-1 0 1)$. The results supporting the first and third predictions establish the conditions necessary to produce position-specific prior list interference. The probe activates the position code in the cued position, which activates items associated with it and its nearby neighbors on the current list and, to a lesser extent, on the prior list. Prior list activation should be strongest at the cued position, so interference should peak at distance = 0. The prior list contrast comparing distance = 0 with distance = ± 1 tests this prediction directly. The contrast for RT was not significant in any experiment (Tables 3 and 4). The contrast for error rate was significant only in Experiment 7 (i.e., 1 out of 16 contrasts), but the difference may be due to the negative $(-2 -1 1 2)$ prior list distance contrast, in which error rate was lower for distances ± 1 than for ± 2 . A contrast comparing error rates at distances ± 2 with distance 0 found no significant difference, $t(31) = 0.000003$, $SEM = 0.0215$, $p = .9999$, $BF_{10} = 0.1888$. On the balance, the data disconfirm the prediction. They challenge the position coding account of position-specific prior list intrusions in serial and cued recall.

Between-Experiment Comparisons

Experiments 3-10 manipulated list type (constrained or unconstrained), probe type (spatial or numerical), and pre-cue delay (0 or 500 ms) between experiments, attempting to

increase the likelihood of activating position codes and to address alternative interpretations. Each experiment involved a single combination of these variables, so their effects were not assessed with the inferential statistics reported so far. Here, we take advantage of the factorial structure of the between-experiment manipulations and evaluate their effects in 2x2x2 between-subject analyses of variance (ANOVAs). We performed one set of ANOVAs on mean RTs and error rates to assess the effects of the manipulations on cued recognition performance. We performed four sets of ANOVAs on the contrasts evaluating the four predictions of the position coding theory, asking whether the effects assessed with the contrasts interact with the between-experiment manipulations. Summary tables for the ANOVAs are presented in Appendix B.

Mean RT and Error Rate. We focused on RTs for “yes” (match) responses. They appeared to change in the same way across experiments as “no” responses to within- and prior-list lures. They were based on more observations than “no” responses (240 vs. 120 for each type of lure) and had not been tested in any of the previous analyses. Averaged across experiments, “yes” RT was 315 ms shorter with a 500 ms pre-cue than without, suggesting that the pre-cue allowed time to focus on the cued position (Logan et al., 2023b), which should increase the activation of position codes. “Yes” RT was 170 ms longer with numeric probes than spatial probes, and not affected by list type (difference = 27 ms). These results were confirmed by significant main effects of pre-cue delay and probe type in the ANOVA on mean RTs. No other effects were significant. Averaged across experiments, error rate on “yes” trials was 0.0281 smaller with a pre-cue than without, 0.0122 smaller with spatial probes than with numeric probes, and 0.0023 smaller with unconstrained lists than with constrained lists. The pre-cue effect was the only significant effect in the analyses. The summary tables for the ANOVAs are presented in Table B1 in Appendix B.

Prediction 1: Within Distance (-2 -1 1 2). There were no significant effects in the ANOVA on the RT contrasts. The effects were consistent across experiments. The null effect of probe delay is consistent with the null interaction between probe delay and distance in Logan et al. (2023b). The only significant effects in the ANOVA on the P(Error) contrasts were the main effects of probe type and probe delay. The contrasts were larger for spatial probes and larger for the 500 ms delay. Summary tables for the ANOVAs are presented in Table B2 in Appendix B.

Prediction 2: Prior Distance (-2 -1 1 2). There were no significant main effects or interactions in the ANOVAs on RT and P(Error). The null prior distance effects were consistent across experiments. Summary tables for the ANOVAs are presented in Table B3 in Appendix B.

Prediction 3: Within vs. Prior Distance (-2 -1 1 2). There were no significant effects in the ANOVA on RT, indicating that within-list distance effects were stronger than prior-list distance effects in each experiment. The effect of probe delay was significant in the ANOVA on P(Error), indicating smaller differences between within and prior distance effects with the 500 ms delay, which may be a floor effect. Summary tables for the ANOVAs are presented in Table B4 in Appendix B.

Prediction 4: Prior Distance (-1 0 1). A sharp peak in interference at distance = 0 is the strongest prediction of the position coding model (Figure 2). There were no significant effects in the ANOVA on this contrast in RT, indicating that the null distance effect replicated consistently across experiments. List type was the only significant effect in the ANOVA on the contrast in P(Error), indicating a smaller contrast value with unconstrained lists. These results disconfirm the prediction and thereby challenge the position coding account of position-specific prior list intrusions. Summary tables for the ANOVAs are presented in Table B5 in Appendix B.

Summary. The ANOVAs provided statistical support for the differences in overall RT and error rate between experiments. There were few differences in the distance contrasts across experiments, suggesting that the contrasts replicated well.

Discussion

Across experiments, overall performance varied with pre-cue delay and probe type but the pattern of distance effects remained the same. Distance had strong effects on within-list lures but null effects on prior-list lures, measured either at (-2 -1 1 2) or (-1 0 1). This pattern of effects confirms Predictions 1 and 3 about within-list lures but disconfirms Predictions 2 and 4 about prior list lures. The results have strong implications for position coding theories of serial order. The experiments established the conditions necessary (in theory) to produce position-specific prior list interference. The large d' values comparing “yes” and within-list “no” responses suggest that the position code for the cued position was activated more than the others. The within-list distance contrast (-2 -1 1 2) suggests that position codes for nearby items were activated in proportion to their distance from the cued position. Within-list distance effects were

stronger than prior-list distance effects, suggesting that position codes activated items in the current list more strongly than items in the prior list. In theory, the activated position codes should activate items from the prior list in proportion to their distance from the cued position. This activation should reduce the drift rate for “no” responses (Equation 7), slowing RT and increasing error rate in (inverse) proportion to their distance from the cued position, but across experiments, distance had no effect on either measure. This key prediction was not confirmed in any of the eight experiments. This challenges the position coding account of position-specific prior list intrusions and position coding theories more broadly. We discuss alternative interpretations and implications in the General Discussion, after reporting the last two experiments.

Experiments 11-12: Cued Recognition with Sequential List Presentation

In all the experiments so far, the memory lists have been presented simultaneously for 1000 ms. In the literature, experiments on position-specific prior list intrusions and serial memory in general usually present the memory list sequentially, one item at a time. Experiments 1 and 2 show that position-specific prior list intrusions can be found with simultaneously presented lists, but the effects may be more robust with sequentially presented lists. Experiments 11 and 12 replicated the cued recognition results with sequentially presented lists to determine whether position-specific prior list interference would occur with those lists.

Experiments 11 and 12 were replications of Experiments 5 and 9 with sequential lists. In both experiments, the position cues were numeric and position and item cues were presented simultaneously (e.g., 5R). Experiment 11 used constrained lists. Experiment 12 used unconstrained lists. Each experiment tested the four predictions of position coding theory.

Method

Subjects

Each experiment recruited 32 subjects from Prolific, using the same selection criteria as the previous experiments. The mean age (standard deviation in brackets) of the subjects was 30.19 (5.29) and 29.09 (5.96) for Experiments 11 and 12, respectively. The gender distribution (male:female:prefer-not-to-say) was 22:10:0 and 18:13:1, respectively.

Apparatus and Stimuli

These were the same as in Experiments 5 and 9 (numeric cues, no pre-cue delay), except that the lists were presented sequentially. Each item appeared in the center of the screen for 500 ms, whereupon it was replaced by the next item. The retention interval, which began after the last item was erased from the screen, was 1000 ms, as in the previous experiments.

Procedure

The procedure was the same as in Experiments 3-10.

Results

Mean RT for correct responses (top) and error rate (bottom) for matches (“yes” response), within-list lures (“no” response), and prior-list lures (“no” response) for each experiment are plotted as a function of distance from the cued position in Figure 5. The results replicated Experiments 5 and 9 closely. There were strong distance effects for within-list lures and null distance effects with prior list lures.

As before, d' comparing hit rates from “yes” trials with false alarm rates from within-list “no” trials showed that subjects were able to focus sharply on the cued position, activating a position code in theory. The d 's (SEM in brackets) were 2.0872 (0.1459) and 2.0642 (0.1549) for Experiments 11 and 12, respectively. We tested the four predictions of position coding theory in each experiment using contrasts presented in Table 6. The contrasts evaluating within-list distance effects (-2 -1 1 2) and the contrasts comparing within-list and prior-list distance effects were highly significant for RT and error rate in each experiment, confirming Predictions 1 and 3. The contrasts evaluating prior list distances did not show evidence of interference, disconfirming Predictions 2 and 4. The contrast for distances (-2 -1 1 2) was not significant for RT or error rate in either experiment, nor was the critical contrast for distances (-1 0 1).

Discussion

The results show that the main findings in cued recognition can be replicated with sequentially presented lists. Thus, the findings generalize to conditions more typical of the literature on position-specific prior list intrusions and the broader literature on serial memory. As in the previous cued recognition experiments, these experiments established the conditions

necessary (in theory) to produce position-specific prior list interference (activating position codes, activating nearby position codes more strongly than remote ones, activating the current list more than the prior list), but none was observed in either experiment. The results confirmed Predictions 1 and 3 but disconfirmed Predictions 2 and 4. Now there are 10 experiments showing that result.

General Discussion

The experiments were designed to test predictions derived from the position coding account of position-specific prior list intrusions. We showed that a position coding model that produces prior list intrusions must also produce position-specific prior list interference in a cued recognition task (when coupled with an appropriate decision process; Figures 1-2). We failed to find such interference in 10 experiments. Figure 7 plots the mean observed RTs and error rates (solid lines) across all 320 subjects in Experiments 3-12 for match responses, within-list lures, and prior-list lures as a function of distance from the cued position. The pattern of the observed data does not resemble the position coding predictions in Figure 2 very closely. The observed pattern is most similar to the predictions with the lowest prior list strength ($sprior = 0.1$). However, these strengths and probabilities are too low to account for the position-specific prior list intrusions we observed in serial recall in Experiments 1 and 2 (cf. predictions in the top left panels of Figure 2). Thus, the results challenge the position coding account of prior list intrusions. They challenge position coding theories more generally because their account of position-specific prior list intrusions is a unique prediction that distinguishes them from other theories of serial memory (Henson et al., 1996; Osth & Hurlstone, 2022).

Each experiment tested four predictions derived from the core assumptions of position coding theory. The theory assumes that position codes are activated in proportion to their distance from the cued location, the position codes activate items associated with them on the current list and the prior list in proportion to their activation, and associations to the current list are stronger than associations to the prior list. This leads to the four predictions, which we tested on the data from all 320 subjects in Experiments 3-12.

Prediction 1. For within-list probes, RT and error rate should both decrease with distance, assessed with contrasts comparing positions (-2 -1 1 2). Prediction 1 was confirmed. The within-list distance effects (-2 -1 1 2) were strong and highly significant. For RT, $t(319) =$

8.9228, $SEM = 8.9228$, $p < 0.0001$, $N > 0 = 269$, $BF_{10} > 1000$; for P(Error), $t(319) = 9.4483$, $SEM = 0.0097$, $p < 0.0001$, $N > 0 = 248$, $BF_{10} > 1000$.

Prediction 2. For prior-list probes, RT and error rate should also decrease with distance over the same set of distances (-2 -1 1 2). Prediction 2 was disconfirmed. The prior-list distance effects (-2 -1 1 2) were null. For RT, $t(319) = 0.1882$, $SEM = 6.7750$, $p = 0.8508$, $N > 0 = 156$, $BF_{10} = 0.0638$; for P(Error), $t(319) = 2.1014$, $SEM = 0.0056$, $p = 0.0364$, $N > 0 = 150$, $BF_{10} = 0.5508$.

Prediction 3. The (-2 -1 1 2) distance effect should be stronger for within-list lures than for prior-list lures. Prediction 3 was confirmed. Within-list distance effects (-2 -1 1 2) were much stronger than prior-list distance effects (-2 -1 1 2). For RT, $t(319) = 12.7175$, $SEM = 11.1316$, $p < 0.0001$, $N > 0 = 248$, $BF_{10} > 1000$; for P(Error), $t(319) = 7.0079$, $SEM = 0.0114$, $p < 0.0001$, $N > 0 = 236$, $BF_{10} > 1000$.

Prediction 4. For prior-list probes, RT and error rate should peak at distance = 0, assessed with contrasts comparing positions (-1 0 1). Prediction 4 was disconfirmed. There was no peak at distance = 0 for prior list lures. For RT, $t(319) = 1.0564$, $SEM = 10.0030$, $p = 0.2916$, $N > 0 = 152$, $BF_{10} = 0.1089$; for P(Error), $t(319) = 0.0177$, $SEM = 0.0072$, $p = 0.2810$, $N > 0 = 118$, $BF_{10} = 0.0627$.

The data confirm Predictions 1 and 3 but disconfirm the critical Predictions 2 and 4, which are the most diagnostic. These results challenge the position coding account of position-specific prior list interference and, by extension, position-specific prior list intrusions.

Model Fits

We sought converging evidence on the four predictions by fitting versions of the position coding model that simulated the predictions in Figure 2 to the data from Experiments 3-12 to test hypotheses about prior list strength and the distance effects. The contrasts in the previous analyses are operational definitions of the strength (*sprior*) and distance (ρ) components of Equation 1. The fits measure these components directly as best-fitting model parameters and confirm the contrast results. The contrast analyses suggested that prior list strength equaled zero in Experiments 3-12. Position coding theory predicts strength greater than zero. We test this hypothesis by comparing the fits of models that fix *sprior* to 0 with models that allow it to vary freely. Position coding theory predicts models with *sprior* free to vary will fit better than models

with *sprior* fixed at 0. The fits with *sprior* free to vary provide estimates of prior list strength. Position coding theory predicts the estimates will be greater than 0. The fitting procedure is described in Appendix D.

We used Equation 1 to generate activation strengths for targets, within-list lures, and prior-list lures, and Equations 6 and 7 to generate drift rates for “yes” and “no” responses. The position similarity gradient ρ , prior list strength *sprior*, and capacity κ were each estimated as free parameters for each subject in each experiment. In addition, we estimated the thresholds for the two accumulators, a residual time parameter, and two scaling parameters for converting activations into drift rates. For each trial experienced by a subject, we used Equations 8 and 9 to compute the likelihood of making the response observed on that trial at the time observed on that trial. We found parameters for each subject in each experiment that maximized the total likelihood of the observed responses and RTs across all trials. As a result, models were fit to the complete joint distributions of correct and error responses and RTs in all conditions.

We fit two models. The first was a *nonzero prior list strength* model, representing the position coding model in Figure 1, which allowed *sprior* to vary between 0 and 1. The second was a *zero prior list strength* baseline model, which fixed *sprior* at 0 to eliminate prior list items from the model. A complete description of the models, their best-fitting parameter values, and their predictions for mean correct RT and error rate in each experiment are presented in Appendix D. The results of a parameter recovery analysis of the models are presented in Appendix E.

The model predictions across all 320 subjects are shown in the left (zero prior list strength) and right (nonzero prior list strength) panels of Figure 7 (dashed lines). The quality of the fits was about the same for the two models (see below) but the nonzero prior list strength model predicted a peak in RT and error rate at distance = 0 for prior-list lures that was not observed in the data (Figure 7). The zero prior strength model correctly predicted the observed flat function.

The four predictions (contrasts) of position coding theory are determined by the combination of prior list strength and distance parameters in Equation 1. Prediction 1 (-2 -1 1 2 distance effect in within-list lures) depends only on the distance parameter ρ in Equation 1. It was greater than zero in every subject, averaging 0.2508 in the zero prior list strength model fits

and 0.2826 in the nonzero prior list strength model fits (Table D1), confirming Prediction 1 in both models.

Predictions 2 and 4 (-2 -1 1 2 and -1 0 1 distance effects in prior list lures) depend on the product of distance and the prior list strength parameter *sprior* in Equation 1. Estimates of *sprior* were greater than zero on average (0.0938; Table D1), as predicted, but they were equal to zero in 169 of the 320 subjects, disconfirming the prediction for those subjects. The low *sprior* values reduce the activation of prior list lures, eliminating the distance effect and disconfirming Predictions 2 and 4.

Prediction 3 compares within-list distance effects, which depend only on the distance parameter, with prior-list distance effects, which depend on the product of the distance parameter and the prior list strength parameter. The relatively high value of the distance parameter accounts for the strong distance effects in within-list lures, but its effect in prior-list lures is diminished by the low value of the prior list strength parameter, predicting the difference in distance effects and confirming Prediction 3.

We tested the importance of the prior list strength parameter underlying predictions 2-4 by comparing the fit of the nonzero prior strength model, which includes the *sprior* parameter, with the fit of the zero prior strength model, which excludes it. The position coding model predicts the nonzero prior strength model will fit better because it allows values of *sprior* > 0. We compared the fit of the two models within each experiment and over all 320 subjects with paired sample *t* tests on four fit measures. AIC and BIC measure the likelihood of the data given the parameters, using different penalty terms for models with greater complexity (*t* tests for each experiment are in Table 6; mean goodness of fit values are in Table C2). Overall, AIC preferred the nonzero prior strength model (565.85) over the zero prior strength model (569.19), $t(319) = -2.3900$, $SEM = 1.3985$, $p = 0.0174$, $BF_{10} = 1.0366$, but the difference was significant only in Experiments 5, 8, and 9. Overall, BIC preferred the zero prior strength model (597.18) to the nonzero prior strength model (599.16), $t(319) = 2.4997$, $SEM = 0.7938$, $p = 0.0129$, $BF_{10} = 1.3457$. The preference for the zero prior strength model was significant in Experiments 4, 7, 8, 10, 11, and 12.

We calculated the squared correlation r^2 between observed and predicted RTs and error rates for each subject in each experiment. It measures the fit of the model to the pattern of the data and uses the same scale for RT and error rate (Table 6). Averaged over subjects and

experiments, the correlation between observed and predicted RTs was larger for the zero prior strength model (0.6758) than for the nonzero prior strength model (0.6696) but the difference was not significant overall, $t(319) = 1.5701$, $SEM = 0.0039$, $p = 0.1174$, $BF_{10} = 0.2117$, or in any experiment. The mean correlation between observed and predicted error rates was larger for the nonzero prior strength model (0.6916) than for the zero prior strength model (0.6797) overall, $t(254) = 3.2518$, $SEM = 0.0116$, $p = 0.0013$, $BF_{10} = 19.1846$, but it was significant only in Experiment 7.

Altogether, the model fits lead to the same conclusions as the contrast analyses of the mean RTs and error rates. They provide little support for the position coding predictions. Estimates of prior list strength were low overall and equal to zero for more than half the subjects. The position coding model with nonzero prior strength did not fit better than the baseline model with zero prior strength. The correlation analyses showed that the baseline model predicted the observed RTs and error rates as well as the more complex model.

The 53% of subjects with best-fitting prior strength values of zero falsify the position coding predictions, but the 47% with values greater than zero may provide some support. We separated the data for the two groups of subjects and plotted them in Figure 8. The $sprior = 0$ group showed no prior list distance effect. The (-1 0 1) distance contrast was not significant for RT, $t(168) = -1.4707$, $SEM = 11.6264$, $p = 0.1432$, $BF_{10} = 0.2469$, or for P(Error), $t(168) = -0.1561$, $SEM = 0.0079$, $p = 0.8762$, $BF_{10} = 0.0868$. However, the $sprior > 0$ group showed a little peak in prior list performance at distance = 0. The (-1 0 1) contrast was significant for RT, $t(150) = 2.4889$, $SEM = 16.4102$, $p = 0.0139$, $BF_{10} = 1.7855$, and for P(Error), $t(150) = 3.1028$, $SEM = 0.0116$, $p = 0.0023$, $BF_{10} = 8.8611$. We compared the magnitude of the (-1 0 1) contrast between groups and found it was significantly larger in the $sprior > 0$ group for both RT, $t(318) = 2.8811$, $SEM = 20.1114$, $p = 0.0042$, $BF_{10} = 6.3267$, and for P(Error), $t(318) = 2.6484$, $SEM = 0.0140$, $p = 0.0085$, $BF_{10} = 3.4546$. The $sprior > 0$ group provides some hope that the position coding account may explain position specific prior list interference, at least in some subjects. Position coding may be an individual difference or a strategy choice. Other approaches may be used by other subjects, either as an individual difference or a choice.

Summary

The cued recognition results (contrasts and model fits in Experiments 3-12) disconfirm Predictions 2 and 4 of the position coding account of position specific prior list interference. Predictions 1 and 3 were confirmed, but they are also consistent with theories of serial memory that do not assume position codes (e.g., item-dependent context theories). By extension, the cued recognition results challenge the position coding account of position specific prior list intrusions in recall tasks, which played a central role in the dominance of position coding theories of serial memory (Henson, 1998; Henson et al., 1996; Lewandowsky & Farrell, 2008). However, the serial recall results (contrasts in Experiments 1-2) confirm Predictions 1-4 of the position coding account of position specific prior list intrusions. Together, the cued recognition and serial recall results present a bigger challenge to position coding theories: They must change somehow to account for both the presence of position specific prior list intrusions in serial recall and the absence of position specific prior list interference in cued recognition.

Alternatives to position coding theories are challenged just as much by our results. Item-dependent theories do not assume position codes and so would predict the null effects of prior list distance we observed in cued recognition but they would also predict no position specific prior list intrusions in serial recall (Logan & Cox, 2023; Osth & Hurlstone, 2022), contrary to the results of Experiments 1 and 2. They too must change somehow to account the whole set of results.

We consider two ways to accommodate our results. First, we consider alternatives to our model of cued recognition that do not require focusing on a position to process prior list lures and so should not activate position codes that cause prior list interference. Then we consider ways to modify item-dependent context theories to produce prior list intrusions in serial recall.

Can Cued Recognition Be Done Without Focusing on Position?

The conclusions about position coding theory rest on the assumption that subjects evaluate prior list lures by using the position cue to retrieve the list item in the cued position and then comparing the retrieved item to the item in the probe (Logan et al., 2021, 2023b). The assumption implies that probing with the cue activates the position code for the cued position, which should produce position-specific prior list interference. The assumption may not be valid.

Subjects may use the item to retrieve position or to make recognition judgments without accessing position.

Using the Item to Retrieve a Position Code. Subjects could perform the cued recognition task with an “item-first” strategy that uses the item to retrieve a position code and then compares the retrieved position code to the one in the probe. This would make exactly the same predictions as our assumed “position-first” strategy that uses the position cue to retrieve an item because both strategies depend on the similarity between the probed position and the position of the item in the list. In the position-first strategy, positional similarity determines the activation of items at different distances from the cued position (Equation 1), and this determines RT and error rate produced by the decision process that compares the items (Equations 6 and 7). In the item-first strategy, positional similarity determines the comparison between the retrieved position and the cued position at different distances (Equation 1), and this determines RT and error rate produced by the decision process (Equations 6 and 7). Consequently, the item-first strategy makes the same predictions as the position-first strategy.

A more challenging possibility is that subjects may use the probe item in an *item recognition* process that compares the probe to all the items in the memory list without focusing on the cued position. The probe could be compared with each item in the memory list in parallel (Ratcliff, 1978) or with a composite representation formed by collapsing the memory list (e.g., by summing item vectors; Anderson, 1973). Neither case involves position information, so RT and error rate would not depend on activating position codes. Prior list distance effects would be null, as observed. Item recognition is a serious alternative that challenges the validity of using prior list lures to measure activation of position codes. We addressed it in five ways.

List Discrimination Strategy. First, we realized that the constrained lists in Experiments 3-6 and 11 allow a *list discrimination strategy*, in which subjects determine whether the probe came from the prior list and say “no” if it did. Items could not repeat from one trial to the next, so this strategy would produce correct “no” responses to prior list lures and predict RTs and error rates that were unaffected by distance. To address this strategy, we ran Experiments 7-10 and 12, replicating the original experiments with unconstrained lists, in which items could repeat from one trial to the next, so membership in the prior list was no longer a valid cue for a “no” response. The results replicated well with unconstrained lists, which disabled the list discrimination strategy. In the between-experiment comparisons of Experiments 3-10 (Tables

B1-B5), list type had no effect on RT or error rate, and no effect on any of the eight analyses assessing distance contrasts in RT and error rate except for the interaction between list type and probe delay in the within-list distance contrast (-2 -1 1 2) in error rate. None of the contrasts differed significantly between Experiment 11 (constrained lists) and Experiment 12 (unconstrained lists; see Table 6). We conclude that the list discrimination strategy was not an important factor in our cued recognition experiments, ruling out one possible item recognition strategy.

Item Recognition Followed by Position Cuing. Second, subjects could use an item recognition process to determine whether the probe item came from the current list and say “no” if it did not. Prior list probes were never present in the current list, so prior list probes could be rejected quickly without accessing position, producing null prior list distance effects. However, if the probe item was in the current list, the position-based cued recognition process would have to be engaged to distinguish matching probes from within-list lures. Subjects would have to focus on the cued position, retrieve the item, and compare it with the probe. This would increase their RTs for matches and within-list probes by an amount roughly equal to prior list lure RT minus motor execution time (Logan et al., 2023a). The data in Figures 4-6 show prolonged RTs to within-list lures, but match RTs were only 18 ms longer than prior-list lures despite wide variation in overall RT across experiments. The difference was significant, $t(31) = 3.0113$, $SEM = 5.8313$, $p = .0028$, $BF_{10} = 7.7859$, but small compared to the prolongation of RTs observed in sequential retrieval decisions in the “psychological refractory period” dual task procedure: Logan et al. (2023a) found a 395 ms prolongation in cued recall dual-task experiments (RT2 for SOA = 100 ms minus RT2 for SOA = 900) and Logan and Delheimer (2001) found a 602 ms prolongation in an item recognition dual-task experiment (with words; RT2 for SOA = 0 ms minus RT2 for SOA = 1000 ms; also see Carrier & Pashler, 1995).

Pre-Cue Effect. Third, we realized that the pre-cue effect distinguishes cued recognition from item recognition. Cued recognition requires focusing on the cued position in the memory list, and the pre-cue allows time to focus before the item part of the probe is presented. This reduces RT in the pre-cue condition relative to the no-pre-cue condition (Logan et al., 2023b). Thus, decisions based on cued recognition should be shorter with a pre-cue than without one. Item recognition does not require focusing on the cued position and so would not benefit from pre-cuing the position. Item recognition can begin only after the item part of the

probe is presented, at the end of the pre-cue delay. RT is measured from the onset of the item part of the probe, so RTs for decisions based on item recognition should be unaffected by pre-cuing. A second, related, prediction is that pre-cues should speed up correct recognition of targets (based on cued recognition) while leaving correct rejection of prior list lures (based on item recognition) unaffected. The difference between RTs to prior list lures and RTs to targets should be larger in experiments with pre-cues than in experiments without pre-cues. On the other hand, if all responses are based on cued recognition, then the pre-cue should speed both “yes” and “no” responses by the same amount. The difference between RTs to prior list lures and RTs to targets should be the same in experiments with and without a pre-cue. The data, plotted in Figures 4 and 5, are more consistent with cued recognition.

We tested the first prediction by comparing prior list lure RTs from experiments with pre-cued probes (4, 6, 8, and 10) and experiments with simultaneous probes (3, 5, 7, and 9). Prior list lure RTs were 292 ms shorter with pre-cued probes than with simultaneous probes, and the difference was significant, $t(254) = 11.0413$, $SEM = 26.4097$, $p < 0.0001$, $BF_{10} > 1000$, disconfirming the item recognition prediction. We tested the second prediction by comparing the difference between prior list lure RT and “yes” RT in experiments with (4, 6, 8, and 10) and without pre-cues (3, 5, 7, and 9). The 24 ms difference of differences only approached significance, $t(254) = 1.9221$, $SEM = 12.2933$, $p = 0.0557$, $BF_{10} = 0.7838$, failing to provide clear support for the item recognition prediction. The results of both comparisons are consistent with our assumption that subjects use cued recognition to evaluate prior list lures.

Cue Type Effect. Fourth, we realized that the same logic applies to the effect of cue type on RT and leads to similar predictions. Cued recognition requires accessing position information in the probe but item recognition does not. Cued recognition RTs will be faster when position information is easier to extract from the cue (spatial cues) than when it is harder (numeric cues). Item recognition does not require position information, so item recognition RT should be unaffected by cue type. We tested this prediction by comparing RTs to prior list probes from experiments with spatial cues (3, 4, 7, 8) with RTs from experiments with numeric cues (5, 6, 9, 10). The difference was 163 ms. It was highly significant, $t(254) = 5.3686$, $SEM = 30.4472$, $p < .0001$, $BF_{10} > 1000$, consistent with our assumption that prior list lures were processed with cued recognition. We tested a second prediction, that “yes” RTs should vary with cue type (because they depend on position) but prior list probe RTs should not (because

they do not depend on position), by comparing the difference between “yes” and prior list probe RTs in experiments with spatial cues (3, 4, 7, 8) and with the difference in experiments with numeric cues (5, 6, 9, 10). The difference of differences was 7 ms, which was not significant, $t(254) = 0.5611$, $SEM = 12.3747$, $p = .5752$, $BF_{10} = 0.1591$, suggesting that prior list lures were processed with cued recognition.

Model Fits. Finally, we used model fits to test the importance of including item recognition in the decision process, comparing models that included item recognition in the decision process with models that did not include it. We assumed that item recognition was not position specific and modeled it by comparing the probe item to each item in the list (following Logan et al., 2021). We assumed this version of item recognition went on in parallel with cued recognition (following Logan et al., 2021). We added value of the item match to the drift rate in the decision process with weight w . The contribution from cued recognition was given weight $1 - w$, so the evidence for “yes” is

$$T_{yes} = (1 - w)(\mathbf{q} \cdot \mathbf{m}_k) + w(\sum_{j=1}^6 \mathbf{q} \cdot \mathbf{m}_j) \quad (10)$$

and the evidence for “no” is

$$T_{no} = (1 - w)\|\mathbf{m}_k\| + w(\sum_{j=1}^6 \|\mathbf{m}_j\|) \quad (11)$$

where k is the cued position in the list. The drift rate for “yes” is

$$v_{yes} = \frac{T_{yes}}{1 + \lambda \kappa T_{no}} \quad (12)$$

and the drift rate for “no” drift becomes

$$v_{no} = \frac{\lambda T_{no}}{1 + \kappa T_{yes}}. \quad (13)$$

We fit two versions of this model. One implemented item recognition in the position coding model with nonzero prior list strength. The other implemented item recognition in the model with zero prior list strength. Values of the best fitting parameters, measures of goodness of fit, and predicted RTs and error rates for each experiment are presented in Appendix D. The mean predicted and observed values across the 320 subjects are presented in Figure 9.

We calculated contrasts comparing goodness of fit measures in models with and without item recognition and got mixed results. Table 7 contains the values for zero prior list length models with and without item recognition within each experiment and over all 320 subjects. AIC favored models without item recognition in four of the 10 experiments but the overall difference was not significant. BIC favored models with without item recognition in five

experiments but the overall difference was not significant. The correlations with RT were significantly higher in models with item recognition in eight experiments and overall. The correlations with error rate were not significantly higher in models with item recognition in any experiment or overall.

Table 8 contains the contrasts comparing goodness of fit measures for nonzero prior list strength models with and without item recognition within each experiment and over all 320 subjects. AIC favored models with item recognition in three experiments but the difference was not significant overall. BIC favored models with item recognition in three experiments but the difference was not significant overall. Correlations with RT were larger with item recognition than without in four experiments and the difference was significant overall. Correlations with error rate were smaller with item recognition than without in one experiment and the difference was significant overall.

In summary, item recognition does not improve the fit of the zero prior list strength model or the nonzero prior list strength model, as measured with AIC and BIC. The correlations with RT improved by adding item recognition, but the correlations with error rate either did not change or reversed. These results converge on the conclusions from the analyses of list discrimination, pre-cue delay, and cue type. They suggest that item recognition is not a viable explanation the results that challenge position coding theory.

Other Accounts of Position-Specific Prior List Intrusions

Taken by themselves, the results of the cued recognition experiments (3-12) support item-dependent context theories of serial memory, which would not predict position specific prior list interference (Botvinick & Plaut, 2006; Lewandowsky & Murdock, 1989; Logan, 2021; Murdock, 1995; Solway et al., 2012). Item-dependent context theories account for prior list intrusions that are semantically related to items in the current list (Loess, 1967; Wickens, 1970) and intrusions that follow an item that repeats from the prior list, intruding the item that followed the repeated item on the prior list (Fischer-Baum & McCloskey, 2015; Kahana et al., 2002; Zaromb et al., 2006). However, item-dependent context theories do not account for the position specific prior list intrusions observed in the serial recall (Experiments 1-2) and cued recall (Logan et al., 2023a) experiments in this article and many others (Conrad, 1959; Henson, 1998; Melton & Von Lackum, 1941; Osth & Dennis, 2015; but see Caplan et al., 2022; Dennis, 2009;

Logan & Cox, 2023). Taken together, our results on serial recall and cued recognition challenge item-dependent context theories as much as position coding theories. Both have to explain why position specific prior list intrusions occur in serial recall and why position specific prior list interference does not occur in cued recognition. As a first step, we consider ways to produce position specific prior list intrusions in item-dependent context models.

Changing Memory Theories. One way to accommodate our results is to develop accounts of position-specific prior list intrusions that do not assume position coding or do not attribute them to serial memory. Accounts of prior list intrusions must assume that items in the prior list are activated at retrieval time along with items in the current list. They must also assume the prior list is activated less than the current list, or else prior list intrusions would dominate correct retrievals. Neither of these assumptions require position coding. In principle, they could be implemented in any theory of serial order, including item-dependent context theories. Accounts of position-specific prior list intrusions must also assume that prior list activation is position specific. The prior list item that is activated most strongly is the one in the list position that is the focus of retrieval in the current list. Position coding theories account for this position specificity in their fundamental assumption that items are associated with position codes and their equally fundamental assumption about activation and distance (Equation 1). It is less clear how item-dependent context theories would account for it.

Osth and Hurlstone (2023) showed important constraints on the ability of item-dependent context theories to produce prior list intrusions, analyzing the Context Retrieval and Updating (CRU) model of serial recall (Logan, 2021). They modified CRU to represent the prior list as well as the current one, and they manipulated similarity between the list contexts to produce intrusions. CRU made position-specific prior list intrusions when list contexts were sufficiently similar, but it did so by switching to the prior list and reporting prior-list items from the intrusion onward. Subjects typically make one prior list intrusion and then return to the current list. There was only one trial in Experiment 1 and one trial in Experiment 2 in which a subject recalled the prior list entirely. We confirmed Osth and Hurlstone's results with our own simulations (Logan & Cox, 2023). List similarity by itself does not seem to be the answer (cf. Dennis, 2009).

Caplan et al. (2022) showed that a simple modification of a classical chaining model could produce position-specific prior list intrusions. The model assumes that the current list *ABCDEF* and the prior list *ghijkl* are both represented as associative chains that link adjacent

items. Retrieval is initiated by activating a start element that is associated with the first item in each chain. The start element is associated more strongly with the current list than with the prior list, so *A* is more likely to be retrieved than *g*. A prior list intrusion occurs when *g* is retrieved instead of *A*. The model prevents perseverating on the prior list (like CRU does) by using the *evidence* retrieved in the decision process as the cue for the next item instead of the *item* retrieved (Lewandowsky & Li, 1994). The evidence for *A* and *g* retrieves evidence for *B* and *h*, continuing both chains. The evidence for *A* is stronger than the evidence for *g* despite *g* winning the competition, so *B* tends to be retrieved next, getting the model back on track. If the retrieved item *g* is used to cue retrieval instead of the evidence that drove retrieval, *h* is more likely to be retrieved than *B*, and the model will perseverate on the prior list, like CRU but unlike human subjects. Caplan et al. (2022) showed that their model could fit position-specific prior list intrusion data, making it an attractive alternative to position coding. They viewed their model as promising but preliminary, as they had not yet fitted it to the range of data and effects in the serial recall literature.

Logan and Cox (2023) tried the Caplan et al. (2022) idea with CRU and found that it showed promise. Using the evidence that drove retrieval instead of the item retrieved to update context, they were able to produce position-specific prior list intrusions but the model still tended to perseverate on the prior list. They developed a version of CRU that updated context with the retrieved item on some trials and with the evidence that drove retrieval on other trials. The updating was adaptive, using the item when it was likely that the retrieved item was correct and using the evidence when it was likely that the retrieved item was an error. This produced position-specific prior list intrusions without perseverating (as much) on the prior list, more like human subjects. However, the simulations are proofs of concept at best, and the changes to the model (adding an error detection component that assesses the likelihood of an error) are extensive, so the extension of CRU requires further investigation.

These models suggest it may be possible to account for position-specific prior list intrusions without position coding, but they are challenged by our experimental results as much as position coding theories are. They must also explain why position-specific prior list interference does not appear in cued recognition.

Intruded Responses, Not Memories. The changes to the models we proposed are based on the assumption that prior list intrusions are produced by retrieval from memory. Position

coding accounts are based on the same assumption. An alternative possibility is that prior list intrusions are produced in the output processes required in recall tasks rather than in the memory system itself. Serial recall requires a sequence of actions to report each of the items (keystrokes in our experiments), and the order of the sequence is controlled by a *motor program* (Keele, 1968; Logan, 2018). Prior list intrusions may result from position coding in the motor program instead of position coding in memory. The motor program might associate keystrokes with positions (“first press the *A* key, second press the *B* key,” etc.) and position-specific intrusions might occur if the motor program used to report the previous list was still available (e.g., “first press the *g* key, second press the *h* key,” etc.) and the keystroke from the prior list wins the competition. This would explain why prior list intrusions occur in serial recall and why prior list interference does not occur in cued recognition. Cued recognition requires a simpler motor program that conveys a single judgment about the probe item, not a sequence of judgments about the identity of every item. There is only one “position” in this motor program, so there is less opportunity for confusion.

The motor program account predicts position specific prior list intrusions in tasks in which subjects execute the motor program without retrieving items from memory. Copy typing is one such task, as it involves reporting a continuously visible list so the information required to respond is available perceptually. Logan (2021) compared copy typing, serial recall, and perceptual report of 5, 6, and 7 letter consonant strings. Each task required the same motor program (typing the letters in order) but varied in its requirements for memory and perception. Copy typing required the motor program but not memory. The motor programming account predicts position specific prior list intrusions in the copy typing task.

We tested this prediction by searching for prior list intrusions in Logan’s (2021) data. The subjects were 24 skilled typists who typed 192 lists in each condition. Serial recall, whole report, and copy typing were run in separate blocks. We identified intrusions in each task and determined whether they came from the previous list. If they did, we calculated the distance between their position in the prior list and their position in the current response. This was complicated by the random variation in list length within blocks. Subjects could encode position from the beginning or the end of the list (Henson, 1998; Fischer-Baum & McCloskey, 2015). Following precedent (Fischer-Baum & McCloskey, 2015), we calculated distance from the beginning of the list and distance from the end of the list and chose the shorter distance. We

summed the frequencies of prior list intrusions at each distance in each task across list lengths and subjects. The frequencies for each task are plotted as a function of distance in Figure 10 (left).

Prior list intrusions were more frequent in report than in recall and least frequent in typing (1979, 1337, and 173 intrusions, respectively), reflecting large differences in overall accuracy, so we replotted the data as the proportion of the total number of prior list intrusions in each task at each distance (Figure 10, right). Both plots show position specific prior list intrusions in copy typing. There is a peak at distance = 0, confirming Prediction 4 of position coding theory (the critical contrast comparing distances -1 0 1). This is consistent with the motor program account, in which position specific prior list intrusions are the product of position coding in the motor program.

Serial recall and whole report also showed position specific prior list intrusions with sharp peaks at distance = 0. The similarity of the proportions in Figure 10 (left) invites the conclusion that the motor system uses position coding but the memory system does not, but we cannot rule out the possibility that the memory system also uses position coding. The difference in frequency between memory and typing could reflect additional memory-based intrusions. We note as well that the motor programming account does not explain prior list intrusions in cued recall (Figure 3; Logan et al., 2023a), where the motor program specifies only one response. At this point, the results invite speculation, not strong conclusions, but the speculation is intriguing. Understanding the role of the motor system and output and decision processes more generally is an important goal for future research (Dendauw et al., 2024).

Implications for Position Coding Theories

The results of our experiments challenge the core assumptions of the position coding account of position-specific prior list intrusions. We showed in theoretical analysis and in simulations that position coding theories that predict position-specific prior list intrusions in recall must also predict position-specific prior list interference in cued recognition. Contrary to this prediction, we failed to find position-specific prior list interference in 10 experiments.

Our experiments challenge the core assumption that prior list items are associated with the same position codes as current list items with lower strength (Figure 1 top row). In recall, the core assumption implies that prior list items can compete with current list items and it predicts

intrusions when prior list items win the competition (Figure 1 second row; Figure 2 top row). The prior list intrusions observed in Experiments 1 and 2 and many others are interpreted as confirming this prediction. In cued recognition, we showed that the assumption predicts interference when prior list items are used as lures (Figure 1 bottom; Figure 2) and we failed to observe such interference (Figures 4-6). The contrast testing the predicted peak in RT and error rate was not significant and the *sprior* parameter that represents prior list strength was 0 in 53% of the subjects and close to 0 in the other 47%.

The challenge to the position coding account of position specific prior list intrusions has broader implications for position coding theories. A major impetus for the development of position coding theories was their ability to account for four error phenomena that classical chaining theories could not explain: recovery from errors, transpositions to earlier list positions, phonological confusion effects, and position-specific prior list intrusions (Henson et al., 1996). Previous research has shown that theories based on different assumptions can account for the first three phenomena (Botvinick & Plaut, 2006; Logan, 2021; Osth & Hurlstone, 2023; Solway et al., 2012), leaving prior list intrusions as the last of the four unique predictions that support position coding theories. Our experiments falsify this prediction when it is extended to cued recognition, leaving position coding models with no unique predictions that distinguish them from other theories of serial memory. This challenges the dominance of position coding theories and encourages renewed attention to other theories and different approaches.

Implications for Serial Memory

The dissociation between position-specific prior list intrusions and position-specific prior list interference challenges all theories of serial memory, not just position coding theories. The theories that account for intrusions must explain why there is no interference from prior-list lures in cued recognition. The theories that account for the lack of interference must explain why intrusions occur in serial and cued recall. We hope our results encourage theorists of all persuasions to rise to the challenge.

Our results highlight the importance of using different retrieval tasks to test assumptions about memory representations (Hintzman, 2011). Most of the work on serial memory has focused on serial recall to the exclusion of other retrieval tasks (Hurlstone et al., 2014). Our results show that the retrieval task matters. The predictions derived from the memory

representations may be the same, but the results differ substantially. Serial recall shows evidence of position-specific prior list activation (Figure 3). Cued recognition does not (Figures 4-6). These results underscore the important point that memory performance is a joint function of the representations and the processes that operate on them (Anderson, 1978; Atkinson & Shiffrin, 1968). Representation and process are confounded in a single task, like serial recall or cued recognition. Their effects can be separated by using different retrieval tasks to access the same representation (e.g., Cox et al., 2018). This has been a productive strategy in global theories of memory, explicating the relations between recognition and recall (Anderson et al., 1998; Gillund & Shiffrin, 1984; Humphreys et al., 1989; Murdock, 1982, 1983). It should also be productive in theories of serial memory and further the goal of integrating those theories with other memory theories (Ward et al., 2010; Ward & Tan, 2023).

Our results show the benefits of using cued recognition to probe serial memory. The position cue requires attention to order information. The item cue probes the state of the memory system, and different cues can be chosen to probe different states. Carefully designed lures have led to important insights into the nature of false recognition (Shiffrin, Huber, & Marinelli, 1995), correct rejection of novel lures (Mewhort & Johns, 2000; Osth et al., 2023), the relationship between item and associative information (Cox & Criss, 2017), the organization of lexical memory (Grainger, 2018), and the organization of semantic memory (Ratcliff & McKoon, 1988; Zbrodoff, 1999). In cued recognition, lures probe the state of current and prior lists at different distances from the cued position, measuring their activation to test theories of serial order.

Strategies and Models of Memory

Our experiments challenge the position coding account of prior list intrusions, which distinguishes them from other theories of serial memory (Henson et al., 1998; Osth & Hurlstone, 2023). The theories are no longer so distinct. This may be frustrating from the usual perspective on modeling, where models are treated as mutually exclusive and the goal is to find the one that fits best, declare it the winner, and discard the rest. Mimicry makes models harder to distinguish. But mimicry can be beneficial. Sometimes it reveals a basic truth in all models that account for the same phenomenon (Anderson, 1978). For serial memory, the basic truth is the exponential distance gradient $a(i|j) = \rho^{|i-j|}$ in Equation 1 that represents the similarity of position codes (Logan

& Cox, 2021; Murdock, 1997). It appears in models that represent order as associations of items to contexts, whether the contexts are independent of (Lewandowsky & Farrell, 2008) or dependent on the items (Logan, 2021). It is the key assumption that allows them to account for a broad range of order phenomena.

One way to deal with mimicry is to embrace it, accepting that different models may fit the data equally well and using other criteria to distinguish models. We might choose models based on how clearly they relate theoretical constructs to observable behavior in a specific domain (Navarro, 2019; Singmann et al., 2022). If the task requires memory for positions, we might choose a position coding model because it provides a clear and direct way to relate the theory to the experiment, not because it provides a better fit (assuming the fits are equivalent). We might also choose models based on the questions they allow us to ask and use them to test hypotheses. We tested hypotheses about the interaction between distance and prior list strength by comparing different versions of the position coding model. Theoretical analyses and simulations of the models allowed us to derive four core predictions about the data. The fits allowed us to measure the distance gradient and prior list strength directly as model parameters (ρ and *sprior*).

Another way to deal with mimicry is to treat models as alternative strategies subjects might employ to represent order instead of different candidates for the One True Model. Different subjects may choose different models for the same task, like our subjects with prior strength = 0 and prior strength > 0. The same subject may choose different models for different tasks or choose different models at different times on the same task (Logan & Cox, 2023). Different models may be better suited to different tasks. We represent position explicitly when keeping track of standings in sports, songs on the hit parade, and birth order of siblings. We represent order with reference to context in understanding events and biographies. These possibilities are enticing, inspiring broader multiple-representation theories of memory and new research on the determinants and consequences of strategy choice and the control processes that make the choices. Mimicry will make it harder to distinguish between strategies in particular cases, but much can be learned from the core assumptions, like the similarity gradient in Equation 1, that all models share.

The call to consider models as strategies emphasizes the role of processing as much as representation. Processing is required to form a memory representation and to extract

information from it at retrieval (e.g., Craik, 2020). Processing is required to control the encoding and retrieval processes, directing them to the relevant parts of the memory representation and controlling the order and timing of their execution (Atkinson & Shiffrin, 1968; Logan et al., 2023a, 2023b). The emphasis on processing raises the value of experimental procedures that allow the processes to be measured directly with RT as well as accuracy, to take advantage of the many process models that predict both measures (e.g., Ratcliff, 1978; Tillman et al., 2020; Usher & McClelland, 2001) and develop more complete models of memory. The emphasis on processing is essential in distinguishing among memory representations. The behavior we measure is the result of processes operating on representations. To make inferences about representation from behavior, we must unravel the interactions between representations and processes, and that requires understanding the processes. Our research has attempted to do that.

Limitations

Our 10 cued recognition experiments are close variations on a common design. They all used lists of six consonants drawn from a pool of 20, presented briefly, and tested after a short 1000 ms retention interval. We chose to vary cue type and cue delay between experiments because of our interest in the relation between memory and attention tasks that manipulate cues in the same way (Logan et al., 2021, 2023a, 2023b). The close variations demonstrate the replicability of the results but they do not demonstrate their generality. It is possible that our results would not replicate with a broader range of materials and more variation in experimental design.

It would be worthwhile adapting our procedure to word lists, which are common materials in studies of interference, and varying the size of the pool of items (Osth & Hurlstone, 2023) and list length over a broad range (Ward et al., 2010; Ward & Tan, 2023). It would also be worthwhile adapting our procedure to simple visual stimuli like color patches or oriented gratings, or to pictures, which are common in studies of visual memory. If our results replicated across these variations, our conclusions would be much stronger.

Our major result, the contrast between strong position-specific intrusions in serial recall and null position-specific interference in cued recognition, was tested between subjects. Each set of subjects performed only one task, and it is possible that they represented serial order differently in ways that were tailored to the tasks they performed. Possibly, subjects let item

activation decay more rapidly in cued recognition, and that produced the null results (if they can control decay, which is not clear). It would be worthwhile replicating our experiments but mixing serial recall and cued recognition randomly and post-cuing the task so the lists would be represented in the same way.

The major result of Experiments 3-12 is the null effect of position-specific prior list interference. We found no evidence of such interference. Indeed, we found no evidence of any kind of prior list interference. Our focus on the predictions of position coding theory led us to an experimental design that maximized the number of targets, within-list lures, and lures from the immediately prior list at each distance (-2 -1 0 1 2). A different control condition is required to demonstrate prior list interference that is not position specific. Items from the immediately prior list would have to be compared with novel items or items from earlier lists. There is a large literature demonstrating such interference in item recognition (Badre & Wagner, 2005; Jonides et al., 1998; McElree & Doshier, 1989; Monsell, 1978; for a review, see Jonides & Nee, 2006). Similar interference might occur in cued recognition.

Failing to find prior list interference in cued recognition could mean that the task does not produce interference. That could be true as an empirical observation, but it would raise the important theoretical question, why not? Why should cued recognition show no prior list interference? Prior list activation may decay faster in cued recognition, but why should that happen? These are the same questions we raise about position specific prior list interference and they would require similar answers. The answers are important and worth obtaining for what they will reveal about the nature of recognition memory and the nature of interference, expanding the insights gained from understanding the lack of position-specific prior list interference.

Finding or failing to find prior list interference that is not position specific is not directly relevant to the specific question that motivated our experiments. We were interested in position-specific prior list interference. We showed that the assumptions shared by all position coding theories predict that prior list items should be activated in proportion to their distance from the current focus of retrieval, and we showed that a plausible decision process that is typical of the literature predicts longer RTs and higher error rates at shorter distances. Whether that particular kind of prior list interference would occur was the question, and that question does not require the existence of any other kind of prior list interference. Indeed, the only kind of prior list

interference predicted by position coding theories is position specific. They say nothing about other kinds of interference.

Conclusions

The ability to predict position specific prior list intrusions has led to the dominance of position coding theories in serial memory. We showed that the assumptions that allow position coding theories to predict position specific prior list intrusions in serial recall also predict position specific interference from prior list lures in cued recognition. We found no such interference in 10 experiments, falsifying the prediction. This challenges the position coding account of position specific prior list intrusions and, by extension, challenges their dominance in research on serial memory. The cued recognition results are consistent with alternatives to position coding theories, which do not assume position codes.

We ran two serial recall experiments that used the same lists and presentation conditions as the cued recognition experiments and found position specific prior list intrusions in both of them, consistent with position coding theories and inconsistent with the alternatives. Together, the results of our cued recognition and serial recall experiments challenge all theories of serial memory, whether or not they assume position coding. All theories must explain why prior list intrusions are position specific while prior list interference is not.

References

- Anderson, J. A. (1973). A theory for the recognition of items from short memorized lists. *Psychological Review*, 80(6), 417–438.
- Anderson, J. R. (1978). Arguments concerning representations for mental imagery. *Psychological Review*, 85(4), 249–277.
- Anderson, J. R., Bothell, D., Lebiere, C., & Matessa, M. (1998). An integrated theory of list memory. *Journal of Memory and Language*, 38, 341–380.
- Anderson, J. R., & Matessa, M. (1997). A production system theory of serial memory. *Psychological Review*, 104(4), 728–748.
- Atkinson, R. C., & Shiffrin, R. M. (1968). Human memory: A proposed system and its control processes. In K. W. Spence & J. T. Spence (Eds.), *The Psychology of Learning and Motivation: Advances in Research and Theory* (Vol. 2, pp. 89–195). Academic Press.
- Badre, D., & Wagner, A. D. (2005). Frontal lobe mechanisms that resolve proactive interference. *Cerebral Cortex*, 15, 2003–2012.
- Benjamin, A. S. (2010). Representational explanations of "process" dissociations in recognition: The DRYAD theory of aging and memory judgments. *Psychological Review*, 117 (4), 1055–1079.
- Botvinick, M. M., & Plaut, D. C. (2006). Short-term memory for serial order: A recurrent neural network model. *Psychological Review*, 113(2), 201–233.
- Brown, G. D. A., Neath, I., & Chater, N. (2007). A temporal ratio model of memory. *Psychological Review*, 114(3), 539–576.
- Brown, G. D. A., Preece, T., & Hulme, C. (2000). Oscillator-based memory for serial order. *Psychological Review*, 107(1), 127–181.
- Burgess, N., & Hitch, G. J. (1999). Memory for serial order: A network model of the phonological loop and its timing. *Psychological Review*, 106(3), 551–581.
- Caplan, J. B., Shafaghat Ardebili, A., & Liu, Y. S. (2022). Chaining models of serial recall can produce positional errors. *Journal of Mathematical Psychology*, 109, Article 102677.
- Carandini, M., & Heeger, D. J. (2012). Normalization as a canonical neural computation. *Nature Reviews Neuroscience*, 13, 51–62.
- Carrier, L. M., & Pashler, H. (1995). Attentional limits in memory retrieval. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 21(5), 1339–1348.

- Clark, S. E., & Shiffrin, R. M. (1987). Recognition of multiple-item probes. *Memory & Cognition*, 15 (5), 367–378.
- Conrad, R. (1959). Errors of immediate memory. *British Journal of Psychology*, 50, 349-359.
- Cox, G. E., & Criss, A. H. (2020). Similarity leads to correlated processing: A dynamic model of encoding and recognition of episodic associations. *Psychological Review*, 127 (5), 792–828.
- Cox, G. E., Hemmer, P., Aue, W. R., & Criss, A. H. (2018). Information and processes underlying semantic and episodic memory across tasks, items, and individuals. *Journal of Experimental Psychology: General*, 147 (4), 545–590.
- Cox, G. E., & Shiffrin, R. M. (2017). A dynamic approach to recognition memory. *Psychological Review*, 124(6), 795–860.
- Craik, F. I. M. (2020). Remembering: An activity of mind and brain. *Annual Review of Psychology*, 71, 1-24.
- De Leeuw, J. R. (2015). jsPsych: A JavaScript library for creating behavioral experiments in a Web browser. *Behavior Research Methods*, 47(1), 1-12.
- Dendauw, E., Evans, N. J., Logan, G. D., Haffen, E., Bennabi, D., Gajdos, T., & Servant, M. (2024). The gated cascade diffusion model: An integrated theory of decision-making, motor preparation, and motor execution. *Psychological Review*, in press.
- Dennis, S. (2009). Can a chaining model account for serial recall? *Proceedings of the Annual Meeting of the Cognitive Science Society*, 31, 2813-2818.
- Doshier, B. A., & Rosedale, G. (1989). Integrated retrieval cues as a mechanism for priming in retrieval from memory. *Journal of Experimental Psychology: General*, 118 (2), 191–211.
- Doshier, B. A., & Rosedale, G. (1997). Configural processing in memory retrieval: Multiple cues and ensemble representations. *Cognitive Psychology*, 33, 209–265.
- Ebbinghaus, H. (1885). *Über das gedächtnis: Untersuchungen zur Experimentellen Psychologie*. Duncker.
- Ebenholtz, S. M. (1963). Serial learning: Position learning and sequential associations. *Journal of Experimental Psychology*, 66, 353–362
- Farrell, S. (2012). Temporal clustering and sequencing in short-term memory and episodic memory. *Psychological Review*, 119(2), 223-271.

- Fisher-Baum, S. & McCloskey, M. (2015). Representation of item position in immediate serial recall: Evidence from intrusion errors. *Journal of Experimental Psychology: Learning, Memory and Cognition*, 41, 1426-1446.
- Gillund, G., & Shiffrin, R. M. (1984). A retrieval model for both recognition and recall. *Psychological Review*, 91(1), 1–67.
- Grainger, J. (2018). Orthographic processing: A ‘mid-level’ vision of reading. *Quarterly Journal of Experimental Psychology*, 71, 335-359.
- Hartley, T., Hurlstone, M. J., & Hitch, G. J. (2016). Effects of rhythm on memory for spoken sequences: A model and tests of its stimulus-driven mechanism. *Cognitive Psychology*, 87, 135-178.
- Healey, M. K. and Kahana, M. J. (2014). Is memory search governed by universal principles or idiosyncratic strategies? *Journal of Experimental Psychology: General*, 143 (2), 575–596.
- Henson, R. N. (1998). Short-term memory for serial order: The Start-End model. *Cognitive Psychology*, 36(2), 73-137.
- Henson, R. N. A., Norris, D. G., Page, M. P. A., & Baddeley, A. D. (1996). Unchained memory: Error patterns rule out chaining models of immediate serial recall. *Quarterly Journal of Experimental Psychology: Human Experimental Psychology*, 49(1), 80 –115.
- Hintzman, D. L. (1984). MINERVA 2: A simulation model of human memory. *Behavior Research Methods, Instruments & Computers*, 16, 96-101.
- Hintzman, D. L. (1988). Judgments of frequency and recognition memory in a multiple-trace memory model. *Psychological Review*, 95(4), 528–551.
- Hintzman, D. L. (2011). Research strategy in the study of memory: Fads, fallacies, and the search for the "coordinates of truth". *Perspectives on Psychological Science*, 6 (3), 253–271.
- Hull, C. L. (1932). The goal-gradient hypothesis and maze learning. *Psychological Review*, 39(1), 25–43.
- Hull, C. L. (1934). The concept of the habit-family hierarchy and maze learning: Part I. *Psychological Review*, 41(1), 33–54.
- Humphreys, M. S. (1978). Item and relational information: A case for context independent retrieval. *Journal of Verbal Learning and Verbal Behavior*, 17, 175–187.

- Humphreys, M. S., Bain, J. D., & Pike, R. (1989). Different ways to cue a coherent memory system: A theory for episodic, semantic, and procedural tasks. *Psychological Review*, 96(2), 208–233.
- Hurlstone, M. J., Hitch, G. J., & Baddeley, A. D. (2014). Memory for serial order across domains: An overview of the literature and directions for future research. *Psychological Bulletin*, 140, 339–373.
- Jonides, J., & Nee, D. E. (2006). Brain mechanisms of proactive interference in working memory. *Neuroscience*, 139, 181–193.
- Jonides, J., Smith, E. E., Marschuetz, C., Koeppe, R. A., & Reuter-Lorenz, P. A. (1998). Inhibition in verbal working memory revealed by brain activation. *Proceedings of the National Academy of Sciences USA*, 95, 8410–8413.
- Kahana, M. J., Howard, M. W., Zaromb, F., & Wingfield, A. (2002). Age dissociates recency and lag recency effects in free recall. *Journal of Experimental Psychology: Learning, Memory and Cognition*, 28, 530–540.
- Keele, S. W. (1968). Movement control in skilled motor performance. *Psychological Bulletin*, 70, 387–403.
- Ladd, G. T., & Woodworth, R. S. (1911). *Elements of physiological psychology: A treatise of the activities and nature of the mind from the physical and experimental point of view*. Charles Scribner's Sons.
- Lashley, K. S. (1951). The problem of serial order. In L. A. Jeffress (Ed.), *Cerebral mechanisms in behavior* (pp. 112–136). Wiley.
- Lewandowsky, S., & Farrell, S. (2008). Short-term memory: New data and a model. In B. H. Ross (Ed.), *The psychology of learning and motivation* (Vol. 49, pp. 1–48). Academic Press.
- Lewandowsky, S., & Li, S.-C. (1994). Memory for serial order revisited. *Psychological Review*, 101(3), 539–543.
- Lewandowsky, S., & Murdock, B. B., Jr. (1989). Memory for serial order. *Psychological Review*, 96(1), 25–57.
- Loess, H. (1967). Short-term memory, word class, and sequence of items. *Journal of Experimental Psychology*, 74(4, Pt.1), 556–561.

- Lo, C.-C., & Wang, X.-J. (2006). Cortico-basal ganglia circuit mechanism for a decision threshold in reaction time tasks. *Nature Neuroscience*, 9, 956-963.
- Logan, G. D. (2002). An instance theory of attention and memory. *Psychological Review*, 109, 376-400.
- Logan, G. D. (2018). Automatic control: How experts act without thinking. *Psychological Review*, 125(4), 453–485.
- Logan, G. D. (2021). Serial order in perception, memory, and action. *Psychological Review*, 128(1), 1–44.
- Logan, G. D., & Cox, G. E. (2021). Serial memory: Putting chains and position codes in context. *Psychological Review*, 128(6), 1197-1205.
- Logan, G. D., & Cox, G. E. (2023). Serial order depends on item-dependent and item-independent contexts. *Psychological Review*.
- Logan, G. D., Cox, G. E., Annis, J., & Lindsey, D. R. B. (2021). The episodic flanker effect: Memory retrieval as attention turned inward. *Psychological Review*, 128(3), 397–445.
- Logan, G. D., & Delheimer, J. A. (2001). Parallel memory retrieval in dual-task situations: II. Episodic memory. *Journal of Experimental Psychology: Learning, Memory and Cognition*, 27, 668-685.
- Logan, G. D., Lilburn, S. D., & Ulrich, J. E. (2023a). Serial attention to serial memory: The psychological refractory period in forward and backward cued recall. *Cognitive Psychology*, 145, 101583.
- Logan, G. D., Lilburn, S. D., & Ulrich, J. E. (2023b). The spotlight turned inward: The time-course of focusing attention on memory. *Psychonomic Bulletin & Review*, 30, 1028-1040.
- McElree, B., & Doshier, B. A. (1989). Serial position and set size in short-term memory: The time course of recognition. *Journal of Experimental Psychology: General*, 118, 346-373.
- Melton, A. W., & Von Lackum, W. J. (1941). Retroactive and proactive inhibition in retention: Evidence for a two-factor theory of retroactive inhibition. *American Journal of Psychology*, 54, 157-173.
- Mewhort, D. J. K., & Johns, E. E. (2000). The extralist-feature effect: Evidence against item matching in short-term recognition memory. *Journal of Experimental Psychology: General*, 129 (2), 262–284.

- Monsell, S. (1978). Recency, immediate recognition memory, and reaction time. *Cognitive Psychology*, 10, 465-501.
- Murdock, B. B. (1982). A theory for the storage and retrieval of item and associative information. *Psychological Review*, 89(6), 609–626.
- Murdock, B. B. (1983). A distributed memory model for serial-order information. *Psychological Review*, 90(4), 316–338.
- Murdock, B. B. (1995). Developing TODAM: Three models for serial-order information. *Memory & Cognition*, 23, 631-645.
- Murdock, B. B. (1997). Context and mediators in a theory of distributed associative memory (TODAM2). *Psychological Review*, 104(4), 839–862.
- Navarro, D. J. (2019). Between the devil and the deep blue sea: Tensions between scientific judgement and statistical model selection. *Computational Brain & Behavior*, 2 (1), 28-34.
- Nielsen H.B., Mortensen S.B. (2016). *Ucminf: General-Purpose Unconstrained Non-Linear Optimization*. R package version 1.1-4, <https://CRAN.R-project.org/package=ucminf>.
- Nosofsky, R. M. (1986). Attention, similarity, and the identification-categorization relationship. *Journal of Experimental Psychology: General*, 115 (1), 39–57.
- Oberauer, K., Lewandowsky, S., Farrell, S., Jarrold, C., & Greaves, M. (2012). Modeling working memory: An interference model of complex span. *Psychonomic Bulletin & Review*, 19(5), 779–819.
- Osth, A. F., & Denis, S. (2015). Prior-list intrusions in serial recall are positional. *Journal of Experimental Psychology: Learning, Memory and Cognition*, 41, 1893-1901.
- Osth, A. F., & Hurlstone, M. J. (2023). Do item-dependent context representations underlie serial order in cognition? Commentary on Logan (2021). *Psychological Review*, 130(2), 513–545.
- Osth, A. F., Zhou, A., Lilburn, S. D., & Little, D. R. (2023). Novelty rejection in episodic memory. *Psychological Review*, 130 (3), 720–769.
- Raaijmakers, J. G. W., & Shiffrin, R. M. (1981). Search of associative memory. *Psychological Review*, 88 (2), 93–134.
- Ratcliff, R. (1978). A theory of memory retrieval. *Psychological Review*, 85(2), 59–108.
- Ratcliff, R., & McKoon, G. (1988). A retrieval theory of priming in memory. *Psychological Review*, 95(3), 385–408.

- Shiffrin, R. M., Huber, D. E., & Marinelli, K. (1995). Effects of category length and strength on familiarity in recognition. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 21 (2), 267–287.
- Singmann, H., Kellen, D., Cox, G. E., Chandramouli, S. H., Davis-Stober, C. P., Dunn, J. C., ... & Shiffrin, R. M. (2023). Statistics in the service of science: Don't let the tail wag the dog. *Computational Brain & Behavior*, 6 (1), 64-83.
- Solway, A., Murdock, B. B., & Kahana, M. J. (2012). Positional and temporal clustering in serial order memory. *Memory & Cognition*, 40(2), 177–190.
- Souza, A. S., & Oberauer, K. (2016). In search of the focus of attention in working memory: 13 years of the retro-cue effect. *Attention, Perception, & Psychophysics*, 78(7), 1839-1860.
- Tillman, G., Van Zandt, T., & Logan, G. D. (2020). Sequential sampling models without random between-trial variability: The racing diffusion model of speeded decision making. *Psychonomic Bulletin & Review*, 27, 911-936.
- Tolman, E. C. (1948). Cognitive maps in rats and men. *Psychological Review*, 55, 189–208.
- Tulving, E. (1962). Subjective organization in free recall of "unrelated" words. *Psychological Review*, 69 (4), 344–354.
- Usher, M., & McClelland, J. L. (2001). The time course of perceptual choice: The leaky, competing accumulator model. *Psychological Review*, 108, 550–592.
- Ward, G., Tan, L., & Grenfell-Essam, R. (2010). Examining the relationship between free recall and immediate serial recall: The effects of list length and output order. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 36(5), 1207–1241.
- Ward, G., & Tan, L. (2023). The role of rehearsal and reminding in the recall of categorized word lists. *Cognitive Psychology*, 143, 101563.
- Wickens, D. D. (1970). Encoding categories of words: An empirical approach to meaning. *Psychological Review*, 77, 1-15.
- Young, R. K. (1961). The stimulus in serial verbal learning. *American Journal of Psychology*, 74(4), 517–528.
- Zaromb, F. M., Howard, M. W., Dolan, E. D., Sirotin, Y. B., Tully, M., Wingfield, A., & Kahana, M. J. (2006). Temporal associations and prior-list intrusions in free recall. *Journal of Experimental Psychology: Learning, Memory and Cognition*, 32, 792-804.

- Zbrodoff, N. J. (1999). Effects of counting in alphabet arithmetic: Opportunistic stopping and priming of intermediate steps. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 25(2), 299–317.
- Zhang, Q., Griffiths, T. L., & Norman, K. A. (2023). Optimal policies for free recall. *Psychological Review*, 130 (4), 1104–1124.

Table 1

The response from the position coding model to a probe at position 2. The first row contains the possible responses. The second row contains the activation of responses given the probe, which is represented as the vector $\mathbf{m}$ in the model. The last three rows contain the vector $\mathbf{q}$, which represents the activation of the possible responses to the probe item. These vectors have 1 in the position of the probe letter and 0 elsewhere, so the dot product of $\mathbf{m}$ and $\mathbf{q}$ is simply 1 times the value of the probe letter in $\mathbf{m}$. Thus, $\mathbf{m} \cdot \mathbf{q}_{yes} = 1.000$, $\mathbf{m} \cdot \mathbf{q}_{within} = 0.500$, and $\mathbf{m} \cdot \mathbf{q}_{prior} = 0.500$.

Position Probe in Position 2												
Responses	“A”	“B”	“C”	“D”	“E”	“F”	“G”	“H”	“I”	“J”	“K”	“L”
$\mathbf{m}$	.500	1.000	.500	.250	.125	.063	.250	.500	.250	.125	.063	.031
Cued Recognition Item Probes												
$\mathbf{q}_{yes}$	.000	1.000	.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
$\mathbf{q}_{within}$	.000	.000	1.000	.000	.000	.000	.000	.000	.000	.000	.000	.000
$\mathbf{q}_{prior}$	.000	.000	.000	.000	.000	.000	.000	1.000	.000	.000	.000	.000

Table 2

Results of contrasts assessing distance effects for within list errors, prior list errors, the difference between within list and prior list errors, and the peak in prior list errors for distance of zero in serial recall in Experiments 1 and 2, and comparisons of effects between experiments. The peak in prior list errors (-1 0 1) assesses position-specific prior list intrusions.

Experiment	<i>t</i>	<i>SEM</i>	<i>p</i>	N > 0	BF ₁₀
Within List Errors (-2 -1 1 2)					
1	12.7841	9.1887	<.0001	32	>1000
2	14.3472	8.2638	<.0001	32	> 1000
1 vs. 2	0.0885	12.3581	.9298	NA	0.2562
Prior List Errors (-2 -1 1 2)					
1	6.7441	1.4364	<.0001	28	>1000
2	4.6384	1.7652	<.0001	24	404.4950
1 vs. 2	0.6591	2.2758	.5123	NA	0.3071
Within List vs Prior List Errors (-2 -1 1 2)					
1	12.3385	8.7353	<.0001	32	>1000
2	14.1163	7.8190	<.0001	32	>1000
1 vs. 2	-0.2212	11.7236	.8256	NA	0.2608
Prior List Error Peak (-1 0 1)					
1	8.3025	4.4942	<.0001	30	>1000
2	8.0475	3.8327	<.0001	31	>1000
1 vs. 2	1.2222	5.9830	.2263	NA	0.4798

Note df = 31 for within-experiment (within-subject) comparisons (Experiments 11 or 12); df = 62 for between-experiment (between-subject) comparisons (Experiments 11-12).

Table 3

Contrasts evaluating the four predictions of position coding theories for current and prior lists in cued recognition (distances compared are in brackets) in Experiments 3-6.

Exp	<i>t</i> (31)	<i>SEM</i>	<i>p</i>	N > 0	BF ₁₀	<i>t</i> (31)	<i>SEM</i>	<i>p</i>	N > 0	BF ₁₀
1. RT Distance Within List (-2 -1 1 2)						1. Error Rate Distance Within List (-2 -1 1 2)				
3	4.9315	30.0478	<.0001	25	867.4468	6.4008	0.0251	<.0001	31	41619.32
4	4.3786	34.4401	0.0001	22	207.3346	7.2050	0.0254	<.0001	28	336411
5	4.2778	20.5961	0.0002	25	160.3881	4.0478	0.0196	0.0003	22	89.8625
6	5.2424	21.5628	<.0001	27	1962.11	4.2424	0.0194	0.0002	24	146.6164
2. RT Distance Prior List (-2 -1 1 2)						2. Error Rate Distance Prior List (-2 -1 1 2)				
3	0.7666	17.7154	0.4491	16	0.2477	0.1341	0.0194	0.8942	14	0.1904
4	1.1039	24.3338	0.2781	18	0.3296	0.3938	0.0165	0.6964	15	0.2029
5	1.6062	19.4070	0.1184	15	0.5997	0.4191	0.0187	0.6780	13	0.2049
6	0.3551	21.6912	0.7249	16	0.2002	0.8212	0.0143	0.4178	14	0.2577
3. RT Distance Within vs Prior (-2 -1 1 2)						3. Error Rate Distance Within vs Prior (-2 -1 1 2)				
3	3.5091	38.3569	0.0014	26	24.2132	4.9108	0.0321	<.0001	27	821.7429
4	3.5596	34.8176	0.0012	22	27.2920	6.1648	0.0287	<.0001	26	22381.54
5	1.6675	34.1431	0.1055	22	0.6537	2.4691	0.0289	0.0193	23	2.5406
6	3.0685	34.3286	0.0044	23	8.8292	3.1934	0.0221	0.0032	24	11.6707
4. RT Peak Prior List (-1 0 1)						4. Error Rate Peak Prior List (-1 0 1)				
3	0.0235	32.9330	0.9814	15	0.1889	1.7545	0.0193	0.0892	16	0.7422
4	0.5569	46.9468	0.5816	13	0.2180	1.1223	0.0220	0.2184	12	0.3357
5	0.6473	20.8574	0.5222	18	0.2292	0.1201	0.0217	0.9052	11	0.1901
6	0.0865	32.4205	0.9317	17	0.1895	0.3848	0.0203	0.7030	10	0.2022

Table 4

Contrasts evaluating the four predictions of position coding theories for current and prior lists in cued recognition (distances compared are in brackets) in Experiments 7-10.

Exp	$t(31)$	SEM	p	$N > 0$	BF_{10}	$t(31)$	SEM	p	$N > 0$	BF_{10}
1: RT Distance Within List (-2 -1 1 2)						1: Error Rate Distance Within List (-2 -1 1 2)				
7	7.2817	23.1304	<.0001	30		5.0076	0.0268	<.0001	28	>1000
8	3.8957	24.6635	0.0005	26		4.2463	0.0250	0.0002	25	148.0722
9	5.8127	18.1504	<.0001	28	8841.26	3.2670	0.0131	0.0027	18	13.7936
10	4.6166	31.8936	0.0001	27	382.312	5.4074	0.0270	<.0001	25	3031.31
2: RT Distance Prior List (-2 -1 1 2)						2: Error Rate Distance Prior List (-2 -1 1 2)				
7	0.4769	21.4020	0.6368	20		-2.7829	0.0150	0.0091	11	4.7773
8	-1.2759	22.4040	0.2115	13		0.5961	0.0153	0.5555	13	0.2226
9	0.3058	19.6308	0.7618	22	0.1972	-0.2981	0.0088	0.7676	10	0.1968
10	-2.3513	20.3793	0.0252	12	2.031	0.7367	0.0230	0.4669	14	0.2426
3: RT Distance Within vs Prior (-2 -1 1 2)						3: Error Rate Distance Within vs Prior (-2 -1 1 2)				
7	5.2315	30.2443	<.0001	26		5.3676	0.0328	<.0001	25	>1000
8	2.7949	29.2231	0.0088	20		3.6148	0.0319	0.0011	22	31.1332
9	4.0090	24.8189	0.0004	27	81.5769	2.8727	0.0158	0.0073	19	5.7731
10	4.9613	39.3365	<.0001	27	937.795	4.1223	0.0313	0.0003	26	108.295
4: RT Peak Prior List (-1 0 1)						4: Error Rate Peak Prior List (-1 0 1)				
7	1.2615	19.9368	0.2615	18		2.1627	0.0193	0.0384	23	1.4426
8	1.2759	22.4040	0.2115	19		1.5081	0.0155	0.1417	22	0.5254
9	0.1760	17.8834	0.8615	17	0.1916	1.4827	0.0105	0.1483	8	0.5083
10	1.6878	31.1194	0.1015	20	0.6731	0.2853	0.0228	0.7773	13	0.1961

Table 5

Contrasts evaluating the four predictions of position coding theories for current and prior lists in cued recognition (distances compared are in brackets) in Experiments 11 and 12.

Exp	<i>t</i>	<i>SEM</i>	<i>p</i>	N > 0	BF ₁₀	<i>t</i>	<i>SEM</i>	<i>p</i>	N > 0	BF ₁₀
1: RT Distance Within List (-2 -1 1 2)						1: Error Rate Distance Within List (-2 -1 1 2)				
11	5.7834	27.7374	<.0001	29	8182.83	7.0412	0.0198	<.0001	28	220612
12	7.4757	23.9611	<.0001	28	672278	7.1550	0.0248	<.0001	29	295826
11-12	-0.5110	36.6153	0.6612	NA	0.2853	-1,1438	0.0328	0.2571	NA	0.4438
2: RT Distance Prior List (-2 -1 1 2)						2: Error Rate Distance Prior List (-2 -1 1 2)				
11	-1.3830	35.7014	0.1765	13	0.4486	0.8915	0.0234	0.3795	17	0.2722
12	-0.9512	18.8222	0.3489	13	0.2861	0.9049	0.0158	0.3725	16	0.2752
11-12	0.4311	28.6265	0.6679	NA	0.2764	0.2325	0.0280	0.8169	NA	0.2613
3: RT Distance Within vs Prior (-2 -1 1 2)						3: Error Rate Distance Within vs Prior (-2 -1 1 2)				
11	5.2983	35.9851	<.0001	25	2273.42	3.8410	0.0309	0.0006	26	53.8626
12	7.9235	24.8662	<.0001	30	2081998	6.2139	0.0308	<.0001	27	25469.5
11-12	0.1456	43.7408	0.8847	NA	0.2577	1.0570	0.416	0.2946	NA	0.4096
4: RT Peak Prior List (-1 0 1)						4: Error Rate Peak Prior List (-1 0 1)				
11	0.1952	35.7014	0.8465	12	0.1922	-1.0553	0.0234	0.2995	11	0.3144
12	-0.7225	24.9867	0.4754	11	0.2403	-2.2424	0.0267	0.322	8	1.6627
11-12	0.6944	39.1651	0.4901	NA	0.3133	1.9842	0.0319	0.0517	NA	1.3208

Note df = 31 for within-experiment (within-subject) comparisons (Experiments 11 or 12); df = 62 for between-experiment (between-subject) comparisons (Experiments 11-12).

Table 6

Contrasts comparing goodness of fit of the position coding models with zero and nonzero prior list strength in Experiments 3-12 (nonzero fit – zero fit).

Exp	<i>t</i>	df	<i>SEM</i>	<i>p</i>	BF ₁₀	<i>t</i>	df	<i>SEM</i>	<i>p</i>	BF ₁₀
AIC*						BIC*				
3	-1.3014	31	3.9523	0.2027	0.4073	-0.2476	31	3.9530	0.8061	0.1943
4	0.0839	31	1.4138	0.9936	0.1894	3.0306	31	1.4143	0.0049	8.1222
5	-2.1235	31	2.5695	0.0418	1.3472	-0.5019	31	2.5701	0.6193	0.2122
6	-0.9192	31	3.2394	0.3651	0.2784	0.3650	31	3.2407	0.7176	0.2009
7	-1.0167	31	11.6346	0.3172		3.3946	31	1.697	<.0001	18.5117
8	11.4163	31	0.1493	<.0001		39.0482	31	0.1504	<.0001	>1000
9	-2.1075	31	4.6019	0.0433	1.3104	-1.2025	31	4.6032	0.2383	0.3649
10	1.3503	31	0.6419	0.1867	0.4313	7.8126	31	0.6442	<.0001	>1000
11	-0.2565	31	1.3063	0.7993	0.1947	2.9222	31	1.3057	0.0064	6.4179
12	-0.5223	31	1.2919	0.6052	0.2143	2.6987	31	1.2922	0.0112	4.0136
3-12	-2.3900	319	1.3985	0.0174		2.4997	319	0.7938	0.0129	1.3457
<i>r</i> RT**						<i>r</i> P(Error)**				
3	-0.4388	31	0.0176	0.6639	0.2065	0.1934	31	0.0185	0.8479	0.1921
4	1.0829	31	0.0107	0.2872	0.3229	1.3326	31	0.0141	0.1924	0.4223
5	-2.0330	31	0.0152	0.0507	1.1544	0.3001	31	0.0056	0.7661	0.1969
6	0.9629	31	0.0123	0.3431	0.289	0.5215	31	0.0073	0.6058	0.2142
7	-0.2559	31	0.0111	0.7997		2.1319	31	0.0090	0.0411	1.3670
8	-1.1341	31	0.0044	0.2655		0.7534	31	0.0038	0.4569	0.2454
9	-0.4316	31	0.0167	0.6690	0.2059	1.17887	31	0.0190	0.0834	0.3558
10	-0.6537	31	0.0118	0.5181	0.2301	1.7762	31	0.0063	0.0855	0.7668
11	-1.6353	31	0.0117	0.1121	0.6246	1.1981	31	0.0109	0.2399	0.3632
12	-0.6727	31	0.0070	0.5061	0.2328	1.0499	31	0.0102	0.3019	0.3128
3-12	-1.5701	319	0.0039	0.1174		3.2518	319	0.0116	0.0013	10.8671

Note: * = negative *t* values indicate preference for the nonzero prior strength model; ** = positive *t* values indicate preference for the nonzero prior strength model.

Table 7

Contrasts comparing goodness of fit of the zero prior list strength position coding models with and without item recognition in Experiments 3-12 (zero prior list strength and item recognition - zero prior list strength and no item recognition).

Exp	<i>t</i>	df	<i>SEM</i>	<i>p</i>	BF ₁₀	<i>t</i>	df	<i>SEM</i>	<i>p</i>	BF ₁₀
AIC*						BIC*				
3	0.0504	31	0.2335	0.9601	0.1890	5.6941	31	0.2334	<.0001	>1000
4	-0.6534	31	0.1958	0.5183	0.2301	6.0820	31	0.1957	<.0001	>1000
5	-4.3915	31	0.2879	0.0001	214.2841	0.1855	31	0.2878	0.8541	0.1919
6	-3.1638	31	0.2615	0.0035	10.9181	1.8649	31	0.2617	0.0717	0.8788
7	-1.1240	31	3.6600	0.2696	0.3363	2.5348	31	0.3479	0.0165	2.8879
8	-0.0848	31	0.1793	0.9323	0.1895	7.2667	31	0.1793	<.0001	>1000
9	-3.3575	31	0.3140	0.0021	16.9845	0.8360	31	0.3138	0.4096	0.2606
10	-2.0873	31	0.2132	0.0452	1.2657	4.0962	31	0.2130	0.0003	101.4303
11	-4.0716	31	0.2868	0.0003	95.3704	0.5070	31	0.2859	0.6158	0.2127
12	0.2941	31	17.9702	0.7706	0.1966	0.3675	31	17.9633	0.7158	0.2011
3-12	0.2050	319	5.7345	0.8377	0.0640	0.7393	319	5.6144	0.4603	0.0822
<i>r</i> RT**						<i>r</i> P(Error)**				
3	2.3828	31	0.0038	0.0235	2.1547	-1.9464	31	0.0023	0.0607	1.0008
4	3.9673	31	0.0060	0.0004	73.5475	0.8072	31	0.0031	0.4257	0.2550
5	2.8023	31	0.0035	0.0087	4.9752	-0.8148	31	0.0037	0.4214	0.2564
6	3.4010	31	0.0059	0.0019	18.7896	0.2111	31	0.0038	0.8342	0.1928
7	1.8425	31	0.0055	0.0750	0.8487	0.7218	31	0.0030	0.4758	0.2402
8	2.4032	31	0.0071	0.0224	2.2395	1.1827	31	0.0035	0.2459	0.3573
9	3.0147	31	0.0042	0.0051	7.8440	-0.2997	31	0.0036	0.7664	0.1969
10	3.5000	31	0.0069	0.0014	23.6986	-1.3297	31	0.0062	0.1933	0.4209
11	1.2186	31	0.0036	0.2322	0.3713	-1.4156	31	0.0046	0.1669	0.4669
12	0.7640	31	0.0095	0.4507	0.2472	-0.3691	31	0.0119	0.7145	0.2012
3-12	7.3881	319	0.0059	<.0001	>1000	-1.0859	319	0.0052	0.0938	0.1123

Note: * = negative *t* values indicate preference for the item recognition model; ** = positive *t* values indicate preference for the no item recognition model.

Table 8

Contrasts comparing goodness of fit of the nonzero prior list strength position coding models with and without item recognition in Experiments 3-12 (nonzero prior list strength and item recognition - nonzero prior list strength and no item recognition).

Exp	<i>t</i>	df	<i>SEM</i>	<i>p</i>	BF ₁₀	<i>t</i>	df	<i>SEM</i>	<i>p</i>	BF ₁₀
AIC*						BIC*				
3	-0.7227	31	0.2694	0.4753	0.2404	4.1675	31	0.2693	0.0002	121.3375
4	1.1520	31	10.9857	0.2581	0.3460	1.2263	31	11.0905	0.2293	0.3744
5	-3.3826	31	0.4835	0.0020	18.0024	-0.6575	31	0.4834	0.5157	0.2306
6	-1.2453	31	9.3170	0.2224	0.3823	-1.1039	31	9.3149	0.2781	0.3296
7	-1.4704	31	0.3696	0.1515	0.5003	2.0952	31	0.3696	0.0444	1.2830
8	-0.2100	31	0.1826	0.8350	0.1927	7.0086	31	0.1826	<.0001	>1000
9	-3.5927	31	0.5099	0.0011	29.5313	-1.0111	31	0.5098	0.3198	0.3017
10	-0.6733	31	9.4167	0.5058	0.2329	-0.5337	31	9.4149	0.5973	0.2155
11	-4.4876	31	0.3599	0.0001	274.1432	-0.8427	31	0.3590	0.4059	0.2619
12	-0.8390	31	20.3149	0.4079	0.2612	-0.8165	31	20.6231	0.4204	0.2568
3-12	1.0601	319	8.089	0.2889	0.1093	-0.6142	319	8.4983	0.5395	0.0756
<i>r</i> RT**						<i>r</i> P(Error)**				
3	0.9633	31	0.0047	0.3429	0.2891	-2.3544	31	0.0036	0.0251	2.0428
4	1.0547	31	0.0117	0.2997	0.3142	-0.5935	31	0.0092	0.5571	0.2223
5	2.3275	31	0.0047	0.0266	1.9430	-1.3134	31	0.0041	0.1987	0.4130
6	2.8030	31	0.0091	0.0087	4.9825	0.3438	31	0.0092	0.7333	0.1995
7	2.8732	31	0.0059	0.0073	5.7792	0.2850	31	0.0039	0.7776	0.1961
8	2.4366	31	0.0071	0.0208	2.3867	1.3352	31	0.0036	0.1915	0.4236
9	2.8420	31	0.0043	0.0079	5.4090	-0.4578	31	0.0045	0.6503	0.2081
10	2.3095	31	0.0105	0.0277	1.8794	-0.6759	31	0.0087	0.5041	0.2332
11	1.2842	31	0.0040	0.2086	0.3993	-1.9044	31	0.0050	0.0662	0.9354
12	-0.7179	31	0.0151	0.4782	0.2396	-1.6658	31	0.0153	0.1058	0.6522
3-12	4.3672	319	0.0086	<.0001	608.8551	-2.1906	319	0.0077	0.0292	0.6641

Note: * = negative *t* values indicate preference for the no item recognition model; ** = positive *t* values indicate preference for the no item recognition model.

Table B1

Summary tables for ANOVAs on Response Time (RT) and error rate (P(Error)) for “yes” responses across Experiments 3-10.

Source	df	Mean Square	<i>F</i>	<i>p</i>	η_p^2
Response Time					
List Type (L)	1	45571.9314	1.1869	.2770	.0048
Probe Type (P)	1	1858346.0481	48.4015	< .0001	.1633
Probe Delay (D)	1	6359450.0220	165.6350	< .0001	.4004
L x P	1	3168.1888	.0825	.7742	.0038
L x D	1	87652.5105	2.2830	.1321	.0091
P x D	1	47771.0235	1.2442	.2657	.0050
L x P x D	1	74845.1045	1.9494	.1639	.0078
Error	248	38394.3716			
P(Error)					
List Type (L)	1	.0037	.3172	.5738	.0013
Probe Type (P)	1	.0095	.8243	.3648	.0033
Probe Delay (D)	1	.0504	4.3790	.0374	.0174
L x P	1	.0017	.1434	.7052	.0006
L x D	1	.0213	1.8482	.1752	.0074
P x D	1	.0121	1.0495	.3066	.0042
L x P x D	1	.0146	1.2688	.2611	.0051
Error	248	.0115			

Note: Significant effects are in bold font.

Table B2

Summary tables for ANOVAs on within list distance contrasts (-2 -1 1 2) in response time (RT) and error rate (P(Error)) across Experiments 3-10.

Source	df	Mean Square	<i>F</i>	<i>p</i>	η_p^2
Response Time					
List Type (L)	1	15306.3291	.5842	.4454	.0024
Probe Type (P)	1	30465.5207	1.1628	.2819	.0047
Probe Delay (D)	1	159885.0207	6.1027	.0142	.0240
L x P	1	117.3160	.0045	.9467	.0000
L x D	1	2023.3129	.0772	.7813	.0003
P x D	1	16856.1535	.6434	.4233	.0026
L x P x D	1	8848.9297	.3378	.5617	.0014
Error	248	26199.2664			
P(Error)					
List Type (L)	1	.0069	.3783	.5391	.0015
Probe Type (P)	1	.0564	3.0753	.0807	.0122
Probe Delay (D)	1	.2542	13.8574	<.0001	.0529
L x P	1	.0063	.3422	.5591	.0014
L x D	1	.0780	4.2507	.0403	.0169
P x D	1	.0044	.2419	.6232	.0010
L x P x D	1	.0044	.2419	.6233	.0010
Error	248	.0183			

Note: Significant effects are in bold font.

Table B3

Summary tables for ANOVAs on prior list distance contrasts (-2 -1 1 2) in response time (RT) and error rate (P(Error)) across Experiments 3-10.

Source	df	Mean Square	<i>F</i>	<i>p</i>	η_p^2
Response Time					
List Type (L)	1	51938.4100	3.5204	.0618	.0140
Probe Type (P)	1	24230.8139	1.6424	.2012	.0066
Probe Delay (D)	1	10070.1225	.6826	.4095	.0027
L x P	1	21708.3389	1.4714	.2263	.0059
L x D	1	3800.7225	.2576	.6122	.0010
P x D	1	35160.9377	2.3832	.1239	.0095
L x P x D	1	65.4077	.0044	.9470	.0000
Error	248	14753.4940			
P(Error)					
List Type (L)	1	.0252	2.6913	.1022	.0107
Probe Type (P)	1	.0150	1.5977	.2074	.0064
Probe Delay (D)	1	.0220	2.3508	.1265	.0094
L x P	1	.0074	.7879	.3756	.0032
L x D	1	.0125	1.3365	.2488	.0054
P x D	1	.0000	.0007	.9786	.0000
L x P x D	1	.0000	.0007	.9786	.0000
Error	248	.0094			

Note: Significant effects are in bold font.

Table B4

Summary tables for ANOVAs on contrasts comparing within-list and prior-list distance effects (-2 -1 1 2) in response time (RT) and error rate (P(Error)) across Experiments 3-10.

Source	df	Mean Square	<i>F</i>	<i>p</i>	η_p^2
Response Time					
List Type (L)	1	123635.7454	3.0473	.0821	.0121
Probe Type (P)	1	109036.1675	2.6875	.1024	.0107
Probe Delay (D)	1	89703.9938	2.2110	.1383	.0088
L x P	1	25017.3535	.6166	.4331	.0025
L x D	1	11370.2235	.2803	.5970	.0011
P x D	1	100707.0557	2.4822	.1164	.0099
L x P x D	1	10435.8994	.2572	.6125	.0010
Error	248	40571.6360			
P(Error)					
List Type (L)	1	.0586	2.1157	.1471	.0085
Probe Type (P)	1	.0133	.4784	.4898	.0019
Probe Delay (D)	1	.4259	15.3697	<.001	.0584
L x P	1	.0000	.0016	.9679	.0000
L x D	1	.0280	1.0102	.3158	.0041
P x D	1	.0041	.1479	.7008	.0006
L x P x D	1	.0041	.1479	.7009	.0006
Error	248	.0277			

Note: Significant effects are in bold font.

Table B5

Summary tables for ANOVAs on prior list distance contrasts (-1 0 1) in response time (RT) and error rate (P(Error)) across Experiments 3-10.

Source	df	Mean Square	<i>F</i>	<i>p</i>	η_p^2
Response Time					
List Type (L)	1	4911.3816	.1494	.6995	.0006
Probe Type (P)	1	20059.4110	.6101	.4355	.0025
Probe Delay (D)	1	9019.0635	.2743	.6009	.0011
L x P	1	8538.9150	.2597	.6108	.0010
L x D	1	293.0516	.0089	.9249	.0000
P x D	1	13825.3504	.4205	.5173	.0017
L x P x D	1	88346.4160	2.6868	.1024	.0107
Error	248	32881.2113			
P(Error)					
List Type (L)	1	.0260	1.7932	.1818	.0072
Probe Type (P)	1	.0002	.0167	.8974	.0000
Probe Delay (D)	1	.0850	5.8570	.0162	.0231
L x P	1	.0022	.1517	.6972	.0006
L x D	1	.0004	.0298	.8630	.0001
P x D	1	.0088	.6050	.4374	.0024
L x P x D	1	.0053	.3661	.5457	.0015
Error	248	.0145			

Note: Significant effects are in bold font.

Table C1

Contrasts evaluating linear and quadratic trends in RT for correct responses as a function of serial position for match (yes) trials, within-list lures, and prior-list lures. Each t test has 31 degrees of freedom.

Trial	Trend	Exp	<i>t</i>	<i>SEM</i>	<i>p</i>	BF ₁₀	Ex p	<i>t</i>	<i>SEM</i>	<i>p</i>	BF ₁₀
Yes	Linear	3	0.3512	229.4812	0.7278	0.2000	7	1.2285	197.3516	0.2285	0.3753
	Quad		6.0704	179.7317	<.0001	>1000		7.8704	167.4703	<.0001	>1000
Within	Linear		1.7801	164.8887	0.0849	0.7714		1.2912	248.0939	0.2062	0.4025
	Quad		8.1089	167.0396	<.0001	>1000		7.4094	223.4416	<.0001	>1000
Prior	Linear		1.4667	138.5296	0.1525	0.4980		2.9715	111.2028	0.0057	7.1392
	Quad		4.7457	116.5855	<.0001	534.3189		3.4544	159.8824	0.0016	21.2875
Yes	Linear	4	-1.6635	169.2890	0.1063	0.6500	8	-0.7377	142.5805	0.4662	0.2428
	Quad		8.8555	136.5208	<.0001	>1000		6.9917	152.3969	<.0001	>1000
Within	Linear		-0.9802	154.2175	0.3346	0.2935		-0.6368	141.8076	0.5289	0.2278
	Quad		7.8751	181.4746	<.0001	>1000		6.6350	172.7026	<.0001	>1000
Prior	Linear		-1.2036	106.6782	0.2378	0.3653		-0.2723	108.9151	0.7872	0.1954
	Quad		4.2685	136.7072	0.0002	156.6463		3.4795	122.6473	0.0015	22.5808
Yes	Linear	5	3.8004	329.3887	0.0006	48.7689	9	1.5716	288.5769	0.1262	0.5719
	Quad		8.2568	209.4112	<.0001	>1000		9.6502	233.5605	<.0001	>1000
Within	Linear		1.8566	268.9361	0.0729	0.8675		4.8283	190.4943	<.0001	662.5154
	Quad		9.8953	136.0669	<.0001	>1000		9.3640	162.4536	<.0001	>1000
Prior	Linear		2.2000	207.5342	0.0354	1.5410		2.6624	140.5042	0.0122	3.7272
	Quad		2.3852	196.4266	0.0234	2.1645		6.1778	136.5590	<.0001	>1000
Yes	Linear	6	3.6919	201.4793	0.0009	37.4755	10	3.3521	249.8308	0.0021	16.7736
	Quad		6.5997	252.3151	<.0001	>1000		7.4755	244.1973	<.0001	>1000
Within	Linear		5.2903	142.2577	<.0001	>1000		2.1750	248.4765	0.0374	1.4742
	Quad		7.7347	246.2249	<.0001	>1000		7.9714	266.8444	<.0001	>1000
Prior	Linear		2.8782	140.4779	0.0072	0.2000		1.2249	179.5029	0.2299	0.3738
	Quad		6.7136	173.1142	<.0001	>1000		5.9659	192.7160	<.0001	>1000
Yes	Linear	11	1.2072	403.9968	0.2365	0.7714	12	2.1768	297.9729	0.0372	1.4789
	Quad		11.0087	209.3930	<.0001	>1000		12.0144	205.2672	<.0001	>1000
Within	Linear		3.7532	259.6868	0.0007	0.4980		4.0314	227.7606	0.0003	86.2591
	Quad		11.2647	194.4027	<.0001	534.3189		7.0327	252.9921	<.0001	>1000
Prior	Linear		1.3253	178.5305	0.1948	0.6500		2.7380	204.1361	0.0101	4.3517
	Quad		7.1728	173.0699	<.0001	>1000		4.9306	160.1743	<.0001	865.4073

Note: Significant effects are in bold font.

Table C2

Contrasts evaluating linear and quadratic trends in proportion of correct responses as a function of serial position for match (yes) trials, within-list lures, and prior-list lures. Each t test has 31 degrees of freedom.

Trial	Trend	Exp	<i>t</i>	<i>SEM</i>	<i>p</i>	BF ₁₀	Exp	<i>t</i>	<i>SEM</i>	<i>p</i>	BF ₁₀
Yes	Linear	3	1.7649	0.2957	0.0874	0.7539	7	-1.0001	0.2742	0.3250	0.2987
	Quad		-3.0199	0.2246	0.0050	7.9338		-4.9276	0.1495	<.0001	858.6436
Within	Linear		-1.7482	0.1466	0.0903	0.7354		-1.4901	0.1017	0.1463	0.5132
	Quad		-3.9957	0.1384	0.0004	78.9219		-6.1795	0.1651	<.0001	>1000
Prior	Linear		-0.5336	0.1142	0.5974	0.2155		-0.4618	0.0812	0.6474	0.2085
	Quad		-0.3613	0.0908	0.7203	0.2006		-2.7368	0.1033	0.0102	4.3410
Yes	Linear	4	-0.5678	0.2133	0.5743	0.2192	8	-0.7481	0.2590	0.4600	0.2445
	Quad		2.1193	11.8699	0.0422	1.3374		-3.9908	0.1913	0.0004	77.9664
Within	Linear		0.8587	0.1019	0.3971	0.2652		0.2639	0.0947	0.7936	0.1950
	Quad		-2.9357	0.1613	0.0062	6.6072		-5.9239	0.1089	<.0001	>1000
Prior	Linear		-0.3603	0.0954	0.7210	0.2006		1.4535	0.0634	0.1561	0.4897
	Quad		-1.1837	0.0805	0.2455	0.3576		-2.1754	0.0725	0.0373	1.4753
Yes	Linear	5	-2.9027	0.4330	0.0068	6.1549	9	1.5274	0.3018	0.1368	0.5390
	Quad		1.8171	0.1836	0.0789	0.8160		-1.0008	0.1077	0.3247	0.2989
Within	Linear		-3.1107	0.1708	0.0040	9.6956		-2.9900	0.1792	0.0054	7.4322
	Quad		-3.6950	0.2089	0.0008	37.7544		-7.4705	0.1627	<.0001	>1000
Prior	Linear		-2.7393	0.1683	0.0101	4.3635		-1.9992	0.1258	0.0544	1.0912
	Quad		-0.6975	0.2352	0.4907	0.2364		-1.7312	0.0939	0.0934	0.7171
Yes	Linear	6	-2.9527	0.3612	0.0060	6.8542	10	-2.0419	0.3872	0.0497	1.1718
	Quad		0.0269	0.2612	0.9787	0.1889		0.0912	0.1970	0.9279	0.1896
Within	Linear		-3.4100	0.0765	0.0018	19.1878		-3.4445	0.1461	0.0017	20.7992
	Quad		-4.6294	0.1458	0.0001	395.1826		-5.4572	0.1549	<.0001	>1000
Prior	Linear		-1.9987	0.0876	0.0545	1.0903		-1.5648	0.0909	0.1278	0.5667
	Quad		-3.5806	0.0847	0.0012	28.6911		-2.7453	0.0882	0.0100	4.4180
Yes	Linear	11	1.7684	0.2567	0.0868	0.7579	12	1.5494	0.2108	0.1314	0.5550
	Quad		-3.8493	0.1372	0.0006	54.9704		22.8041	0.1824	<.0001	>1000
Within	Linear		0.3466	0.1533	0.7313	0.1997		-0.3037	0.0926	0.7634	0.1971
	Quad		-5.6375	0.2001	<.0001	>1000		-5.2350	0.2337	<.0001	>1000
Prior	Linear		2.8955	0.0863	0.0069	6.0608		1.4146	0.1381	0.1672	0.4663
	Quad		-3.8128	0.1025	0.0006	50.2692		-3.3548	0.0880	0.0021	16.8787

Note: Significant effects are in bold font.

Table D1

Mean parameter values for model fits in Experiments 3-12.

Expt	ρ	Msmth	κ	Bound	Bias	Scl RT	Residual	<i>sprior</i>	ω_{item}
Zero Prior List Strength									
3	0.2225	0.6587	0.3819	3.5704	0.5367	4.6853	0.3241		
4	0.1247	0.5333	0.4872	2.9956	0.5378	6.2008	0.2070		
5	0.3383	0.6776	0.3027	4.9415	0.5445	3.7936	0.2621		
6	0.2071	0.5625	0.4484	3.4991	0.5312	5.0891	0.1608		
7	0.2441	0.5961	0.4337	3.7335	0.5286	4.9420	0.3261		
8	0.1371	0.5521	0.5981	3.0091	0.5276	6.0719	0.2183		
9	0.3624	0.6765	0.2921	5.0483	0.5512	3.9276	0.2558		
10	0.1932	0.5071	0.7616	3.2817	0.5321	5.4502	0.1972		
11	0.3284	0.7200	0.2686	5.3314	0.5658	3.8052	0.2753		
12	0.3508	0.6655	0.3283	5.3019	0.5563	4.2929	0.2482		
Mean	0.2508	0.6149	0.4303	4.0713	0.5412	4.8259	0.2475		
Nonzero Prior List Strength									
3	0.2780	0.5499	0.3839	3.7374	0.5308	5.0790	0.3090	0.1636	
4	0.1410	0.5056	0.4894	3.0250	0.5363	6.3203	0.2049	0.0528	
5	0.4240	0.5403	0.3229	5.3506	0.5311	4.8966	0.2257	0.2085	
6	0.2271	0.5173	0.4613	3.6473	0.5276	5.3816	0.1497	0.0857	
7	0.2585	0.5615	0.4352	3.7699	0.5270	5.0137	0.3237	0.0594	
8	0.1411	0.5421	0.6003	3.0157	0.5273	6.1024	0.2178	0.0151	
9	0.3855	0.6269	0.2965	5.1239	0.5485	4.0078	0.2495	0.0712	
10	0.2035	0.4797	0.7643	3.3100	0.5309	5.4906	0.1942	0.0421	
11	0.3838	0.6218	0.2845	5.7388	0.5558	4.6619	0.2343	0.1627	
12	0.3839	0.6120	0.3337	5.3678	0.5527	4.3867	0.2431	0.0764	
Mean	0.2826	0.5557	0.4372	4.2086	0.5368	5.1341	0.2352	0.0938	
Zero Prior List Strength and Item Recognition									
3	0.1717	0.6790	0.3917	3.5904	0.5374	4.8299	0.3226		0.0439
4	0.0850	0.5343	0.5127	2.9979	0.5384	6.5128	0.2075		0.0318
5	0.1738	0.7611	0.3036	5.0613	0.5465	3.8063	0.2513		0.1495
6	0.0944	0.6003	0.4837	3.5389	0.5326	5.4687	0.1579		0.0980
7	0.1511	0.5013	0.4543	3.7407	0.5297	5.1595	0.3267		0.0742
8	0.0897	0.4696	0.6329	3.0233	0.5279	6.4195	0.2174		0.0389
9	0.1396	0.8200	0.2889	5.1902	0.5547	3.8781	0.2438		0.1955
10	0.1342	0.5900	0.8431	3.2920	0.5334	5.9884	0.1970		0.0770
11	0.1300	0.8169	0.2645	5.4740	0.5688	3.8088	0.2635		0.1787
12	0.1396	0.7603	0.3498	5.3859	0.5602	4.4834	0.2439		0.1723
Mean	0.1309	0.6533	0.4525	4.1295	0.5429	5.0355	0.2432		0.1060
Nonzero Prior List Strength and Item Recognition									
3	0.1849	0.5561	0.3963	3.8551	0.5300	5.3521	0.2967	0.2050	0.0923
4	0.0858	0.5093	0.5192	3.0312	0.5369	6.6807	0.2050	0.0588	0.0446
5	0.1769	0.6693	0.2860	5.9494	0.5568	4.6514	0.2154	0.2114	0.1938
6	0.1144	0.5320	0.5009	3.7088	0.5283	5.8238	0.1448	0.1057	0.1052
7	0.1798	0.4645	0.4587	3.7842	0.5278	5.2670	0.3237	0.0723	0.0692
8	0.0897	0.4554	0.6372	3.0346	0.5276	6.4696	0.2164	0.0194	0.0434
9	0.1317	0.7603	0.2946	5.3046	0.5519	3.9815	0.2335	0.0935	0.2285
10	0.1381	0.5620	0.8546	3.3316	0.5322	6.0929	0.1926	0.0489	0.0852
11	0.1769	0.6693	0.2860	5.9494	0.5568	4.6514	0.2154	0.2114	0.1938
12	0.1338	0.7009	0.3562	5.4980	0.5564	4.6102	0.2352	0.1066	0.2147
Mean	0.1412	0.5879	0.4590	4.3447	0.5405	5.3581	0.2279	0.1133	0.1271

Table D2

Measures of goodness of fit for each model fit in Experiments 3-12

Expt	ZPL	NZPL	IRO	IRPL	ZPL	NZPL	IRZ	IRNZ
	Akaike Information Criterion				Bayesian Information Criterion			
3	519.50	514.36	519.54	513.74	548.66	547.68	552.86	551.23
4	252.63	252.75	252.23	292.77	281.81	286.09	285.57	329.10
5	738.02	732.57	734.03	727.40	767.19	765.90	767.36	764.90
6	493.53	490.56	490.92	453.87	522.66	523.84	524.20	491.32
7	569.19	565.85	565.94	563.02	597.18	599.16	599.25	600.49
8	340.08	341.78	340.03	341.66	369.25	375.12	373.37	379.17
9	762.73	753.03	759.40	747.24	791.87	786.34	792.70	784.71
10	647.90	648.76	646.49	628.71	677.06	682.09	679.82	666.20
11	694.89	694.56	691.20	689.45	723.95	727.76	724.40	726.80
12	706.02	705.34	722.73	651.45	735.15	738.64	756.03	685.39
Mean	544.31	541.97	544.19	533.43	573.46	575.28	577.50	536.51
	Correlation RT				Correlation P(Error)			
3	0.6769	0.6692	0.7056	0.6834	0.7149	0.7185	0.7010	0.6920
4	0.6760	0.6876	0.7512	0.7265	0.5869	0.6056	0.5949	0.5883
5	0.7701	0.7391	0.8007	0.7735	0.7594	0.7611	0.7499	0.7441
6	0.6175	0.6294	0.6807	0.7097	0.6818	0.6856	0.6843	0.6956
7	0.6962	0.6911	0.7407	0.7373	0.6932	0.7036	0.6877	0.6950
8	0.5767	0.5696	0.6198	0.6121	0.7643	0.7661	0.7791	0.7822
9	0.7285	0.7213	0.7686	0.7603	0.7075	0.7415	0.7041	0.7351
10	0.6857	0.6780	0.7615	0.7544	0.6533	0.6645	0.6273	0.6460
11	0.8188	0.7996	0.8328	0.8157	0.7242	0.7373	0.7034	0.7073
12	0.8072	0.8025	0.8301	0.7682	0.7849	0.7955	0.7710	0.7148
Mean	0.7152	0.7088	0.7623	0.7490	0.6903	0.6979	0.6812	0.6760

Note: ZPL = zero prior list strength; NZPL = nonzero prior list strength; IRZ= item recognition with zero prior list strength; IRNZ = item recognition with nonzero prior list strength.

Appendix A: Simulation Methods

We conducted two sets of simulations. One varied the strength of associations to the prior list (*sprior*). The other varied the probability of using position coding (*pprior*). Each simulation generated a list of five items in which the middle position was cued, creating distances $\{-2 -1 0 1 2\}$. The activation of current and prior list items was generated using Equation 1 for each distance. These activation values were used to generate drift rates for the limited-capacity racing diffusion model using Equation 2 for recall and Equations 6 and 7 for cued recognition. Thus, the same position codes and representations of order were used to simulate recall and cued recognition. In both simulations, $\rho = .3$, recall threshold = 10.0, recognition thresholds = 2.8 for “yes” and 3.0 for “no,” $\kappa = 1.0$, and $\lambda = 0.8$. In the simulations that varied list probability, prior list strength was greater than zero (*sprior* > 0) on *pprior* proportion of the trials (when position coding was engaged) and set equal to zero (*sprior* = 0) on $1 - \textit{pprior}$ proportion of the trials (when position coding was not engaged).

On each trial, the simulation used drift rates defined in Equation 2 or Equations 6 and 7 and a threshold (10 for recall; 2.8 for “yes” and 3.0 for “no” in cued recognition) to sample a random value from a Wald distribution (the finishing time distribution for a diffusion to a single bound) for each response category (10 current and prior list items for recall; “yes” vs. “no” for cued recognition), and the simulation chose the category with the shortest simulated RT. Each condition (recall vs. recognition x 10 current- and prior-list items) was simulated 100000 times. Response probabilities and mean RTs were calculated for each response category as a function of the cued position in the current or prior list. The results are plotted in Figures 2 (*sprior* varied) and A1 (*pprior* varied).

To simulate recall, the program stepped through the 10 items in the current and the prior lists, using Equation 5 to calculate the probability of recalling the items in each list given their activation and strength of association to the position code (1 for the current list; *sprior* for the prior list) when trying to recall the item in position 3 in the current list. To simulate cued recognition, the program stepped through the same 10 items in the current and prior lists, using Equation 8 and 9 to simulate the probability and response time (RT) for “yes” and “no” decisions, respectively. To evaluate the effects of the strength of prior associations, the simulation was run five times with *sprior* = .1, .2, .3, .5, and .7 to cover the range where the

changes were most dramatic. To evaluate the effects of the probability of using position coding, the simulation was run five times with $p_{prior} = .1, .2, .3, .5$, and $.7$ with s_{prior} fixed at $.5$.

Matlab code for the simulations and the simulation results can be found on the Open Science Framework at <https://osf.io/j4z7a/>

Appendix B: Between-Experiment ANOVAs

We compared Experiments 3-10 in 2 (precue vs no precue) x 2 (spatial vs numeric cues) x 2 (constrained vs unconstrained lists) between-subject ANOVAs on RT and error rate for “yes” responses (Table B1), within-list distance contrasts (-2 -1 1 2; Table B2), prior list distance contrasts (-2 -1 1 2; Table B3), contrasts evaluating the difference between within-list and prior-list distance contrasts (Table B4), and contrasts evaluating the peak in prior-list distance effects at distance = 0 (Table B5).

Appendix C: Serial Position Effects

Mean RTs for correct responses and error rates for match (“yes”), within-list lures (“no”), and prior-list lures (“no”) in Experiments 3-12 are plotted as a function of the serial position of the probe in Figure C1. Contrasts evaluating linear and quadratic trends in the serial position effects in these data are presented in Tables C1 (RT) and C2 (proportion correct). The raw data and the means Figure C1 depicts are available on the Open Science Framework at <https://osf.io/j4z7a/>.

The contrasts can be interpreted as measures of the direction of sequential access to list items (Logan et al., 2023a): The linear trend reflects sequential access from the beginning of the list (positive slope) or from the end of the list (negative slope). The quadratic trend reflects access from both ends of the list, as if subjects start at the end of the list that is nearest to the probed position. Of course there are other interpretations of the serial position effects, including interference (greater for middle positions) and encoding differences (early items may be encoded better than later items).

Appendix D: Model Fitting Methods

The models we fit to the data from each experiment are simplified versions of the models Logan et al. (2021) fit to their episodic flanker task. We assume that memory for the current list

is represented in the form of a matrix $\mathbf{M}$. The matrix $\mathbf{M}$ has N rows and 6 columns. N is the total number of unique items in the stimulus set (for the consonants used in our experiments, $N = 20$). The six columns correspond to the six locations in which items are presented. The entry m_{ij} in the matrix $\mathbf{M}$ gives the degree to which item i is activated by the position code for location j (i.e., $m_{ij} = a(i|j)$ in Equation 1). Let $\mathbb{C}_i$ be an indicator variable that equals 1 if item i was on the current list and zero otherwise and let $\mathbb{P}_i$ be an indicator variable that equals 1 if item i was on the previous list and zero otherwise. Then m_{ij} is given by $m_{ij} = (\mathbb{C}_i + \textit{s prior} \times \mathbb{P}_i)\rho^{|i-j|}$ where parameters ρ and *s prior* are as defined in the main text.

Each trial of cued recognition involves a probe item and a cued location k . The probe item is represented using a vector $\mathbf{q}$ with a 1 in the entry corresponding to the probe item and zeros elsewhere. The degree to which the probe item is activated by the code for position k is given by the dot product between the vector $\mathbf{q}$ and the k th column of $\mathbf{M}$. The k th column of $\mathbf{M}$, written as $\mathbf{m}_{.k}$, is equivalent to the vector of item activations $\mathbf{m}$ described in the main text. The only difference is that, in the main text, only the elements of $\mathbf{m}$ corresponding to items that were in either the current or prior list are depicted; all other elements of $\mathbf{m}$ have activations of zero (since, for any item i not in either the current or prior list, $\mathbb{C}_i = \mathbb{P}_i = 0$).

As described in the main text, a recognition decision is modeled as the outcome of a race between a “yes” accumulator and a “no” accumulator. The input to the “yes” accumulator is a function of the degree to which the contents of the recognition probe match the contents of memory. The input to the “no” accumulator is a function of the maximum possible match value. As such, a subject will be more willing to make a “yes” response, and to do so more quickly, to the extent that the degree of match is large relative to how large it could be. In total, we fit four different models to each of our cued recognition experiments. The four models represent a factorial combination of the presence or absence of two potential contributors to the recognition process: prior-list representations and item recognition. The simplest model, with no additional contributors, assumes that the inputs to the “yes” and “no” accumulators depend only on a comparison between the probe and the memory representation for the cued position in the current list, that is, the column of $\mathbf{M}$ corresponding to the cued position. We first describe the simplest model and its implementation.

Cued Recognition

On any given trial, the probe consists of an item and cued location k , which together are used to construct a vector $\mathbf{q}$ which serves as a retrieval cue. The vector $\mathbf{q}$ has 6 entries, one for each possible position. If the probe item had been presented at position i in the current list, then the vector $\mathbf{q}$ has a 1 in its i th position and zeros elsewhere. Otherwise, vector $\mathbf{q}$ consists of all zeros, although this is merely a shorthand for the idea that the probe item does not have a corresponding row in the memory matrix $\mathbf{M}$. The joint item-position match is the dot product between the k th column of $\mathbf{M}$ and the cue vector $\mathbf{q}$. Because $\mathbf{q}$ has zeros everywhere except for the entry corresponding to the position in which the probe item had been studied (if it had been), this dot product is simply $m_{ik} = \rho^{|I-k|}$, i.e., the degree to which the item studied in position I is associated with cued location k . This match value is multiplied by a scaling parameter A ($A > 0$) to yield T_Y , the total input to the “yes” accumulator:

$$T_Y = A(\mathbf{q} \cdot \mathbf{M}_{\cdot k}) = A\rho^{|i-k|}$$

The maximum possible match is the product of the magnitudes of $\mathbf{q}$ and $\mathbf{M}_{\cdot k}$, which would occur if they had exactly the same values in each of their entries. By design, the magnitude of $\mathbf{q}$ is $\|\mathbf{q}\| = 1$, so the maximum possible match is determined by the magnitude of $\mathbf{M}_{\cdot k}$. The magnitude of $\mathbf{M}_{\cdot k}$ is the square root of the sum of the squared entries in column k of matrix $\mathbf{M}$, i.e.,

$$\|\mathbf{M}_{\cdot k}\| = \sqrt{\sum_{j=1}^6 \rho^{2|j-k|}}.$$

The maximum match is multiplied by both the scaling parameter A from above as well as an additional weighting factor λ ($\lambda > 0$) to yield the total input to the “no” accumulator:

$$T_N = A\lambda(\|\mathbf{q}\| \times \|\mathbf{M}_{\cdot k}\|) = A\lambda\|\mathbf{M}_{\cdot k}\| = A\lambda \sqrt{\sum_{j=1}^6 \rho^{2|j-k|}}$$

where the λ parameter acts to give different degrees of weight to mismatch information. When $\lambda \geq 1$, the total input to the “yes” accumulator will never exceed that to the “no” accumulator since, by definition, the input to the “no” accumulator is based on the maximum possible match. The λ parameter therefore reflects how large a match needs to be relative to its maximum before the match is seen as strong enough to favor a “yes” response. For example, if $\lambda = 0.5$, then the total

input to the “yes” accumulator would exceed that of the “no” accumulator if the degree of match were at least half of its maximum possible value.

The activation level of each accumulator is assumed to evolve over time according to a Wiener process with infinitesimal variance of 1. The drift rates for each accumulator are functions of the total input to each accumulator along with feedforward inhibition from the input to the other accumulator, the strength of which is governed by parameter κ ($\kappa > 0$):

$$d_Y = \frac{T_Y}{1 + \kappa T_N}$$

$$d_N = \frac{T_N}{1 + \kappa T_Y}$$

where d_Y and d_N are the drift rates for the “yes” and “no” accumulators, respectively.

Each accumulator has a threshold, θ_Y for the “yes” accumulator and θ_N for the “no” accumulator. Both accumulators start with zero activation at the beginning of a trial and the first accumulator to reach its threshold determines the response as well as the response time. We parameterize the thresholds in terms of a “response caution” parameter B ($B > 0$) and a “bias” parameter w ($0 < w < 1$). “Response caution” is the sum of the thresholds, i.e., $B = \theta_Y + \theta_N$, and reflects the total amount of memory evidence a subject generally requires before responding. “Bias” reflects the degree to which the threshold for the “yes” accumulator is lower than that for the “no” accumulator, thereby favoring a “yes” response. The two thresholds are given by

$$\theta_Y = B(1 - w)$$

$$\theta_N = Bw$$

such that the thresholds are unbiased when $w = 0.5$, are biased in favor of “yes” responses when $w > 0.5$, and are biased in favor of “no” responses when $w < 0.5$. The total response time on a given trial is the time needed for the first accumulator to reach its threshold, plus a residual time R that includes the time needed to detect and orient to the probe, to focus on the cued position, and to execute the response associated with the winning accumulator. In the present models, we simply assume that R is a constant.

To summarize, the simplest model we consider has seven free parameters: The position association gradient ($0 < \rho < 1$), the scaling parameter for converting matches to accumulator inputs ($A > 0$), the mismatch weight parameter ($\lambda > 0$), the feedforward inhibition between accumulators ($\kappa > 0$), response caution ($B > 0$), response bias ($0 < w < 1$), and residual time ($R >$

0). For models assuming no contribution from the prior list, *sprior* is not a free parameter because it is fixed at *sprior* = 0. For models that allow for prior list representations to contribute to cued recognition, *sprior* ($0 < sprior < 1$) is an additional free parameter to be estimated.

Item recognition

We also explored models that included an additional form of match/mismatch process corresponding to simple item recognition. Item recognition was modeled by matching the probe vector $\mathbf{q}$ not just to the k th column of $\mathbf{M}$, but to all columns of $\mathbf{M}$ and summing the result. This amounts to item recognition because the resulting match represents the degree to which the probe item matches anything that had been studied recently, regardless of location. This is accomplished by summing the dot products between the cue vector $\mathbf{q}$ and all 6 columns of the memory matrix $\mathbf{M}$. The total input to the “yes” accumulator is then a weighted sum of the joint item-position match and the item recognition match, where the parameter ω ($0 < \omega < 1$) represents the relative weight of the item recognition match:

$$T_Y = A \left[(1 - \omega)(\mathbf{q} \cdot \mathbf{M}_{.k}) + \omega \left(\sum_{j=1}^6 \mathbf{q} \cdot \mathbf{M}_{.j} \right) \right]$$

The maximum possible item recognition match, which contributes to the input to the “no” accumulator, is the sum of the maximum possible item-position joint match across all positions (columns) in the memory matrix $\mathbf{M}$. The contribution of the maximum possible item recognition match to the “no” input is weighted by the same factor as the contribution to the “yes” input:

$$T_N = A\lambda \left[\omega(\|\mathbf{q}\| \times \|\mathbf{M}_{.k}\|) + (1 - \omega) \left(\sum_{j=1}^6 \|\mathbf{q}\| \times \|\mathbf{M}_{.j}\| \right) \right]$$

The rest of the model is unchanged and operates exactly as described above. Thus, modeling the contribution of item recognition involves adding only one free parameter, the weight ω given to item recognition as opposed to joint item-position recognition.

Prior list representations

Just like the current list is represented in memory with the matrix $\mathbf{M}$, the previous list is represented in another matrix $\mathbf{L}$ with the same structure (i.e., six columns corresponding to the six locations in the prior list and six rows corresponding to the six items presented in the prior

list). If a probe item was present in the prior list, the degree to which it activates its representation in the prior list is given by parameter p , which ranges between 0 and 1. For models that assume no contribution of prior list representations, p is assumed to be fixed at zero. If p is greater than zero, then the prior list representation contributes to both the joint item-position match as well as the item recognition match. In addition, the maximum possible values of both types of match are higher, reflecting the additional contribution of prior-list representations.

Let $\mathbf{q}_L$ denote a cue vector constructed in an analogous manner to the one for the current list ($\mathbf{q}$). The vector $\mathbf{q}_L$ has all zeros except for the entry corresponding to the position in which the probe item appeared in the prior list (as above, this vector is all zeros if the item was not present in the prior list). Then the total match value is given by

$$T_Y = A \left\{ (1 - \omega)[\mathbf{q} \cdot \mathbf{M}_{\cdot k} + p(\mathbf{q}_L \cdot \mathbf{L}_{\cdot k})] + \omega \left[\sum_{j=1}^6 \mathbf{q} \cdot \mathbf{M}_{\cdot j} + p \left(\sum_{j=1}^6 \mathbf{q}_L \cdot \mathbf{L}_{\cdot j} \right) \right] \right\}$$

while the maximum total match value is given by

$$T_N = A\lambda \left[(1 - \omega)(\|\mathbf{M}_{\cdot k}\| + p\|\mathbf{L}_{\cdot k}\|) + \omega \left(\sum_{j=1}^6 \|\mathbf{M}_{\cdot j}\| + p \sum_{j=1}^6 \|\mathbf{L}_{\cdot j}\| \right) \right]$$

Note that, because the matrices for each list are constructed in an identical manner,

$$\sum_{j=1}^6 \|\mathbf{M}_{\cdot j}\| = \sum_{j=1}^6 \|\mathbf{L}_{\cdot j}\|.$$

Model fitting

We fit a total of four models to the data from each subject in each experiment, finding the parameters of each model that maximized the total log-likelihood of the choices and response times produced by each subject in each experiment. Let $d_Y[n]$ and $d_N[n]$ denote the drift rates for the “yes” and “no” accumulators on trial n , which are determined by the study items and cues on trial n as described above. The likelihood that the “yes” accumulator reaches its threshold at time t is given by the probability density function of an inverse Gaussian (Wald) distribution

$$f_Y(t|d_Y[n], \theta_Y) = \frac{\theta_Y}{\sqrt{2\pi t^3}} \exp \left[-\frac{(\theta_Y - t d_Y[n])^2}{2t} \right]$$

where we assume that the infinitesimal variance of the diffusion process is one (since this amounts to a scaling parameter). The likelihood that the “no” accumulator reaches its threshold

at time t is defined analogously, replacing $d_Y[n]$ with $d_N[n]$ and θ_Y with θ_N . Then, according to the racing diffusion decision process we employ, the likelihood of making response $Q[n]$ (either Y for “yes” or N for “no”) on trial n with response time $RT[n]$ is

$$L[n] = f_{Q[n]}(RT[n] - R|d_{Q[n]}[n], \theta_{Q[n]}) \times \left[1 - F_{\bar{Q}[n]}(RT[n] - R|d_{\bar{Q}[n]}[n], \theta_{\bar{Q}[n]})\right]$$

where $\bar{Q}[n]$ denotes the response that was *not* made on trial n and R is the residual time. The total log-likelihood of choices and response times is then given by

$$LL = \sum_{n=1}^{N_T} \log L[n]$$

where N_T is the total number of trials observed.

When fitting these models, we noticed some numerical problems that arose when certain parameters were allowed to take extremely large or small values, which interfered with the parameter search routines we used (discussed shortly). To address this issue, we introduced a set of regularization terms that encouraged model parameters to stay within a reasonable range. These terms amount to prior information about the scales of particular model parameters and were expressed in terms of simple probability distributions. For the bias w and position similarity gradient ψ , both of which range between 0 and 1, we imposed a weak Beta prior with both shape parameters set to 1.5. The intent of this prior was to prevent these parameters from being exactly zero or exactly one, both of which are *a priori* implausible anyway. For competition κ , boundary separation B , drift scale A , and “no” scale λ , all of which must be nonnegative, we imposed a weak Gamma prior with shape 1.05 and rate 0.05, corresponding to a mode of 1 and a standard deviation of 20. The effect of this was to avoid extremely large values while also preventing these parameters from being exactly zero; again, both of these situations are implausible regarding any of these parameters. Notice that no regularization was applied to either the prior-list strength parameter p or the item recognition weight parameter ω . This was to avoid introducing any bias into the model comparisons that might arise from favoring particular values for these parameters. As a result, the total quantity to be maximized during model fitting is given by

$$V = \sum_{n=1}^{N_T} \log L[n] + \log \text{Beta}(w|1.5, 1.5) + \log \text{Beta}(\psi|1.5, 1.5) \\ + \log \text{Gamma}(\kappa|1.05, 0.05)$$

$$+ \log \text{Gamma}(B|1.05, 0.05) + \log \text{Gamma}(A|1.05, 0.05) \\ + \log \text{Gamma}(\lambda|1.05, 0.05)$$

To find the model parameters that maximized V , we first ran the Nelder-Mead Simplex routine starting from a generic starting point for 500 iterations. The set of parameters from the final step of the Simplex search was then used as the initial seed value for a more sophisticated nonlinear optimization routine implemented in the ``ucminf`` R package (Nielsen & Mortensen, 2016).

The predicted and observed RTs in each experiment are presented in Figure D1. The predicted and observed error rates in each experiment are presented in Figure D2.

Appendix E: Parameter Recovery

In the main text, we used model fits to test the hypothesis that subjects did not activate prior list representations in cued recognition. The alternative hypothesis is that subjects did activate prior-list representations. These two hypotheses are embodied by model variants that either fix $sprior = 0$ (no prior-list activation) or allow for $sprior$ to take any value between 0 and 1 (allowing for prior-list activation). We conducted a parameter recovery exercise to understand how well the design of our cued recognition experiments would be able to distinguish between these two hypotheses. For example, it may be difficult to distinguish a subject with a very small value of $sprior$ from one with $sprior = 0$. For such a subject, model selection metrics like AIC or BIC might favor the simpler model (with $sprior$ fixed to zero) not because this subject actually had $sprior = 0$, but because the improvement in model fit is overwhelmed by the penalty for introducing an additional free parameter. These parameter recovery exercises were designed to understand how often we might expect that to occur, both at the level of individual subjects and when these metrics are aggregated across groups of subjects. For example, AIC or BIC might favor the simpler model for a single subject with a small value of $sprior$, but if many subjects have a small value of $sprior$, the more complex---and more correct---model may be identified when comparisons are based on summed or average AIC/BIC across subjects. Therefore, as part of this parameter recovery exercise, we examined not just how well AIC/BIC could distinguish between individual subjects with $sprior = 0$ vs. $sprior \neq 0$, but how well summed AIC/BIC could distinguish between groups of subjects, all of whom have $sprior = 0$ vs. $sprior \neq 0$.

We simulated 6 groups of 10000 subjects each. Within each group, each subject had the same value of *sprior*, which could take one of six different values: 0, 0.1, 0.2, 0.3, 0.5, or 0.7 (matching the values used for our initial simulations in the main text). Because these simulations did not include an item recognition component, there were seven other parameters that were randomly sampled for each simulated subject. These were sampled from the probability distributions summarized in Table E1, which were chosen to roughly match the mean and standard deviation of the estimated parameters values across all 10 cued recognition experiments described in the main text. For each subject, we simulated choice and RT in 480 trials of cued recognition. Those 480 trials had exactly the same frequency of trial types as in each of our experiments. As such, each simulated subject engaged in the same number of target, within-list lure, and prior-list lure trials across different cued locations and lags as was experienced by each actual subject. To simulate the outcome of a trial, we used the sampled parameter values for each simulated subject to compute the drift rates of the “yes” and “no” accumulators on each trial and drew random samples from the resulting Wald distributions to represent the time needed for each accumulator to reach its threshold on each trial. The simulated choice on each trial was given by which accumulator had the shortest simulated time-to-threshold. The simulated RT was how long it took the fastest accumulator to reach its threshold, plus the simulated subject’s residual time.

After simulating data from each simulated subject, we fit both the constrained model (with *sprior* fixed at zero) and the unconstrained model (with *sprior* as a free parameter) to the data from each simulated subject. To do so, we used exactly the same fitting procedure as was used for the real subjects (described in Appendix D). As such, our parameter recovery methods exactly matched the methods we used to apply these models, simply exchanging data produced by real subject with data produced by simulated subjects.

Figure E1 shows the fits of each model to the simulated data from each group of subjects. The unconstrained model that allows for nonzero prior list strength is able to fit the error rates and RTs for each group, although there is a slight tendency for this model to predict higher error rates and RTs for prior list lures at lag zero even when the data are simulated assuming *sprior* = 0. Note that this is not a consequence of the regularizing priors described in Appendix C, since no such regularization was applied to *sprior* (such regularization could push estimates of *sprior*

away from zero). On the other hand, the constrained model that assumes zero prior list strength is clearly unable to fit the data produced by subjects with $sprior > 0$.

But while this discrepancy is apparent when looking at averages over 1000 simulated participants, does it also result in model comparison metrics that favor the appropriate model? This question is addressed by Figure D2, which shows the proportion of simulated samples of different sizes (from 1 participant up to 320 participants) that resulted in summed AIC (left panel) or summed BIC (right panel) favoring the unconstrained model. For each sample size, we simulated 10000 samples by sampling with replacement from the pool of 1000 simulated participants. For single participants, AIC is more likely to favor the correct model regardless of the true value of $sprior$ (see the individual points on the left side of the left panel of Figure D.3). On the other hand, BIC is more conservative at the individual participant level, only favoring the more complex model when $sprior > 0.1$ (see the individual points on the left side of the right panel of Figure D.3). When aggregating across participants in each sample, both AIC and BIC are more likely to favor the correct model. Sample sizes of 32 and 320 are highlighted in Figure D.3 with vertical bars, since these reflect the sample size of each of our experiments (each of which had 32 actual participants) as well as the sample size across all ten cued recognition experiments (320 total participants). With a sample size of 32, summed AIC favors the correct model essentially the whole time, regardless of the value of $sprior$ used to simulate the data. With a sample size of 32, summed BIC favors the correct model almost always except when $sprior = 0.1$, in which case it correctly favors the unconstrained model in 63% of simulated samples. When aggregating across 320 participants---equivalent to aggregating across each participant across all 10 of our cued recognition experiments---both summed AIC and summed BIC favor the correct model the vast majority of the time; summed BIC favors the correct model in 94% of samples of size 320 when $sprior = 0.1$. To be sure, these results are optimistic in the sense that the models being used to fit the data have the same structure as the models used to produce the data. Moreover, each group of simulated participants has the same value of $sprior$ when actual participants would not be so homogeneous. Nonetheless, these results verify that our experimental designs have sufficient power to distinguish between participants with different values of $sprior$ on the basis of relative model fit. Moreover, considerable power might be achieved by aggregating across participants.

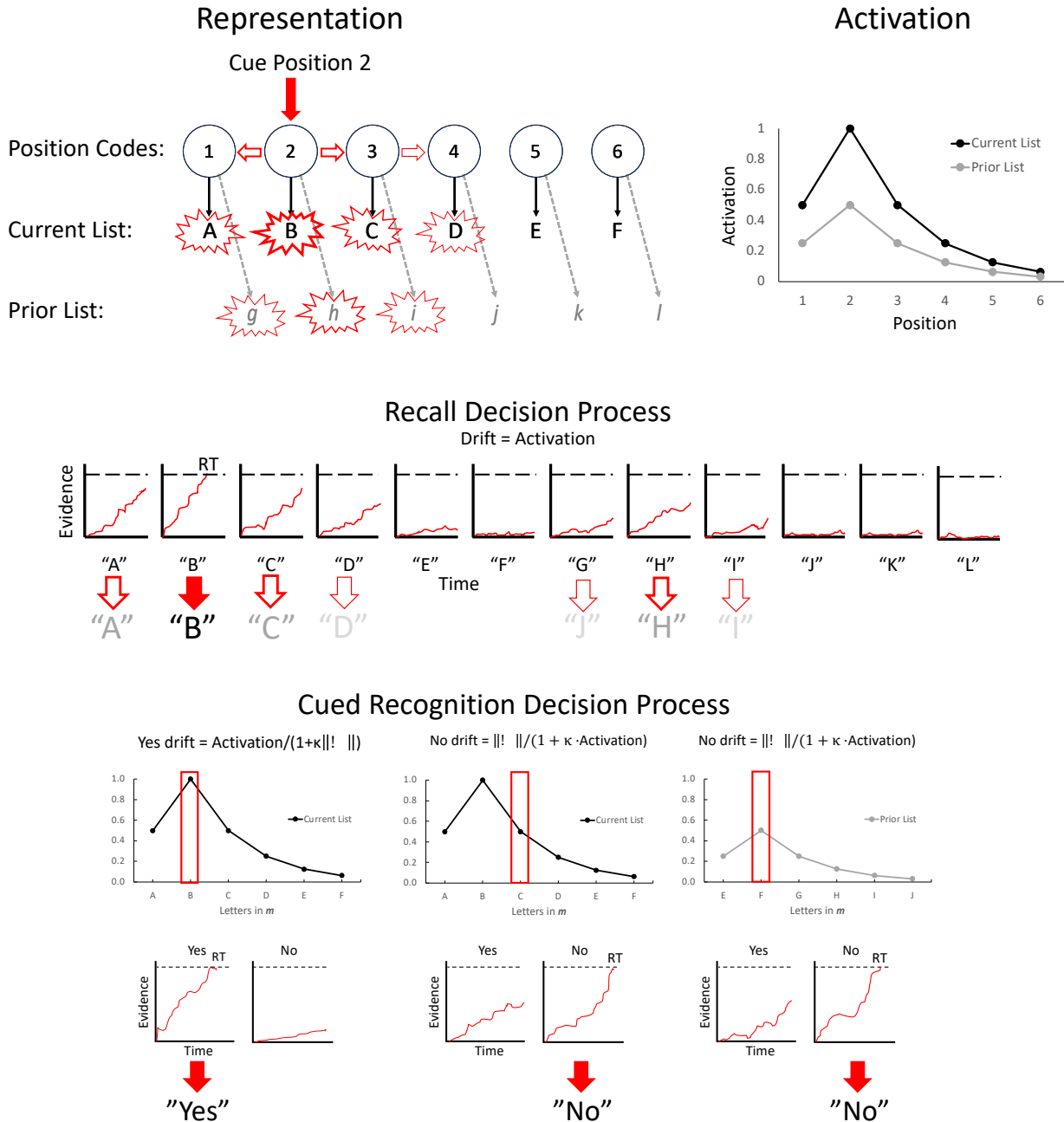
Figure 1: Position Coding Model

Figure 1 caption: The simple position coding model. The top row shows its representation (left) and activation (right) assumptions, illustrating a probe cuing the second position. The probe activates position code 2 and its neighbors, and they activate items on the current and prior lists that were associated with them. Activation peaks at the cued position and decreases with distance for both the current list and the prior list, but prior list activation is weaker because the associations are not as strong. The second row shows the decision process for recall. The

activations produced by the probe become drift rates in separate diffusion processes, each with a single boundary. The first to reach its boundary determines the response and its response time. The third row and fourth rows show the decision process in cued recognition. The probe item is compared with the activated items by taking the dot product of a vector representing activation of possible responses and a vector representing the activation of the probe letter in the probe. There is only one letter in the probe, so the vector has activation = 1 in that position and 0 everywhere else. Consequently, the dot product is simply 1 times the activation of the probe letter in the memory lists. This is illustrated by the red boxes on the activation functions in the third row. The activation increases drift rate for “yes” responses and decreases drift rate for “no” responses. The graded activation of current and prior list lures predicts distance effects for both lists and position-specific interference for prior list lures.

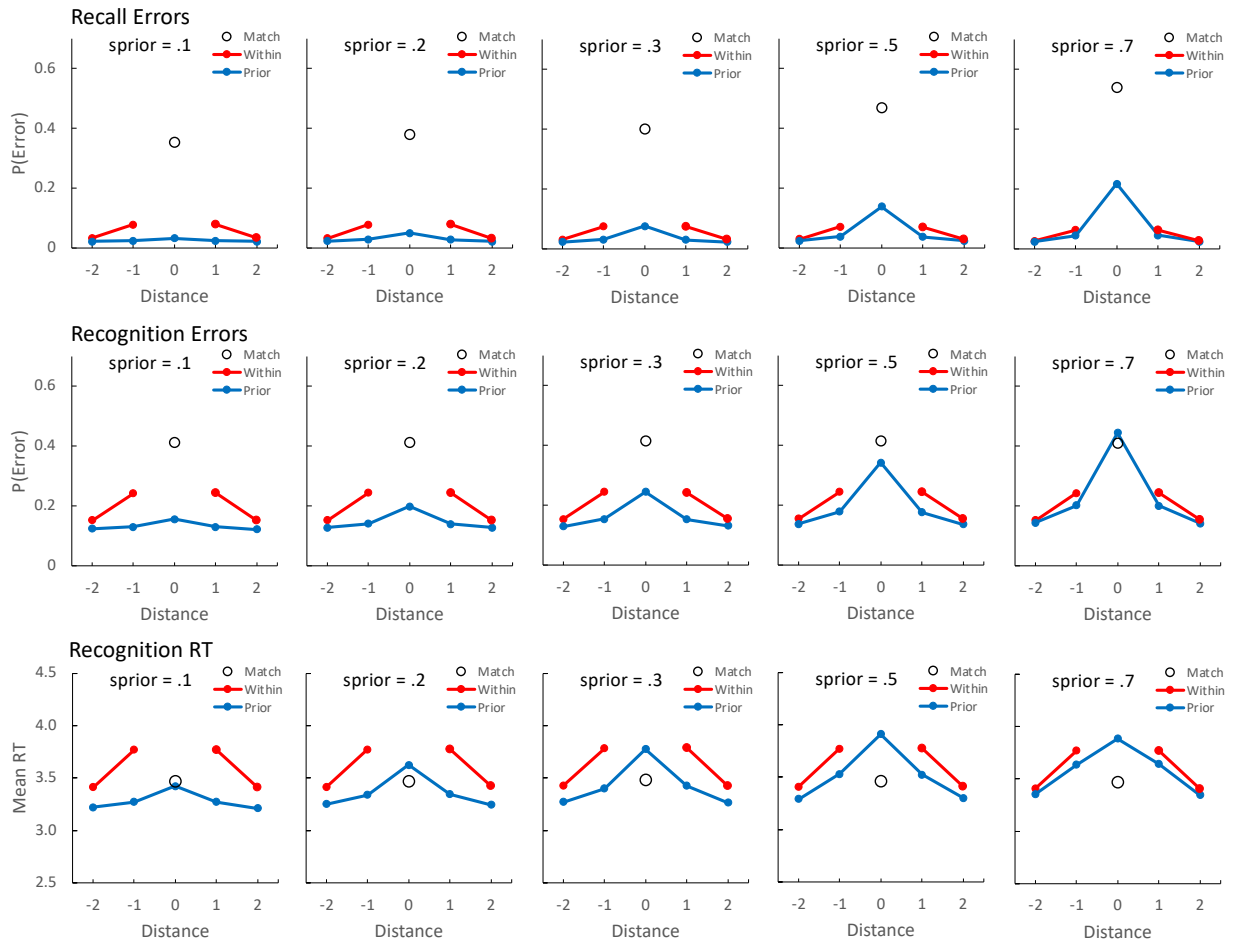
Figure 2: Position Coding Predictions

Figure 2 caption: Simulated predictions of within- and prior-list distance effects in response time (RT) and response probability from the position coding model in Figure 1. The same representations of position are used in each panel. The columns represent different values of prior list strength (.1-.7) relative to current list strength (1.0). The top row presents serial and cued recall error rates, the middle row presents cued recognition task error rates, and the bottom row presents cued recognition response times (RT) in arbitrary units. Prior list distance effects are observed in recall error rates for list strengths $\geq .2$. They are observed in cued recognition error rates and RTs for all prior list strengths.

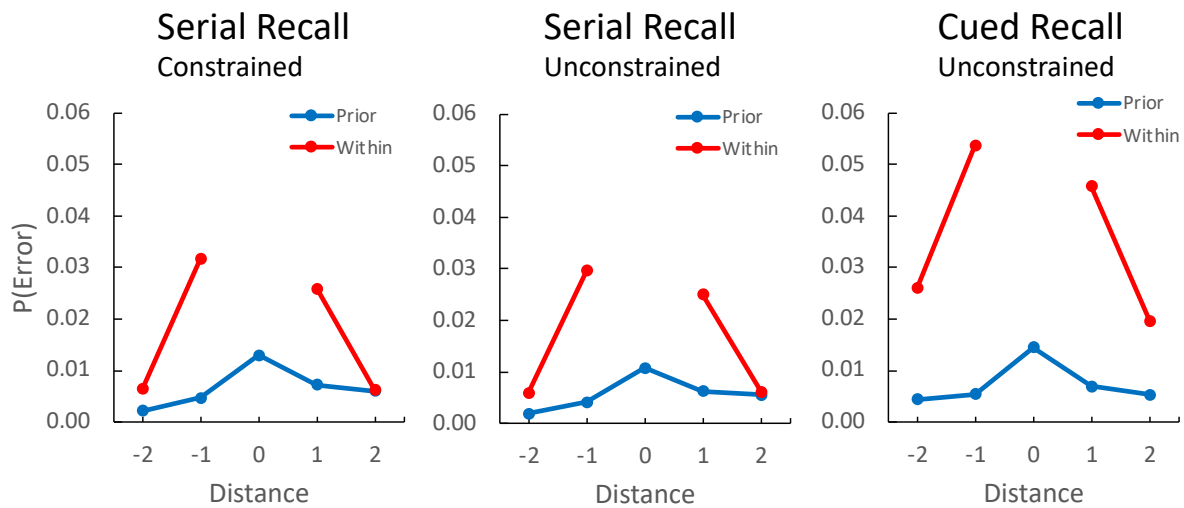
Figure 3: Recall

Figure 3 caption. Within-list (red) and prior-list (blue) intrusions as a function of distance from the correct position. The left and middle panels contain results from Experiments 1 and 2, respectively. The right panel shows results from cued recall experiments reported by Logan et al. (2023a), which used the same list length, exposure duration, and retention interval as the present experiments.

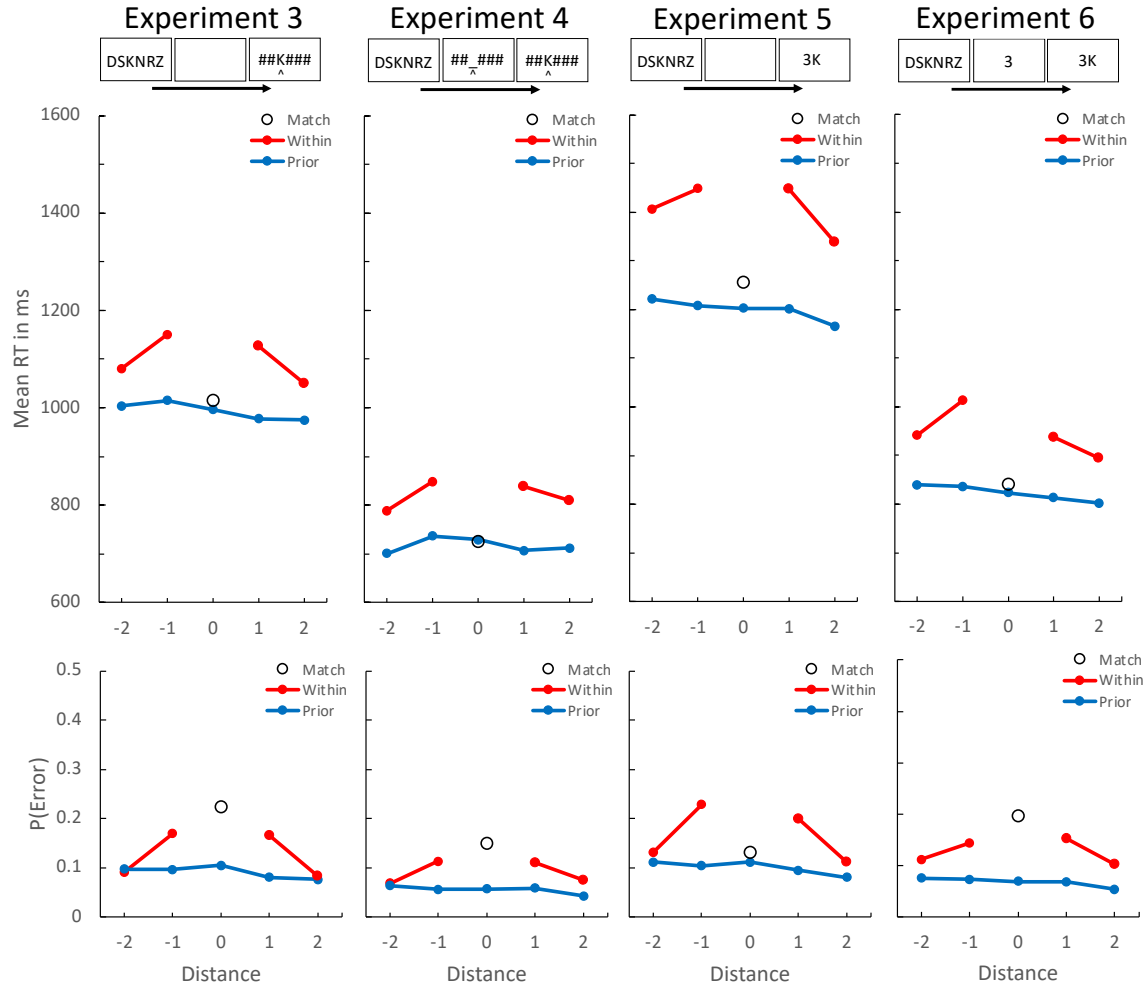
Figure 4: Cued Recognition Unconstrained Lists

Figure 4 caption: Mean response times (RTs; top panels) and error rates (bottom panels) as a function of distance between the cued position and the position of the probed item in the current (within) or prior list for responses to Matches (“yes”) and responses to within-list and prior-list lures (“no”) in Experiments 3-6. The cuing procedure for each experiment is illustrated at the top of each column (list → retention interval → probe). In Experiments 3 and 6, the position cue is presented before the probe item. Experiments 3-6 used lists that were constrained not to repeat letters from the immediately previous list.

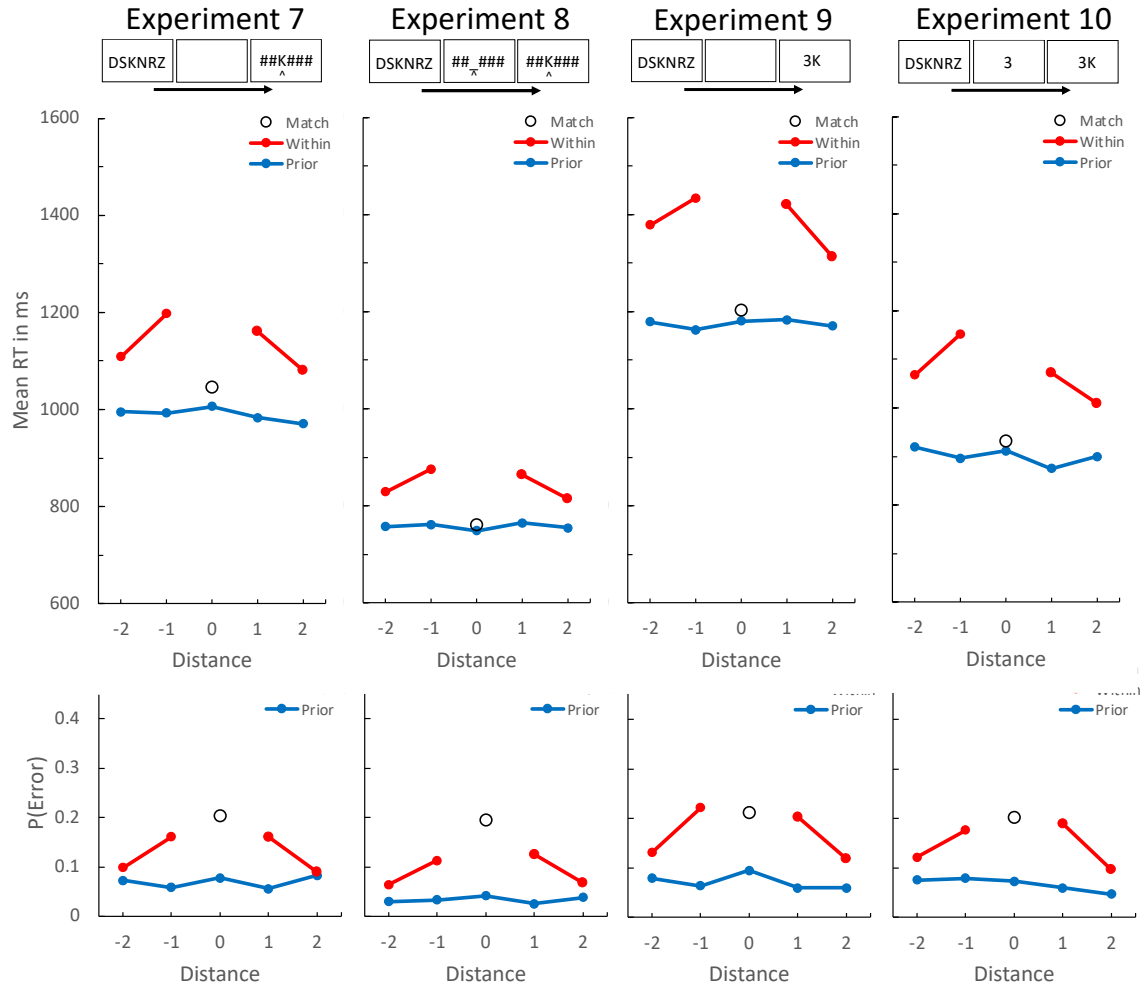
Figure 5: Cued Recognition Constrained Lists

Figure 5 caption: Mean response times (RTs; top panels) and error rates (bottom panels) as a function of distance between the cued position and the position of the probed item in the current (within) or prior list for responses to Matches (“yes”) and responses to within-list and prior-list lures (“no”) in Experiments 7-10. The cuing procedure for each experiment is illustrated at the top of each column (list → retention interval → probe). In Experiments 8 and 10, the position cue is presented before the probe item. Experiments 8-10 used unconstrained lists, in which letters from the immediately previous list were allowed to repeat.

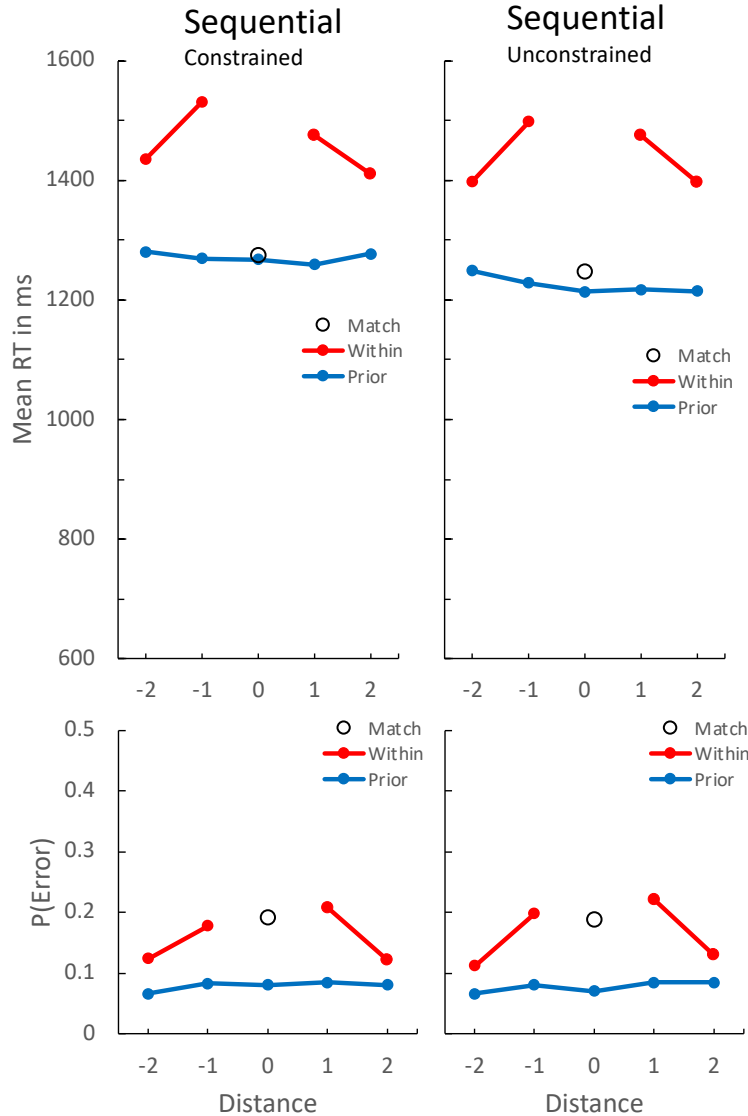
Figure 6: Sequential Lists

Figure 6 caption: Mean response times (RTs; top panels) and error rates (bottom panels) as a function of distance between the cued position and the position of the probed item in the current (within) or prior list for responses to Matches (“yes”) and responses to within-list and prior-list lures (“no”) in Experiments 11 and 12. Both experiments presented the memory lists sequentially and both used simultaneous numeric probes to cue recognition (e.g., 5D). Experiment 11 used constrained lists. Experiment 12 used unconstrained lists.

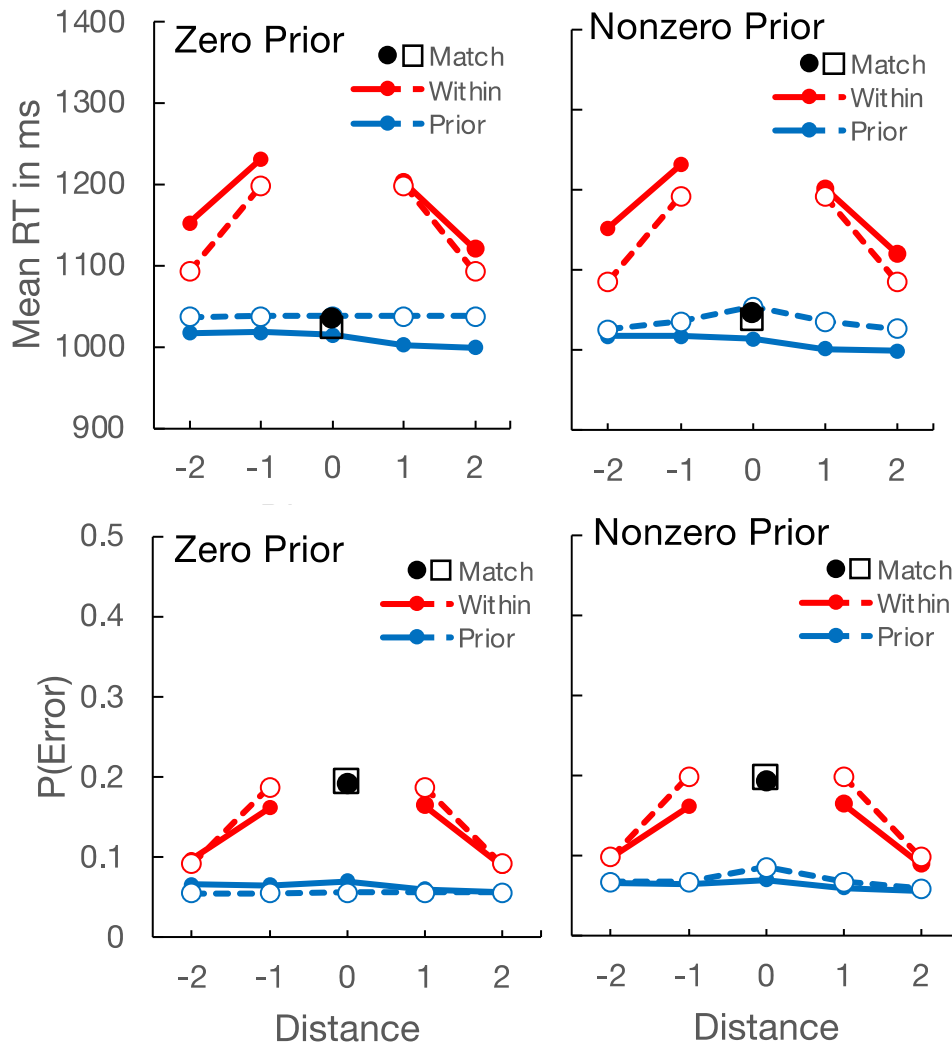
Figure 7: Observed and Predicted Performance Across Experiments 3-12

Figure 7 caption: Observed and predicted performance from the zero prior list strength model (left panels) and the nonzero prior list strength model (right panels) across Experiments 3-12. Solid lines and filled circle: Observed mean RT (top) and error rate (P(Error), bottom) across all 320 subjects in the cued recognition experiments (3-12) for match trials (circle), within-list lures (red), and prior-list lures (blue) as a function of their distance from the cued position. The observed data are repeated in the left and right panels to illustrate fits of different models. Dashed lines and empty square: Predicted mean RT and P(Error) for the zero prior list strength model (left panels) and the nonzero prior list strength model (right panels).

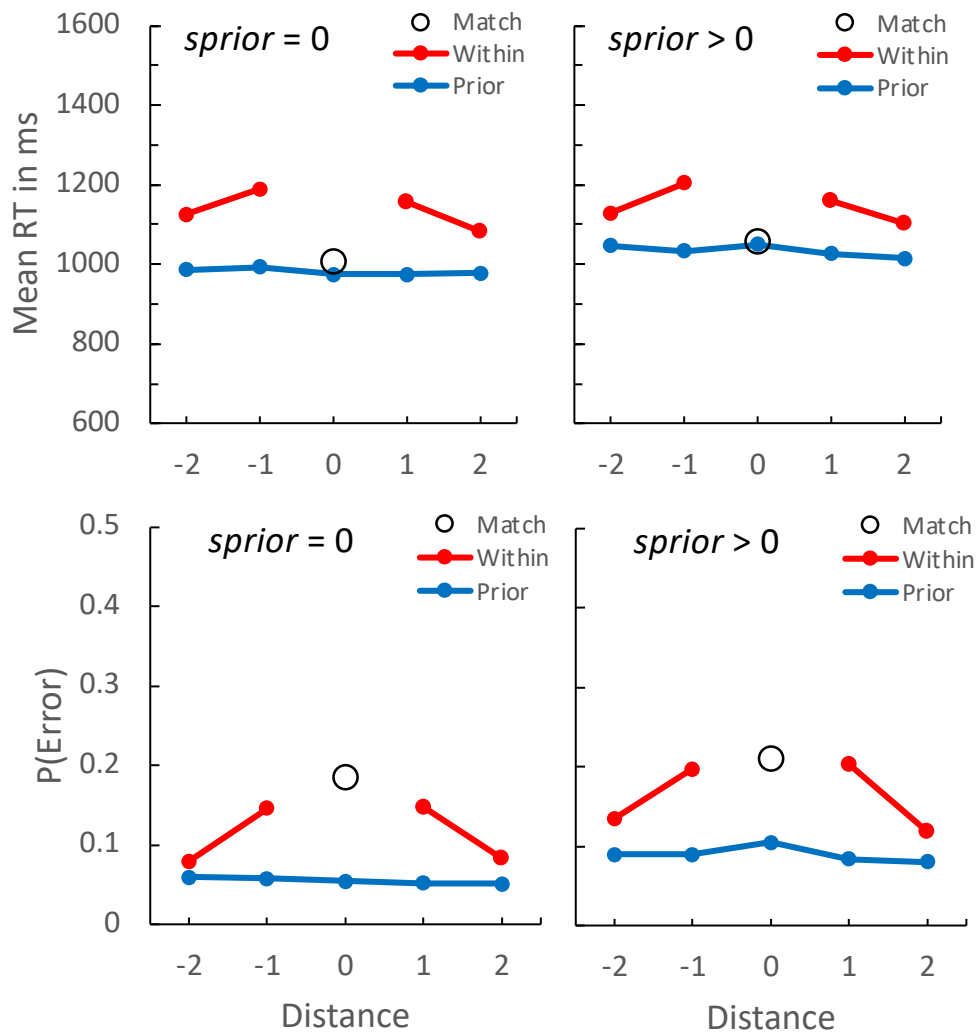
Figure 8: Differences in Prior List Strength

Figure 8 caption: Mean RT (top panels) and error rate (P(Error), bottom panels) for subjects with estimated prior list strength parameters equal to zero (left panels) and greater than zero (right panels).

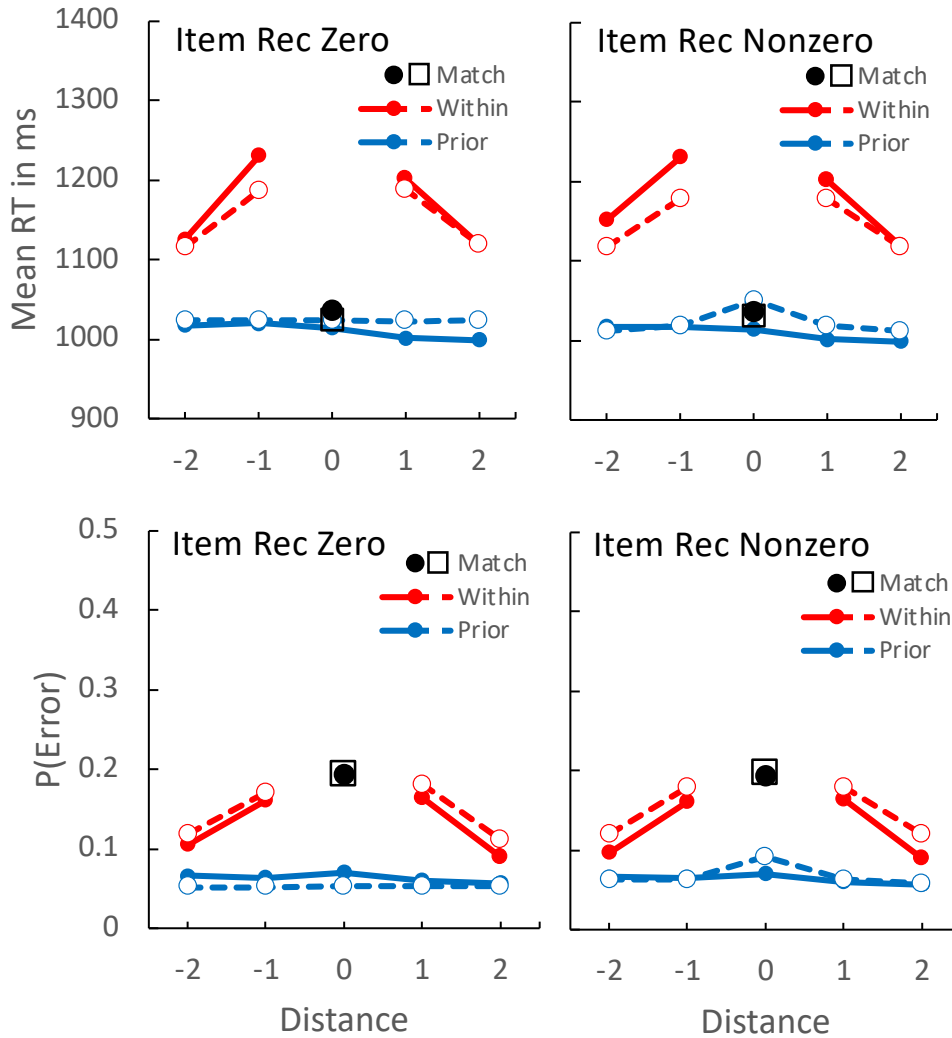
Figure 9: Observed and Predicted Performance Across Experiments 3-12

Figure 9 caption: Observed and predicted performance from the item recognition model with prior list strength = 0 (left panels) and the item recognition with prior list strength > 0 (right panels) across Experiments 3-12. Solid lines and filled circle: Observed mean RT (top) and error rate (P(Error), bottom) across all 320 subjects in the cued recognition experiments (3-12) for match trials (circle), within-list lures (red), and prior-list lures (blue) as a function of their distance from the cued position. The observed data are repeated in the left and right panels to illustrate fits of different models. Dashed lines and empty square: Predicted mean RT and P(Error) for the item recognition only model (left panels) and the item recognition plus prior list strength model (right panels).

Figure 10: Prior List Intrusions in Typing, Recall, and Report

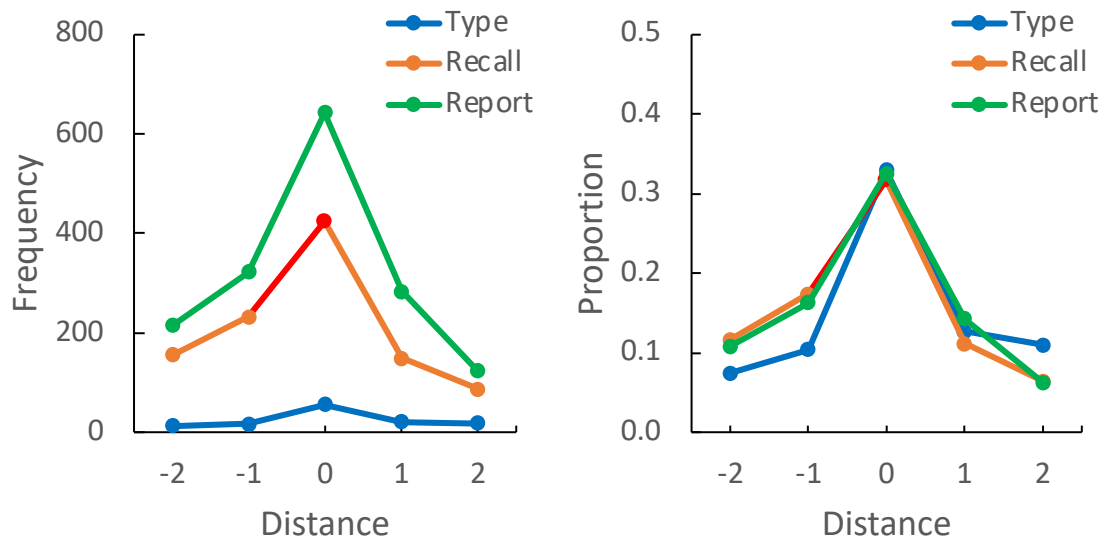


Figure 10 caption: Prior list intrusions from copy typing, serial recall, and whole report tasks from Logan (2021). The left panel contains frequency counts of the number of intrusions across list lengths (5, 6, 7 letters) and subjects ($N = 24$). The right panel converts the frequencies to proportions of the total number of prior list intrusions.

Figure A1: Using Position Coding Probabilistically

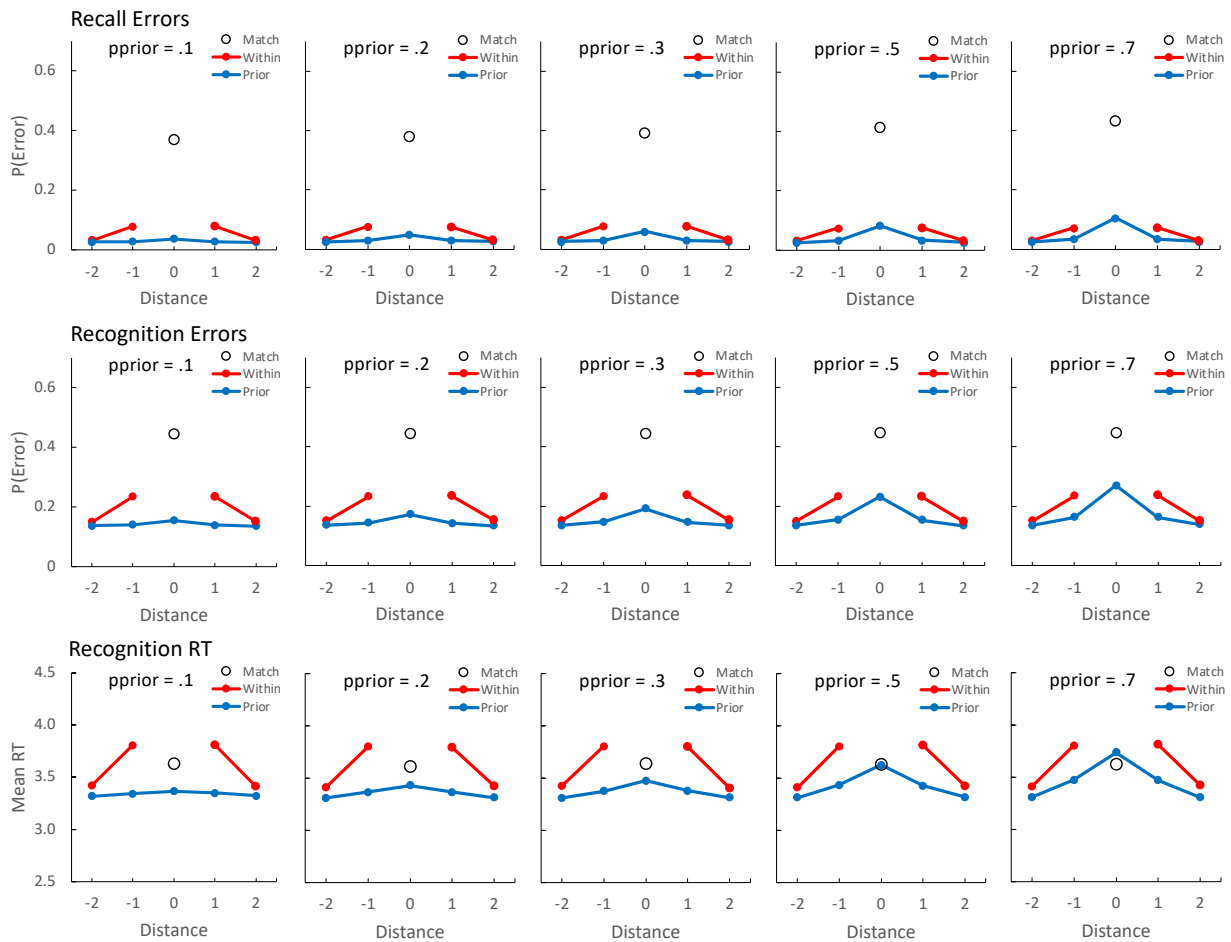


Figure A1 caption: Simulated predictions of within- and prior-list distance effects in response time (RT) and response probability from the position coding model. The same representations of position are used in each panel. The columns represent different probabilities ($pprior$) of using position coding to represent lists. The top row presents serial and cued recall error rates, the middle row presents cued recognition task error rates, and the bottom row presents cued recognition response times (RT). Prior list distance effects are observed in recall and cued recognition for $pprior$ values greater than or equal to 0.2.

Figure C1: Serial Position Curves

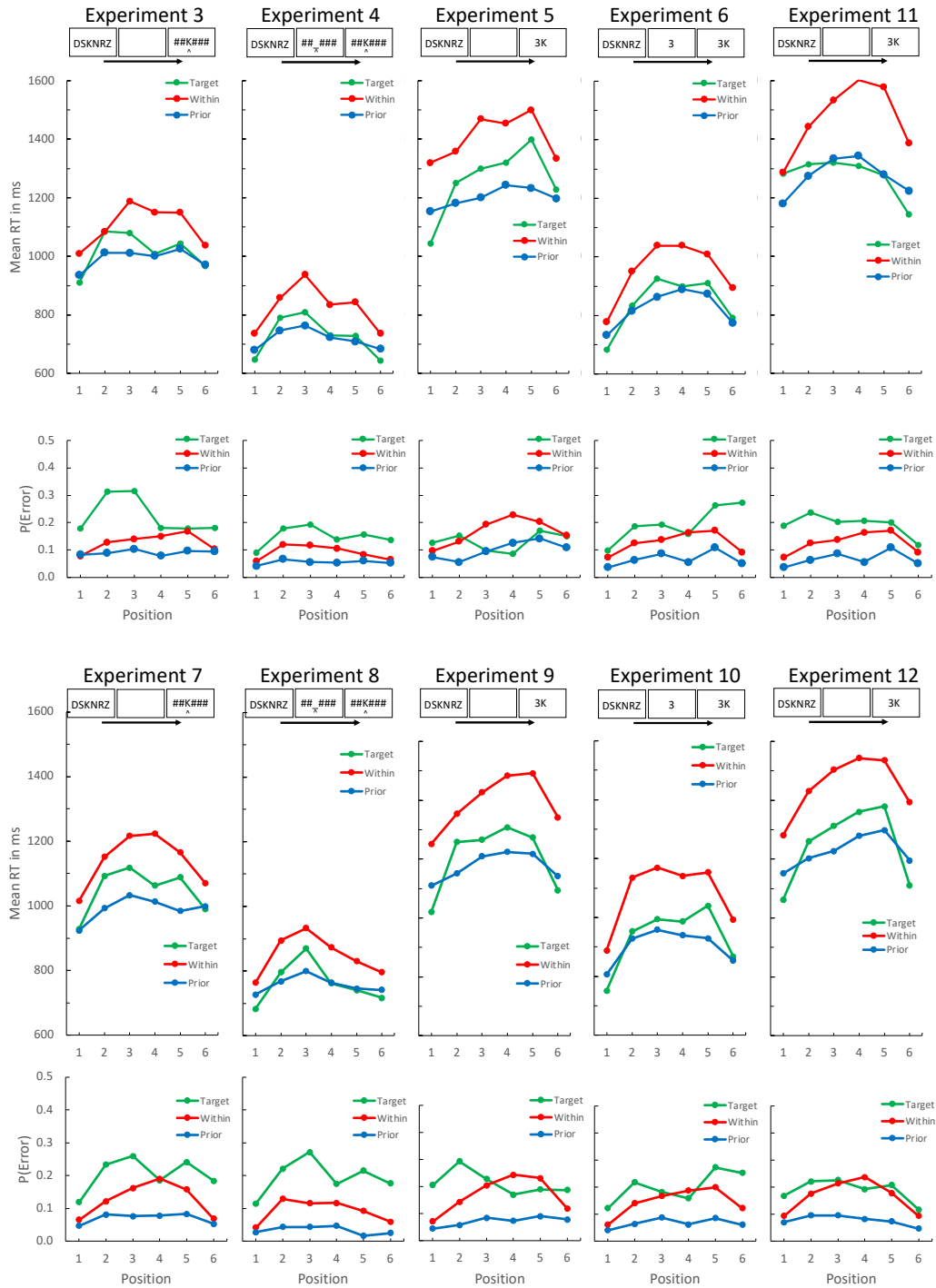


Figure C1 caption: Mean RT (rows 1 and 3) and mean error rate (rows 2 and 4) for targets (match), within-list lures (within), and prior-list lures (prior) as a function of serial position in Experiments 3-12. Rows 1 and 2 show data from constrained lists. Rows 3 and 4 show data from unconstrained lists.

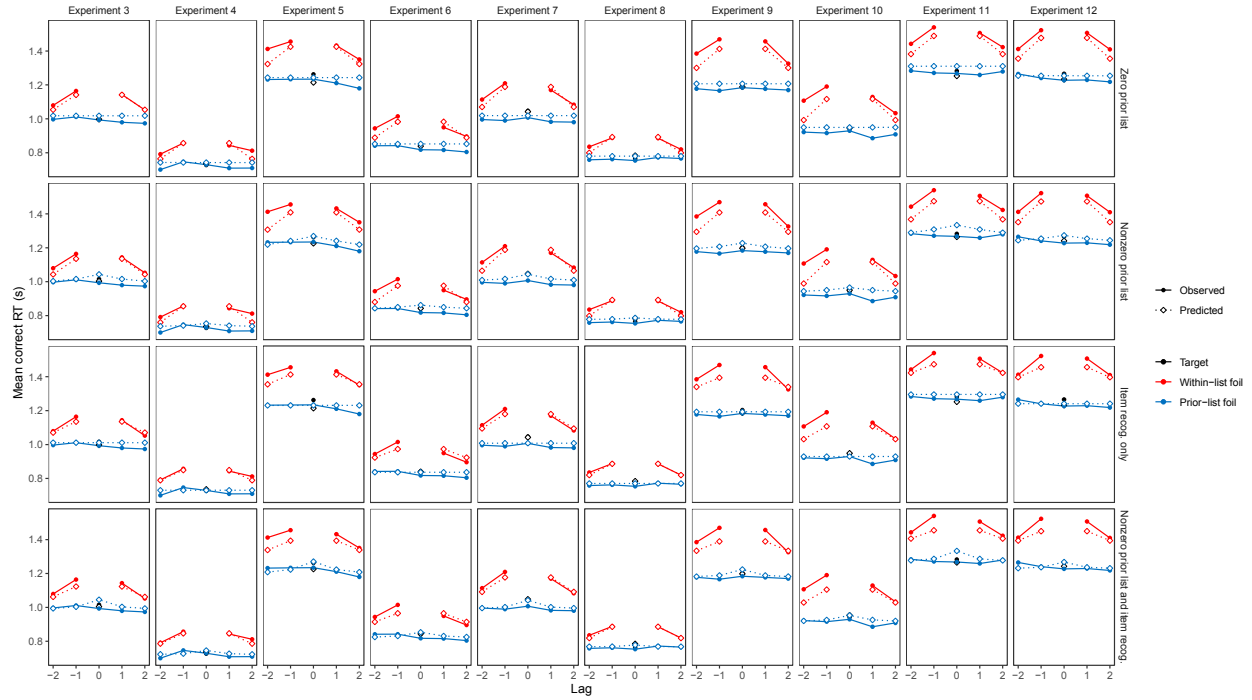
Figure D1: Observed and Predicted RTs

Figure D1 caption: Observed (solid lines) and predicted (dashed lines) response times (RTs) in Experiments 3-12 (columns) for each model (rows).

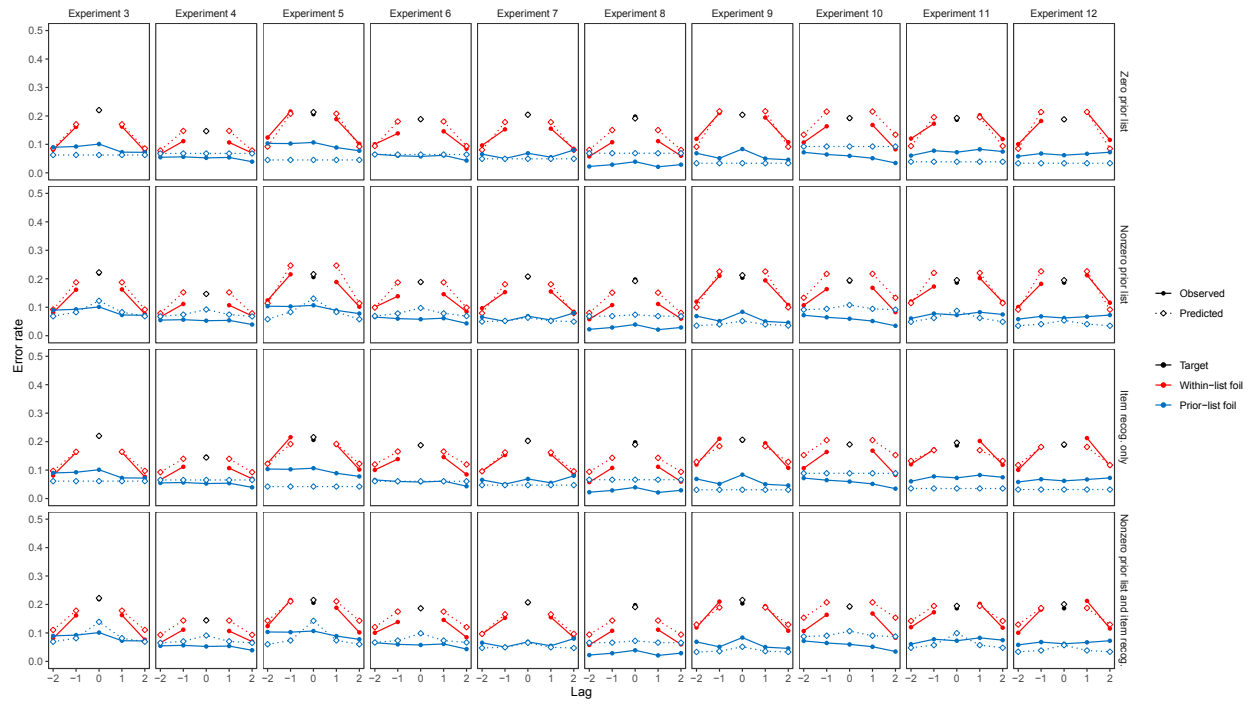
Figure D2: Observed and Predicted Error Rates

Figure D2 caption: Observed (solid lines) and predicted (dashed lines) error rates in Experiments 3-12 (columns) for each model (rows).

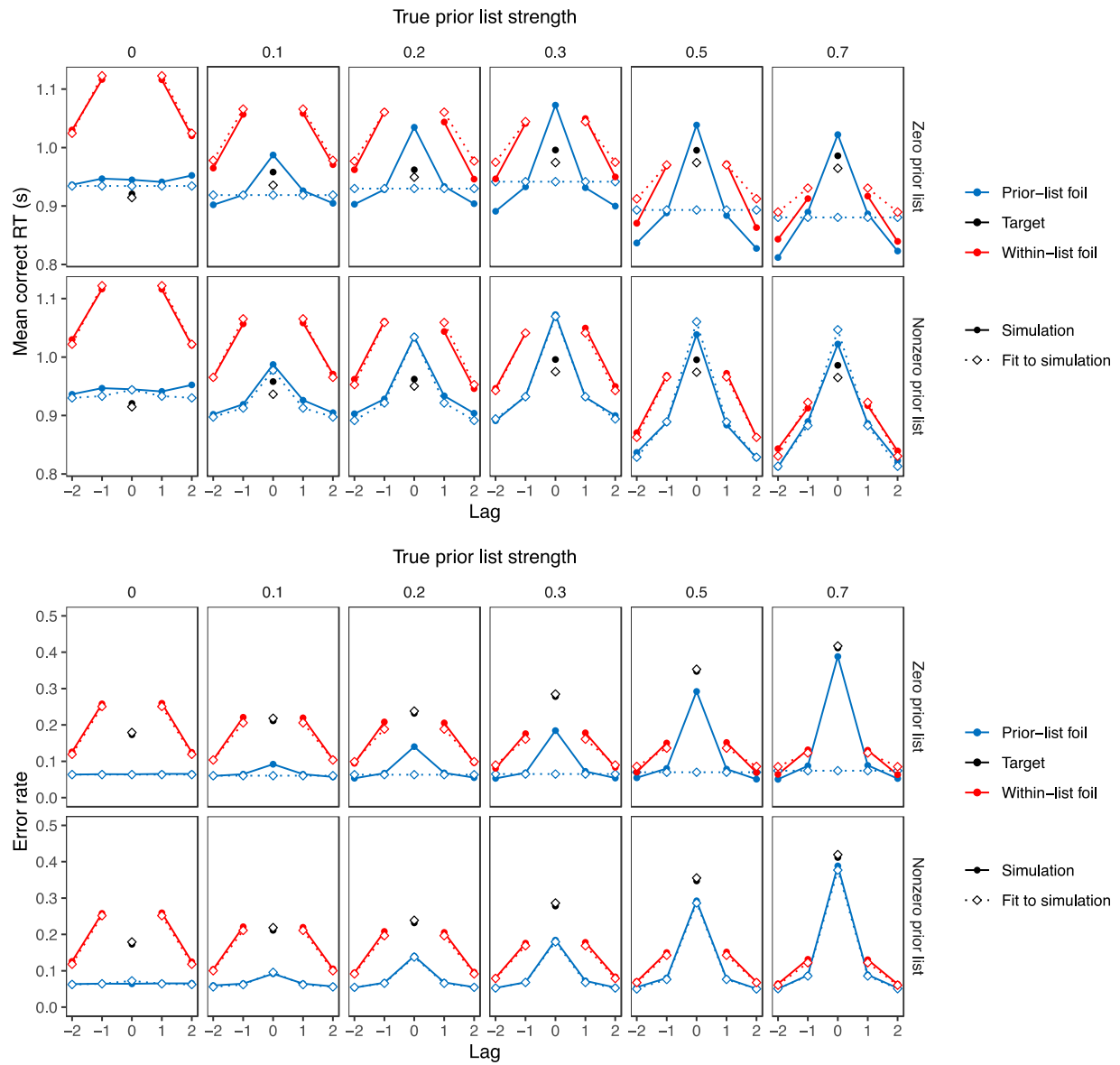
Figure E1: Simulated RTs and Error Rates and Fits to Simulated RTs and Error Rates

Figure E1 caption: Mean simulated RTs (top pair) error rates (bottom pair) across different probe types at different lags (solid lines), and mean predicted error rates (dashed lines) from models fit to simulated data. Each column represents a different value of the *sprior* parameter, representing “true” prior list strength, used to generate the simulated data in each column. The top row in each pair shows fits of the model constrained to have zero prior list strength (i.e., the estimated value of *sprior* was constrained to be zero). The bottom row in each pair shows fits of the model allowing for nonzero prior list strength (i.e., *sprior* was a free parameter).

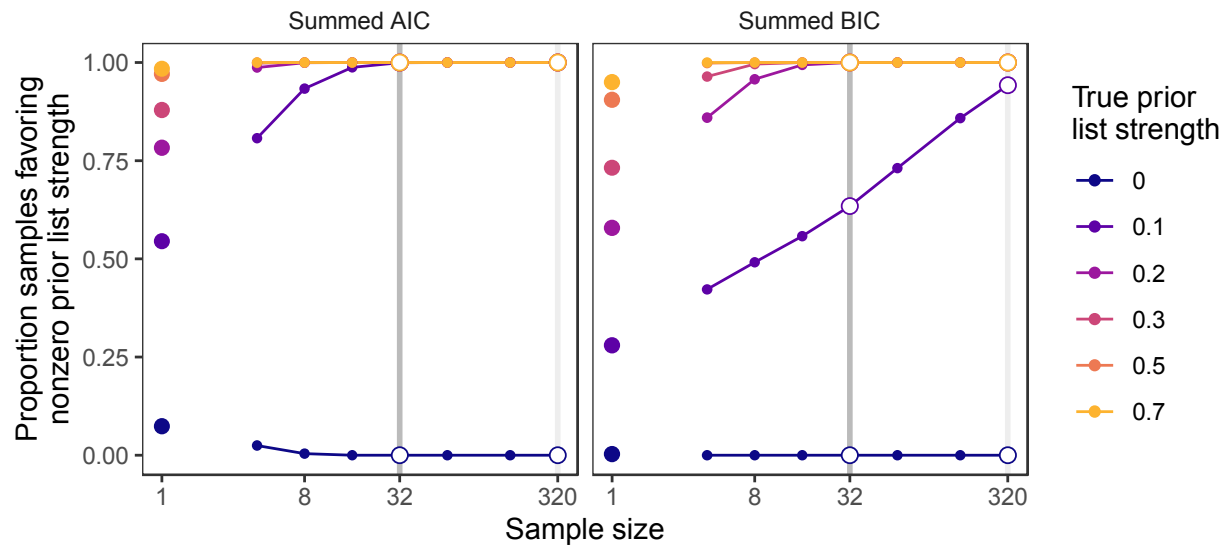
Figure E2: Proportion of Simulated Samples Favoring Nonzero Prior List Strength

Figure E2 caption: Each point corresponds to 10000 simulated samples of each size and gives the proportion out of those 10000 simulated samples in which AIC (left panel) or BIC (right panel) summed across all subjects in each simulated sample favors the unconstrained model that allows nonzero prior list strength. Highlighted sample sizes at 32 and 320 correspond to the sample size for each of the 10 cued recognition experiments in the main text (each of which had 32 subjects) as well as the sample size that would result from aggregating across all experiments (320 subjects total).