



Original Article

A general instance-based learning framework for studying intuitive decision-making in a cognitive architecture



Robert Thomson*, Christian Lebiere, John R. Anderson, James Staszewski

Carnegie Mellon University, United States

ARTICLE INFO

Article history:

Received 22 October 2013

Accepted 11 June 2014

Available online 19 June 2014

Keywords:

Cognitive architecture

Cognitive modeling

Decision-making

Heuristics

ABSTRACT

Cognitive architectures (e.g., ACT-R) have not traditionally been used to understand intuitive decision-making; instead, models tend to be designed with the intuitions of their modelers already hardcoded in the decision process. This is due in part to a fuzzy boundary between automatic and deliberative processes within the architecture. We argue that instance-based learning satisfies the conditions for intuitive decision-making described in Kahneman and Klein (2009), separates automatic from deliberative processes, and provides a general mechanism for the study of intuitive decision-making. To better understand the role of the environment in decision-making, we describe biases as arising from three sources: the mechanisms and limitations of the human cognitive architecture, the information structure in the task environment, and the use of heuristics and strategies to adapt performance to the dual constraints of cognition and environment. A unified decision-making model performing multiple complex reasoning tasks is described according to this framework.

© 2014 Society for Applied Research in Memory and Cognition. Published by Elsevier Inc. This is an open access article under the CC BY-NC-ND license (<http://creativecommons.org/licenses/by-nc-nd/4.0/>).

This article describes how computational models of intuitive decision-making are expressed within the constraints of the *ACT-R cognitive architecture* (Anderson et al., 2004). These models are noteworthy for their ability to explain a variety of heuristics and biases in terms of the processes and representations that produce them. These phenomena have largely been captured and defined as results of experimental manipulations (Kahneman & Tversky, 1996) but not in terms of process models justified by a cognitive architecture (Dimov, Marewski, & Schooler, 2013). A concern of modeling intuitive decision-making behavior using cognitive architectures is confounded by the explicit decisions encoded by the modelers. This criticism can be described as: instead of modeling intuitive behavior per se, cognitive models make explicit the intuitions of their designers (Cooper, 2007; Lewandowsky, 1993; Schultheis, 2009; Cooper, 2007; Lewandowsky, 1993; Schultheis, 2009; Shallice & Cooper, 2011). We address this criticism by showing that the instance-based learning mechanisms in the ACT-R cognitive architecture (Gonzalez, Lerch, & Lebiere, 2003) exhibit the characteristics of intuitive decision-making as described in Kahneman and Klein, (2009), and provide a clearer distinction between automatic and implicit (System 1) processes and

deliberative and explicit (System 2) processes. In addition, we specifically address this *modeler selection* criticism by showing that the explicit strategies of the models instantiate the theories of the model designer and thus are a mechanism for theory evaluation rather than a confounding factor in model development.

In making this argument, we recommend adopting a tripartite explanation of decision-making and biases that illustrates the critical role of the task environment in the decision-making process. We argue that decision-making should be understood in terms of: (1) the mechanisms and limitations of the architecture; (2) the information structure in the task environment; and (3) the use of heuristics and strategies to adapt performance to the dual constraints of cognition and environment. From examples of existing models, we show that simulating behavior within a cognitive architecture is a useful methodology for the study of the mechanisms, variables, and time-course in complex decision-making processes that are impossible in experimentation due to exploding combinatorics.

1. What is intuitive decision-making?

Simon (1992) characterized intuitive decision-making skill as “nothing more and nothing less than recognition” (p. 155). In their seminal work on expertise, Chase and Simon (1973) identified that chess experts require upwards of a decade of study

* Corresponding author. Tel.: +1 4126890536.
E-mail addresses: thomsonr@andrew.cmu.edu,
noblephoenix@gmail.com (R. Thomson).

to retain 50,000–100,000 distinct and rapidly accessible patterns of chess positions. Intuitive decision-making has been studied in both the *naturalistic decision-making* and *heuristics and biases* literature, with the former generally focused on the successes of intuitive reasoning, while the latter generally focused on its failures (Kahneman & Klein, 2009). A distinguishing feature of intuitive decision-making is that a single plausible solution rapidly ‘comes to mind’ in its entirety without explicit or conscious awareness of the causal factors entering into the decision (i.e., not being consciously derived in a piecemeal, step-by-step, or in a ‘deliberative’ manner; Newell & Simon, 1972; Simon, 1995). As such, intuitive reasoning is considered System 1. For example, Klein, Calderwood, and Clinton-Cirocco (1986) found that fire marshals tended to make rapid decisions by generating a single alternative, mentally simulating its outcome, and either making minor revisions or adopting the next closest alternative. Effectively, fire marshals were pattern-matching based on their prior experiences. This strategy has been termed *recognition-primed decision-making*.

Conversely, deliberative decision-making is often characterized as strategic, effortful, slow, and rule-oriented (Klein, 1998), and as such is considered System 2 thinking (Kahneman & Frederick, 2005; Stanovich & West, 1999). Interestingly, the act of verifying an intuition is generally seen as optional, effortful, and thus a function of System 2 (Kahneman & Klein, 2009).

In order to gain intuitive expertise, two conditions first need be met. The first condition is that people receive extensive practice in a task environment that is sufficiently stable and provides causal or statistical cues/structures that may at least theoretically be operationalized (Hogarth, 2001; Hogarth, 2001; Brunswik, 1957). This need not be deterministic (e.g., playing poker is a probabilistic but stable environment; Kahneman & Klein, 2009). The second condition is that there must be sufficient feedback from the task environment which provides people an opportunity to learn the relevant cues/structures. In other words, feedback must be sufficient to generate a relevant internal problem space. This requirement of feedback and interaction with the task environment drove our adoption of the tripartite level of description.

2. Why use a cognitive architecture?

Cognitive architectures model behavior using a set of common mechanisms and processes (i.e., *the architecture*) whose goal is to not only explain human behavior, but the underlying structures and representations subsuming cognition as a whole. These mechanisms should be both psychologically and neurally plausible to account for human behavior. This level of description is not generally captured by either mathematical or informal models of decision-making. Before getting into further details of mechanisms, models, and results; there is an important argument to be made for the role of cognitive architectures in general, which is best characterized by Herbert Simon (in 1971, no less):

The programmability of the theories is the guarantor of their operability, an iron-clad insurance against admitting magical entities into the head. A computer program containing magical instructions does not run, but it is asserted of these information-processing theories of thinking that they can be programmed and will run. They may be empirically correct theories about the nature of human thought processes or empirically invalid theories; [but] they are not magical theories. (p. 148)

In modern terms, simulations using a cognitive architecture provide a falsifiable methodology for the study of cognitive processes and representations, a particularly important characteristic when studying largely implicit processes such as intuitive decision-making. They serve several theoretical functions including:

organizing and relating a substantial number of cognitive mechanisms, making testable predictions, and explaining the cognitive processes underlying human performance. In many cases, cognitive models can perform tasks too complex to analyze with traditional experimentation due to the combinatorics of the possible decision space. As will be described, a single ACT-R model has explained anchoring and adjustment, confirmation, and probability matching biases across a range of complex geospatial intelligence tasks using a common instance-based learning approach (Lebiere et al., 2013). Similarly, Marewski and Mehlhorn (2011) were able to specify 39 process models studied in decision-making using a smaller subset of 5 ACT-R models. In short, cognitive architectures allow for theories to be constrained by scientifically established mechanisms and (hopefully) easily describable processes.

This is not to argue that cognitive architectures are a panacea for studying decision-making (or psychology in general), but we do claim that they are a valuable tool in the generation and exploration of theories (c.f., models) which may be too complex for traditional piecemeal experimental methods. In particular, intuitive decision-making tends to be cognitively ‘opaque’ with little observable evidence, and what little evidence there is coming from highly fallible introspection. As such, many descriptions of intuitive decision-making are inherently qualitative or are characterized using relatively simple experimental results (Dimov et al., 2013). An advantage of cognitive architectures is not only their ability to objectively explain accuracy and response times in terms of the operation of both symbolic elements and their sub-symbolic activation strengths (and in the case of ACT-R, links to neural structure), but also the ability to go ‘under the hood’ and actually *look inside the model to explicitly examine causal processes*. Such computational cognitive models make testable predictions of what is going on *inside the mind* of someone performing intuitive decision-making.

One measure for validating *inside the mind* predictions is to perform model tracing. Model tracing is a technique where a model is forced to respond with some or all of the same values as a human participant, and then the internal states of the model are examined to determine the influence of these ‘forced’ decisions. By examining the commonalities between the model’s internal states and human behavior, modelers are potentially able to make causal claims about the nature of mental processes within participants; that is, to explain how human performance is produced by various cognitive mechanisms and their interaction. This performance includes traditional measures such as accuracy and response time, but also predictions of fMRI bold response for specific brain areas associated with the functional modules of the cognitive architecture (Anderson, 2007).

The benefits of cognitive architectures can be seen as bridging or synthesizing formal mathematical theories (such as Bayesian modeling) and knowledge-level strategies (e.g., heuristics). As such, cognitive architectures act as a link between Marr’s (1982) computational and algorithmic levels, with the benefits of a corresponding bridge to the physical level (i.e., neural) implementation. Bayesian models belong to a broad class of abstract models that formally (i.e., mathematically) explain human behavior in terms of processes computing probabilities over a set of possible decisions. While Bayesian (and related probabilistic) models do provide an explanation of behavior, it is not generally accepted to be a cognitively (i.e., psychologically) plausible one as the underlying mechanisms driving the processes are somewhat vague or not tractable (Bowers & Davis, 2012). As such, Bayesian theories belong at the computational level of Marr’s hierarchy. This is not a criticism specific to Bayesian models, but can also be applied to other mathematical theories such as prospect theory (Kahneman & Tversky, 1979), decision theory (Berger, 1985), and quantum probability theory (Busemeyer, Pothos, Franco, & Trueblood, 2011). Similarly, explanations in the

form of heuristics (i.e., knowledge-level explanations) tend to be vague as to the underlying processes leading to the biased behavior. To give another example, while fast-and-frugal heuristics (see Gigerenzer, Todd, & the ABC Research Group, 1999) are explained in terms of process models, it is contested whether the underlying cognitive mechanisms required are psychologically plausible (e.g., Dougherty, Franco-Watkins, & Thomas, 2008; but see Gigerenzer, Hoffrage, & Goldstein, 2008).

We would like to briefly discuss what some consider are limitations of the cognitive architecture approach, but we maintain that these limitations are outweighed by the advantages of this method of theorizing. Perhaps the main criticism of cognitive architectures is the *degrees of freedom* argument (e.g., Roberts and Pashler, 2000; Griffiths, Chater, Norris, & Pouget, 2012). This argument relies on the notion that the numerous parameters of the architecture (and the extensibility of the parameter space) add unnecessary complexity while offering little constraint on the kinds of models that may be generated to solve a given problem. While it is true that cognitive architectures such as ACT-R are parameterized, these parameters are well-documented, transparent in the model, and are generally dependent on each other such that it limits their capability to over-fit to the data (i.e., they have a constrained effect on data). Furthermore, this same degrees of freedom argument can and has been applied to probabilistic models (Bowers & Davis, 2012). In fact, we argue that informal theories also fall prey to the same degrees of freedom criticism (Griffiths et al., 2012), however, these theories benefit from their lack of formalization by having less obvious points of criticism (i.e., no obvious parameters to point at). One other concern with informal theories is that they tend toward binary oppositions (e.g., System 1 vs System 2 or explicit vs implicit processes) without a clear understanding of where the boundaries lie or how they may be behaviorally or neurally instantiated within the mind (see Kruglanski & Gigerenzer, 2011 for a similar argument). In our opinion, this lack of formalization can stifle progress by leading our theorizing toward relatively arbitrary boundaries rather than forcing theorizing to occur in a testable and predictive environment.

To more directly address the degrees of freedom argument, we argue for two techniques to mitigate its concerns: the first is to limit the freedom of architectural parameters through the use of scientifically justified default values, and the second is to develop models which are relatively insensitive to parameter values. For instance, the *base-level learning* (i.e., memory decay) rate of .5 in ACT-R has been justified in over 100 published models (see the ACT-R website: <http://act-r.psy.cmu.edu>). In addition to architectural parameters, there are ways of controlling 'knowledge' parameters (e.g., knowledge representation) through approaches like instance-based learning which use a common knowledge representation determined directly by the interaction with the external environment. Finally, there have been efforts to use a single, more general ACT-R model to perform a series of decision-making tasks (e.g., Lebiere et al., 2013; Marewski & Mehlhorn, 2011; Stewart, West, & Lebiere, 2009; Taatgen & Anderson, 2008, 2010). Using more general models of cognition reduces the degrees of freedom in the architecture by aligning a common set of mechanisms, parameters, and strategies into a cohesive modeling paradigm.

In summary, descriptions at either the formal or informal level are under-constrained. Cognitive architectures tie together both formal and informal levels of description (Gonzalez & Lebiere, 2013) by combining sub-symbolic algorithms akin to formal theories with symbolic knowledge structures controlling those processes in a heuristic manner. As such, they provide not only explanations of existing behavior, but predictions of the mechanisms, knowledge, and behaviors of future actions based on a model's prior experience (Thomson & Lebiere, 2013).

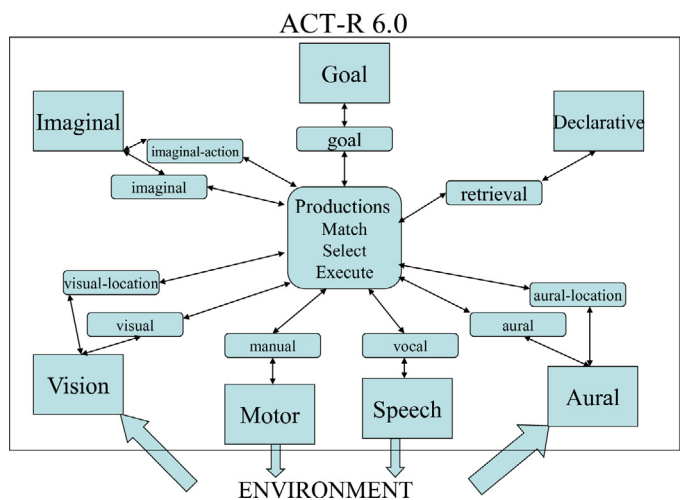


Fig. 1. An overview of ACT-R's default modules and their dependent buffers.

3. Why use ACT-R?

3.1. What is ACT-R?

The following is a brief technical description of ACT-R, and while it is helpful to understand our later arguments, readers not interested in technical descriptions should feel free to skip this section. For those interested in a more in-depth description of ACT-R, please see Anderson et al. (2004) and the ACT-R website: <http://act-r.psy.cmu.edu/>.

ACT-R is a cognitive architecture defined as a set of modules which are integrated and coordinated through a centralized production system (see Fig. 1). Each module is assumed to access and deposit information into buffers associated with the module, and the production system only responds to the contents of the buffers, not the internal processing of the modules. The declarative memory and production system modules, respectively, store and retrieve information that corresponds to *declarative knowledge* and *procedural knowledge*. Declarative knowledge is the kind of knowledge that a person can attend to, reflect upon, and usually articulate in some way (e.g., by declaring it verbally or by gesture). Procedural knowledge consists of the skills we display in our behavior, generally without conscious awareness.

Declarative knowledge in ACT-R is represented formally in terms of *chunks*. The information in declarative memory corresponds to episodic and semantic knowledge that promotes long-term coherence in behavior. Chunks have an explicit type, and consist of slot-value pairs of information (see Fig. 2). Chunks are retrieved from long-term declarative memory by an activation process. Each chunk has a base-level activation that reflects its recency and frequency of occurrence. Activation spreads from the current focus of attention through associations among chunks

```
(chunk-type player name value mission)
```

```
Red-Player
  ISA      player
  name     aggressive
  value    .8
  mission  2
```

Fig. 2. An example of a chunk in ACT-R. The name of the chunk is Red-Player, it is of type *player*, and has three slots: *name*, *value* and *mission* containing values *aggressive*, *.8* and *2*, respectively. This chunk represents the name of a hypothetical opponent who the model assumes is aggressive with a probability of .8. In this example, the value of .8 is derived by the model based on evidence accumulated during the task and is derived by the proportion of trials in which attacks occurred.

Table 1

The list of sub-symbolic mechanisms in the ACT-R architecture.

Mechanism	Equation	Description
Activation	$A_i = B_i + S_i + P_i + \varepsilon_i$	B_i : Base-level activation reflects the recency and frequency of use of chunk i S_i : Spreading activation reflects the effect that buffer contents have on the retrieval process P_i : Partial matching reflects the degree to which the chunk matches the request ε_i : Noise value includes both a transient and (optional) permanent component (permanent component not used by the integrated model)
Base-level	$B_i = \ln \left(\sum_{j=1}^n t_j^{-d} \right) + \beta_i$	n : The number of presentations for chunk i t_j : The time since the j th presentation d : A decay rate (not used by the integrated model) β_i : A constant offset (not used by the integrated model)
Spreading activation	$S_i = \sum_k \sum_j W_{kj} S_{ji}$ $S_{ji} = S - \ln(fan_{ji})$	k : Weight of buffers summed over are all of the buffers in the model j : Weight of chunks which are in the slots of the chunk in buffer k W_{kj} : Amount of activation from sources j in buffer k S_{ji} : Strength of association from sources j to chunk i S : The maximum associative strength (set at 4 in the model) fan_{ji} : A measure of how many chunks are associated with chunk j
Partial matching	$P_i = \sum_k PM_{ki}$	P : Match scale parameter (set at 2) which reflects the weight given to the similarity M_{ki} : Similarity between the value k in the retrieval specification and the value in the corresponding slot of chunk i The default range is from 0 to -1 with 0 being the most similar and -1 being the largest difference
Declarative retrievals	$P_i = \frac{e^{A_i/s}}{\sum_j e^{A_j/s}}$	P_i : The probability that chunk i will be recalled A_i : Activation strength of chunk i $\sum A_j$: Activation strength of all of eligible chunks j s : Chunk activation noise
Blended retrievals	$V = \arg \min \sum P_i (1 - Sim_{ij})^2$	P_i : Probability from declarative retrieval Sim_{ij} : Similarity between compromise value j and actual value i
Utility learning	$U_i(n) = U_i(n-1) + \alpha [R_i(n) - U_i(n-1)]$ $P_i = \frac{e^{U_i/s}}{\sum_j e^{U_j/s}}$	$U_i(n-1)$: Utility of production i after its $n-1$ st application $R_i(n)$: Reward production receives for its n th application $U_i(n)$: Utility of production i after its n th application P_i : Probability that production i will be selected U_i : Expected utility of the production determined by the utility equation above U_j : is the expected utility of the competing productions j

in declarative memory. These associations are built up from experience, and they reflect how chunks co-occur in cognitive processing. Chunks are compared to the desired retrieval pattern using a partial matching mechanism that subtracts from the activation of a chunk its degree of mismatch to the desired pattern, additively for each component of the pattern and corresponding chunk value. Noise is added to chunk activations to make retrieval a probabilistic process governed by a Boltzmann (softmax) distribution.

While the most active chunk is usually retrieved, a blending process (i.e., a *blended retrieval*; see Lebiere, 1999; Wallach & Lebiere, 2003a) can also be applied that returns a derived output reflecting the similarity between the values of the content of all chunks, weighted by their retrieval probabilities reflecting their activations and partial-matching scores (see Table 1 for a list of sub-symbolic activations involved in chunk retrieval and production selection). This process enables not just the retrieval of previously encountered symbolic values but also the generation of continuous values such as probability judgments in a process akin to weighted interpolation.

Production rules are used to represent procedural knowledge in ACT-R. They specify procedures that represent and apply cognitive skill in the current context, including how to retrieve and modify information in the buffers and transfer it to other modules. In ACT-R, each production rule is a set of conditions and actions which are analogous to an IF-THEN rule. Conditions specify structures that are matched in buffers, and correspond to information from the external world or other internal modules. Actions represent requests and modifications to the contents of the buffers, including queuing perceptual-motor responses (e.g., speaking, typing, or looking to given location). Matching production rules effectively means: if the conditions of a given production match the current state of

affairs (i.e., the state of the modules and contents of the buffers) then perform the following actions (see Fig. 3).

ACT-R uses a mix of parallel and serial processing. Modules are encapsulated and may process information in parallel within one another. However, there are two serial bottlenecks in processing. First, only one production may be executed at a time. Second, a buffer can only contain one chunk at a time. In general, multiple production rules can match, but only one can be active – in ACT-R parlance *fired* – at any point. *Production utilities*, learned using a

(p	get-feedback	English Description
=goal>		If the goal chunk is
isa	feedback	of the type <i>feedback</i>
value	=v	there is a <i>value</i> we will call =v
name	=name	there is a <i>name</i> we will call =name
mission	=m	there is a <i>mission</i> we will call =m
?retrieval>		If the retrieval module has
state free		a state of free (is not busy)
=>		Then
=goal>		Change the goal
name	nil	to set the value of <i>name</i> to <i>nil</i>
+retrieval>		and request a retrieval
isa	player	of a <i>player</i> chunk
name	=name	with a <i>name</i> of =name
value	=v	with a <i>value</i> of =v
mission	=m	with a <i>mission</i> if =m
)		

Fig. 3. An example of a production in ACT-R. The name of the production is get-feedback, and it tests the goal buffers (the = sign) and queries (the ? sign) the state of the retrieval module on the left-hand side of the equation (everything before the =>). The production also modifies the chunk in the goal buffer (the = sign) and makes a request (the + sign) for a new chunk to be placed in the retrieval buffer on the right-hand side of the equation.

reinforcement learning scheme, are used to select the single rule that fires. As for declarative memory retrieval, production selection is a probabilistic process. Based on experience and matching certain criteria, two production rules may be automatically compiled together into a new and more-efficient rule, which accounts for proceduralization of behavior.

3.2. What does ACT-R have to do with intuitive decision-making?

Cognitive model development in ACT-R is in part derived from the rational analysis of the task (Anderson, 1982) and information structures in the external environment (e.g., the design of the tasks being simulated), the constraints of the ACT-R architecture, and guidelines from previous models of similar tasks (Taategen, Lebiere, & Anderson, 2006). A successful *design pattern* in specifying cognitive process sequencing in ACT-R is to decompose a complex task to the level of *unit tasks*. Card, Moran, and Newell (1983) suggested that unit tasks control immediate behavior. Unit tasks empirically take about 10 s. To an approximation, the structure of behavior above the unit task level largely reflects a rational structuring of the task within the constraints of the environment, whereas the structure within and below the unit task level reflects cognitive and biological mechanisms, in accordance with Newell's (1990) bands of cognition. Accordingly, in ACT-R, unit tasks are implemented by specific *goal types* that control a set of productions which represent the cognitive skills for solving those tasks.

There are a broad range of ACT-R models studying problem solving, decision-making (including intuitive decision-making; Kennedy & Patterson, 2012), and implicit learning (see <http://act-r.psy.cmu.edu/publications> for examples of each; also, see Anderson (2007) and Lebiere and Anderson (2011) for an overview). Specific examples (all using instance-based learning) include a model of how batters predict baseball pitch speed (Lebiere, Gray, Salvucci, & West, 2003), a model predicting risk aversion in a repeated binary choice task (Lebiere, Gonzalez, & Martin, 2007), a model of sequence learning (Lebiere & Wallach, 2001), and a model of playing Paper Rock Scissors (West & Lebiere, 2001). Another model of repeated binary choice also won the Technion Prediction competition over machine-learning algorithms (Stewart et al., 2009). These models all work by storing problem-solving instances in declarative memory, then they make decisions by retrieving those instances by leveraging the cognitive architecture's activation processes to extract regularities in the task environment.

There is a misconception that intuitive processes in ACT-R – as implicit – are governed using only *procedural* memory processes, while deliberative processes – as explicit – are governed by only *declarative* memory processes (this implicit/explicit distinction has been mistakenly attributed to Wallach & Lebiere, 2003b). In fact, while each declarative chunk is usually considered a piece of conscious knowledge, the sub-symbolic activations that control the retrieval process (e.g., base-level activations and strengths of associations) are consciously inaccessible and constitute the implicit knowledge of the model (Gonzalez & Lebiere, 2005; Lebiere, Wallach, & Taategen, 1998). In essence, the activation calculus involved in retrieving a chunk is the implicit part of the declarative system, while the contents of the chunk itself are the explicit part of the declarative system.

An interesting interplay between System 1 and System 2 processes occurs during a retrieval request. When a production makes a retrieval request it specifies the type of chunk to retrieve and potentially a set of slot-value pairs from which to match, which is essentially the *specification* of what to retrieve. While the production system is generally seen as an implicit (System 1) process, the constraints in matching the retrieval request come from explicitly setting which slot-value pairs to match against. Since this is something coded by the modeler, it could be argued that it is a

totally explicit strategy (i.e., System 2) based on the modelers *intuition* (c.f., theory) of how the retrieval should function. This is a bit of a false dichotomy because every model is a blend of both System 1 and System 2 processes, and the retrieval request links both strategic (e.g., requesting a specific chunk) and implicit processes (e.g., spreading activation). In terms of a cognitive architecture and the *no-magic doctrine* (Anderson & Lebiere, 1998), we argue that the retrieval specification is instead best described as an implicit heuristic (as opposed to a conscious/strategic heuristic), albeit still being effectively the modeler's theory of how the retrieval process should unfold.

It is possible to make retrievals driven more by implicit processes by using a technique that we call 'open' retrievals. A retrieval is considered *open* when only the type of chunk is requested and *no* (or in a relaxed case, *minimal*) slot-value pairs are used in the specification of the retrieval request. An example would be when one is given a set of indirect clues about the identity of a person and the name of the person pops up in one's mind from the convergence of the clues rather than any specific information retrieval process. Effectively, *open* retrievals are a kind of context-driven free association. By using *open* retrievals, the model is relying more on sub-symbolic activations – which are driven by experience – to control the retrieval process. For instance, performing a retrieval by specifying only the context and doing free association on the outcome allows the model to match the best outcome based on the recency and frequency of prior outcomes and spreading activation from the current context. This stands in contrast to specifying a particular outcome in the retrieval request, which is more analogous to the model engaging in a more strategic retrieval strategy.

A similar theme between System 1 and System 2 processes is the nature of heuristics in decision-making. Are heuristics explicit because of their symbolic nature, or implicit because the decision maker is often unaware of them? The choice of which simplifying heuristics are available to the model tends to be a conscious strategy of the modeler (as opposed to being chosen by the model; Lewandowsky, 1993). This may be considered *explicit*, although this is an uncharitable view of the modeler's selection of heuristics/strategies (e.g., the modeler's theory of which heuristics are available; Taategen & Anderson, 2008). It is possible to view task instruction as a kind of heuristic imposed by the task environment, thus by parsing the task instructions the model chooses from a set of (generally) implicit heuristics (e.g., recognition heuristic; fluency heuristic; Marewski & Link, 2014) and performs the appropriate action. Two such theories are the related notions of *cognitive niches* (Marewski & Schooler, 2011) and the *adaptive toolbox* (Gigerenzer & Selten, 2002). A cognitive niche is a simplifying framework which describes how only a limited number of applicable strategies may be considered in a given situation based on the interplay between available strategies, limited human capacities, and the task environment. Similarly, the adaptive toolbox is a psychologically validated set of heuristics from which the model may select. Of course, these solutions still leave open to theory the underlying basis for strategy selection of the model (Marewski, Gaissmaier, Schooler, Goldstein, & Gigerenzer, 2010).

One of the difficulties in providing a more complete answer comes from the traditional dichotomy of automatic versus deliberative (or procedural versus declarative) being insufficient to explain the source of heuristics (mainly from the task environment), which is a key indicator of whether the heuristic should be seen as primarily implicit, explicit, or both.

3.3. How does instance-based learning tie into intuitive decision-making?

Instance-based learning theory (Gonzalez et al., 2003; Taategen et al., 2006) is the claim that implicit expertise is gained through the

accumulation and recognition of experienced events or instances. Unlike instance-based machine learning algorithms (Gagliardi, 2011) that are essentially strict exemplar models of categorization applied to big data (Erickson & Kruschke, 1998), instance-based learning theory allows for generalization and the bootstrapping of learning with weak methods. Weak methods are relatively knowledge-free heuristic methods of action and exploration (such as random choice) that are procedurally driven when there is insufficient domain knowledge (i.e., instances) to make effective decisions. Once enough instances are stored, these weak methods are supplanted by the retrieval of decisions based on these prior instances.

Similar to theories of intuitive expertise (Kahneman & Klein, 2009), instance-based learning theory argues for the necessity of receiving effective feedback. Feedback is required to determine the relative payoffs from not only expected outcomes, but from the actual outcome. This effectively tunes instances to real experiences as opposed to simply existing in the cognitive realm of expectations. The combination of contextual information and the current goal, the selected action, and the outcome of that action result in a common condition → action → outcome representational structure. This structure reflects the necessary requirements for effective learning and subsequent performance. Supporting this structure, Lebiere, Gonzalez, and Warwick (2009) have shown how Klein's (2009) recognition-primed decision-making and instance-based learning use similar mechanisms and make similar predictions in the context of naturalistic decision-making. Instance-based learning, having been formulated within the principles and mechanisms of cognition in ACT-R, makes use of the dynamics of chunk retrieval to recall instances and also makes use of blended retrievals to generalize knowledge. This *instance + generalization* process provides an additional level of explanation and predictive power to complement the process specified in Klein's analysis. As such, recognition-primed decision-making and other similar naturalistic processes can be seen as a macrocognitive substrate that naturally complements the microcognitive mechanisms of a cognitive architecture (Lebiere & Best, 2009).

The main claim of instance-based learning is that implicit knowledge is generated through the creation of instances. These instances are represented in chunks with slots containing the conditions (e.g., a set of contextual cues), the decision made (e.g., an action), and the outcome of the decision (e.g., the utility of the decision). Before there is sufficient task-relevant knowledge, decision-makers implicitly evaluate alternatives using heuristics (e.g., random choice, minimize loss, maximize gain). Once a sufficient number of instances are learned, decision-makers retrieve and generalize from these instances to evaluate alternatives, make a decision, and execute the task. The process of feedback involves updating the outcome slot of the chunk according to the post hoc generated utility of the decision. Thus, when decision-makers are confronted with similar situations while performing a task, they gradually abandon general heuristics in favor of improved instance-based decision-making processes (Gonzalez and Lebiere, 2005).

Comparing instance-based learning with the necessity claims of intuitive decision-making from Klein and Kahneman (2012), both consider intuitive knowledge to be learned via instances. Also, in both cases decisions are made by pattern-matching over prior instances (and/or supplemented by heuristics) and then retrieving the best fit. In the case of instance-based learning, however, this best fit is computed using a generalization across the closest neighbors using partial matching or blended retrievals. Both require the task environment to be sufficiently regular to be able to implicitly learn the statistical correlations between condition, action, and – through either internal or external feedback – outcome. However, instance-based learning offers constraints on explanation by

grounding implicit learning within the mechanisms of a cognitive architecture. For instance, the dynamics of an instance's sub-symbolic activations (e.g., frequency and recency in the base-level learning equation) provide a theoretically grounded mechanism for determining which instances are likely to be retrieved for a given situation, and also can explain *why* they were retrieved and what factors came into play. This provides a much more rigorous explanation of intuitive decision-making than case-studies and introspection of experts.

IBL – as instantiated in ACT-R – also provides for a clearer distinction between automatic and deliberative processes. The act of encoding and retrieving instances is a fully automatic process, guided by either (hopefully *open*) retrieval or implicit heuristics. However, once the retrieval is completed, what the model does with the retrieved chunk (e.g., the structure of the subsequent productions) is an explicit heuristic/strategy. For instance, while the retrieved chunk might provide a recommended action, it is up to the model (through the production system) to determine whether to verify the action, discard the action, perform the action, or simulate possible other outcomes.

In ACT-R, over the past 10 years models related to decision-making and problem-solving have seen increasing use of instance-based learning (whether explicitly referred-to as such or otherwise; e.g., Kennedy & Patterson, 2012) to learn intuitive knowledge structures. This is unsurprising given that ACT-R's declarative memory module and chunk structure is an excellent match for the storage and retrieval of instances, which effectively guides people to some form of instance-based learning. In other words, the design and constraints of the architecture lead people to adopt an instance-based learning-like approach by using the architecture in the most direct and intuitive way.

4. Why a tripartite description?

An essential feature in being able to explain how a model performs decision-making is to examine not only the sources of generating expertise (e.g., the role of instance-based learning in naturalistic decision-making), but also to examine both *where* heuristics come from and *how* they are applied; and how they potentially lead to biased behavior. The implicit versus explicit argument ignores the question of where heuristics may come from – such as the structure of the task environment – something which is essential for the implicit learning of expertise. We argue that a tripartite description is sufficient to explain the source of both successes and failures in decision-making. These three levels include a description of the mechanisms and limitations of the architecture, the information structure in the task environment, and the use of heuristics and strategies to adapt performance to the dual constraints of cognition and environment.

We are not the first to argue for three levels of description. For instance, in the *fast-and-frugal* heuristics (Gigerenzer & Selten, 2002; Gigerenzer et al., 1999; Todd & Gigerenzer, 2003) framework, heuristics were determined from the interplay of basic cognitive capacities and the environment. Fast-and-frugal heuristics exploit commonalities in the structure of the task environment to not only support capacity limitations in the mind (e.g., short-term memory), but show how these limitation may be adaptive given the environment.

Similarly, the general argument of ecological rationality is that rationality is context- and situation-dependent, thus something considering irrational in one context may be fully rational in another. For instance, if the goal is to make accurate inferences, then using the recognition heuristic is ecological rational in environments where one's recognition of an object (e.g., a city name) correlates with the criterion to be inferred (e.g., city size). Ecological rationality is composed of three focuses: a focus on the mind, a

focus on the world (the study of regularities and constraints in the environment), and a focus on putting mind and world together (the study of ecological rationality). What we have done is adopt a similar framework for the study of cognitive architectures, and focus on explaining behavior in terms of the more transparent cognitive constraints of the architecture (in that you can point to the underlying mechanism in the architecture), task environment, and model structure. However, we argue that by situating our tripartite level of description within a cognitive architecture, we may make more causal claims as to the underlying structure of the heuristics and how they may result in biased behavior based on the complexity of the environment.

The first level of description entails an understanding of the constraints imposed by the mechanisms and limitations of the cognitive architecture. In ACT-R, these include an understanding of the impact of recency and frequency of the likelihood of an instance being retrieved, which also influences the ability of the model to generalize to new situations when using blended retrievals to generate a derived output rather than a specific instance. Other sources of constraint include the serial nature of the production system, only a single chunk being in a buffer at a time, and matching human time-course of responses. A common source of biased behavior in instance-based learning decision-making models is the use of blended retrievals, which have a tendency to retrieve values that are pulled toward the mean of all values in memory. This common mechanism can lead to both anchoring and confirmation biases based on how far the anchored value varies from the mean across all instances in memory (Lebiere et al., 2013). It is important to note that this wholly implicit process is not consciously available to the model.

The second level of description entails an understanding of the constraints imposed by the task environment. This kind of description has been somewhat neglected in discussions of the validity of cognitive models; however, it is a critical feature in understanding both the consistency of learning and the nature of biases. An understanding of the statistical and quantifiable regularities within the task environment drives the overall ability and rate of learning, and the nature of environmental feedback provides further evidence. Using an example from Simon (1990); if you want to study the movement of an ant across the beach you need look no further than the hills and valleys in the sand to determine its path. To further push this issue, Simon (1990) argued that “[h]uman rational behavior ... is shaped by a scissors whose two blades are the structure of the task environments and the computational capabilities of the actor” (Simon, 1990, p. 7). Some of the mechanisms of ACT-R were created following a rational analysis (Anderson, 1990), which assumes that cognitive processes are optimally adapted to the environment. Therefore, if we are able to capture the essential structure of the environment in ACT-R, we should be able to predict what kind of heuristics may be available or used in a given situation. This level of description is also the level of the *unit task* (Card et al., 1983), and is generally captured in ACT-R by specific goal types that drive a set of productions that represent the cognitive skills required for solving the task.

The third level of description entails an understanding of how the joint constraints of architecture and task environment influence the kinds of heuristics and strategies available to the model. In ACT-R terms, this is the explanation of the selection and sequence of productions firing. This level is most important to describe as it entails most of the choices of the modeler in designing the model. In strategy selection, even simple heuristic structures can greatly influence the output of the model, which in turn could overly constrain decision-making while also making complex problems tractable. In other words, the detection of affordances (Gibson, 1977) provided by the task environment influence the kinds of

information that the model can accumulate and the actions that the model may perform.

One use of affordances in ACT-R is to think of them in terms of cognitive niches (Marewski & Link, 2014; Marewski & Schooler, 2011). As previously discussed, cognitive niches are affordances that constrain the set of available strategies that the model may choose from based on the interplay between prior experiences and the task environment. The selection of which heuristics to apply to a given context may also be learned by the model through production utilities via reinforcement learning, with early learning occurring as a trial-and-error process until sufficient reinforcement occurs through experience or is inferred by explicit task instruction. This is analogous to how, in instance-based learning, the model transitions from reasoning via heuristics to reasoning via instances with enough experience. The mechanism for how the model moves from heuristic- to instance-based reasoning can be seen as a kind of metacognitive awareness (which itself does not totally escape the strategy selection argument).

Now that we are armed with a theory (instance-based learning) and a means of describing model output (the tripartite description), we can delve into an example.

5. Intuitive decisions in sensemaking

Rather than provide an overview of many examples, we would like to focus on an in-depth analysis of a single ACT-R model of sensemaking that uses instance-based learning to perform six complex geospatial intelligence tasks and provides both an explanation of the origin of biases and a close fit to human data (see Lebiere et al., 2013 for a more complete description of the tasks and for quantitative model fits). Sensemaking is a concept that has been used to define a class of activities and tasks in which there is an active seeking and processing of information to achieve understanding about some state of affairs in the world, which has also been applied in organizational decision-making (Weick, 1995).

Our sensemaking model is composed of three related components. The first (Tasks 1–3) learns statistical patterns of events and then generates probability distributions of category membership based on the spatial location and frequency of these events (e.g. how likely does a given event belong to each of the categories). The second (Tasks 4–6) applies probabilistic decision rules in order to generate and revise probability distributions of category membership (e.g., if a given feature is present at an event, then that event is twice as likely to belong to category A). The third (Tasks 1–6) involves making decisions about the allocation of resources based on the judged probabilities of the causes of perceived events, and is effectively a metacognitive measure of confidence in one's judgment.

The remainder of this section describes the methods and results of the six tasks from Lebiere et al. (2013). We use the term *participants* to reflect both human subjects and the ACT-R model performing the task.

For Tasks 1–3, the flow of an average trial proceeded according to the following general outline (see Fig. 4 for an example of the task interface). First, participants perceived a series of events labeled according to which category the event belonged. After perceiving the series of events, participants were asked to generate a center of activity (e.g., prototype) for each category's events, reflect on how strongly they believed the probe belonged to each category, and generate a probability estimate for each category (summed to 100% across all groups). Scoring was determined by comparing participants' distributions to an optimal Bayesian solution. Using these scores it was possible to determine certain biases. For instance, participants' probability estimates that exhibited lower entropy than a fully rational Bayes model would be considered to exhibit a

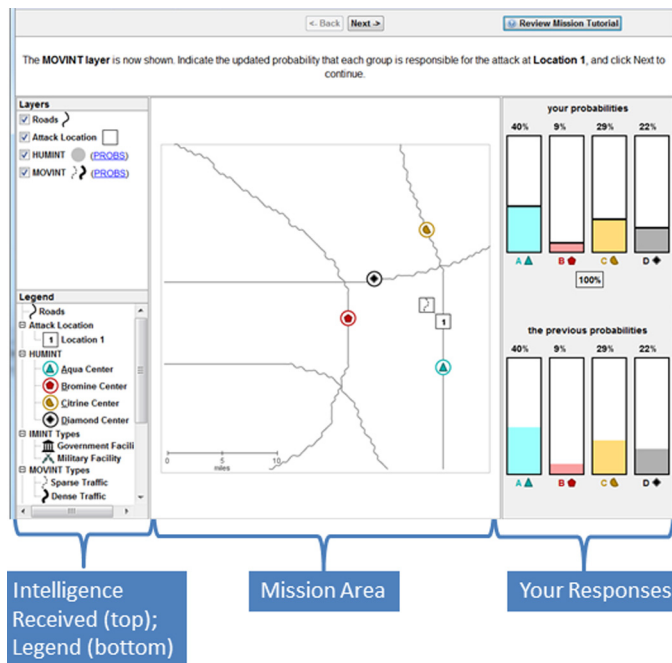


Fig. 4. A sample of the task interface. To the left is a legend explaining all the symbols on the map (center). To the right are the probability distributions for the four event categories. The pane across the top provides step-by-step instructions for participants.

confirmation bias, while probability estimates having higher entropy than an optimal Bayes model would be considered to exhibit an anchoring bias. The model was compared trial-by-trial against humans to determine whether both expressed the biases or otherwise.

Participants were then asked to allocate resources to each category with the goal of maximizing their resource allocation score, which was the amount of resources allocated to the correct category. For Tasks 1–3, the resource allocation response was a forced choice decision to allocate 100% of their resources to a single category, and participants received feedback whether or not their categorization was correct.

For Tasks 4–6, the flow of an average trial was structurally different as intelligence ‘features’, governed by probabilistic decision rules, were presented in sequential layers on the display. These Tasks required reasoning based on rules concerning the relation of observed evidence to the likelihood of an unknown event belonging to each of four different categories. Participants updated their beliefs (i.e., likelihoods) after each layer of information (i.e., feature) was presented. For instance, in Task 4, after determining the center of activity for each category (similar in mechanism to Tasks 1–3) and reporting an initial probability estimate, the SOCINT (SOCIAL INTeLLIGENCE) layer would be presented by displaying color-coded regions on the display representing each category’s boundary, where the likelihood of the event belonging to a given category is twice as likely if it was within that category’s boundary. After reviewing the information presented by the SOCINT layer, participants were required to update their likelihoods based on this information and the corresponding probabilistic decision rule.

When all the layers have been presented (two layers in Task 4, five layers in Task 5, and four layers in Task 6), participants were required to generate a resource allocation. In these Tasks, the resource allocation response was produced using the same interface as probability estimates. For instance, assuming that resources were allocated such that $\{A=40\%, B=30\%, C=20\%, D=10\%\}$, if the probe belonged to category A (i.e., that A was the ‘ground truth’)

then the participant would receive a score of 40 out of 100, whereas if the probe instead belonged to category B, they would score 30 points.

Two separate exams were collected from two separate populations. The model performed the first exam and was compared against 45 participants who were employees of the MITRE Corporation. All participants completed informed consent and debriefing questionnaires that satisfied IRB requirements. Without going into extensive detail over the results, the ACT-R model significantly predicted many of the trial-by-trial variations in human performance, and not only the presence or absence of a bias, but also the quantity of the bias metric, reflected in an overall $r^2 = .645$ for negentropy scores across all tasks.

The results of the model were then compared to the results of a novel sample gathered from 103 students at Penn State University. This new data set was not available before the model was run, and no parameters or knowledge structures were changed to fit this data set. Unlike the original 45-participant dataset, the Penn State sample used only people who had taken course credit toward a graduate Geospatial Intelligence Certificate. The model correctly predicted both the presence and degree of biases on every trial in Tasks 1–3, and followed similar trial-by-trial trends for biases in Tasks 4–5. The quantitative fits of the model were also similar, with an overall $r^2 = .591$ for negentropy scores across all tasks.

6. We will now provide an overview of the model function

6.1. Biases in group center generation

In the first task component, the flow of an average trial began with participants perceiving a series of events labeled according to which category the event belonged, each corresponding to a group icon on the central map, after which a probe was displayed. Participants were then required to generate a center of activity for each category’s events, and generate a probability estimate for each category (summed to 100%).

When group centers were generated directly from a retrieval of events represented in memory, the blended retrieval process in ACT-R reflected a disproportionate influence of the most recent events given their higher base-level activation. A strategy to combat this recency bias consisted of generating a final response by performing a blended retrieval over all the group centers (both current and past centers generated for previous trials) stored in memory, thereby giving more weight to earlier events by compounding the influence of earlier centers over the subsequent blended retrievals. This second-order blended retrieval is done for each category across their prior existing centers, which we refer to as the generation of a *centroid-of-centroids*. This effectively implements an anchoring-and-adjustment process where each new estimate is a combination of the previous ones together with the new evidence.

A fundamental difference with traditional implementation of anchoring-and-adjustment heuristics is that this process is entirely constrained by the architectural mechanisms (especially blending) and does not involve any additional degrees of freedom. Moreover, because there are an equal number of centroid-of-centroids chunks (one per category created after each trial), there is no effect of base-rate on the model’s later probability judgments, even though the base-rate for each category is implicitly available in the model based on the number of recallable events. This illustrates the metacognitive nature of heuristics in our tripartite organization: given that the nature of cognitive mechanisms gives rise to a recency bias that is incompatible with the task environment (assuming a stable distribution), the centroid-of-centroids heuristic is used to give more weight to older instances and circumvent the recency bias. Note that the bias toward recency in architectural

mechanisms arose because it indeed reflected the nature of many environments (Anderson & Schooler, 1991), making it well adapted to those settings. There is no such thing as suboptimal bias: just a mismatch between assumptions and environment that occasionally needs to be supplemented with the proper heuristic adjustment.

6.2. Biases in probability adjustment

In this task component, event features – such as the location or context of events – were presented in sequential layers on the display. Initial distributions for each category were provided to participants, after which participants updated their beliefs after each feature was revealed. Beliefs were updated based on a set of provided probabilistic decision rules: e.g., if the MOVINT (*movement intelligence*) feature shows dense traffic, then groups A and C are four times as likely as groups B and D. When all the layers were presented, participants were required to allocate resources to each category.

To leverage an instance-based learning approach for probability adjustment, the ACT-R model's memory was seeded with a range of instances consisting of triplets: an initial probability, an adjustment factor, and the resulting probability. The factor is set by the explicit rules of the task. When the model is asked to estimate the resulting probability for a given prior and multiplying factor, it simply performs a blended retrieval specifying prior and factor, and then outputs the posterior probability that represents the blended consensus of the seeded chunks.

When provided with linear similarities between probabilities (and factors), the primary effect is an underestimation of the adjusted probability for much of the initial probability range (i.e., an anchoring bias), with an overestimation on the lower end of the range (i.e., confirmation bias). While the magnitude of the biases can be modulated somewhat by architectural parameters, the effects themselves are a priori predictions of the architecture, in particular its theoretical constraints on memory retrieval.

A simpler and more *implicit* model of probability adjustment can be produced by representing the various hypotheses as chunks in memory and using their activation as an estimate of their strength of support. When evidence is received, it is matched against patterns linking it to various hypotheses and the best matching one is retrieved, leading to a boost in activation. If contradictory evidence starts accumulating, two biases will emerge. First, new evidence will sometimes be misinterpreted because the current dominant hypothesis is most active and can overcome some degree of mismatch. Second, even if the evidence is correctly interpreted and the correct hypothesis reinforced, for the new hypothesis to attain primacy it will take some time to sufficiently build activation and for the activation of the previously dominant hypothesis to sufficiently decay over time. This process has been given a number of names, from anchoring bias to persistence of discredited evidence.

A number of structured analytic techniques have been proposed to remedy these biases emerging from the dynamics of our cognitive system (Heuer & Pherson, 2010). The most prominent one might be Analysis of Competing Hypotheses, which proposes a process by which all competing hypotheses are evaluated against each piece of evidence and the sum of their support only computed and compared at the end. This is done to prevent the early emergence of a favored hypothesis and the resulting biases. An analog to the Analysis of Competing Hypotheses has been implemented in our model and can be shown to directly affect the activation dynamics described above. Each hypothesis chunk receives a rehearsal at each step, equalizing the influence of base-rate from their activation and preventing a winner-take-all dynamic. The result is that their activation over time will simply reflect the degree of support that they have received. In this example, structured analytic

techniques can also be seen as metacognitive heuristics that leverage the beneficial aspects of cognitive mechanisms while defeating or at least limiting their potential biases and thus provide external aids to our intuitive decision-making.

6.3. Biases in resource allocation

Resource allocation makes use of the same instance-based learning paradigm as probability adjustment. This unified mechanism has no explicit strategies, but instead learns to allocate resources according to the outcomes of past decisions. The model generates a resource allocation distribution by focusing on the leading category and determining how many resources to allocate to that category. The remaining resources are divided amongst the remaining three categories in proportion to their assigned probabilities. Representation of a trial instance consists of three parts: a decision context (i.e., the probability of the leading category), the decision itself (i.e., the resource allocation to the leading category), and the outcome of the decision (i.e., the payoff).

The model's control logic takes a hybrid approach between choice (Lebiere & Anderson, 2011) and decision models (Wallach & Lebiere, 2003a), involving two steps of access to experiences in declarative memory rather than a single one. When determining how many resources to apply to the lead category, the model initially has only the probability assigned to that category. The first step is done by performing a blended retrieval on chunks representing past resource allocation decisions using the probability as a cue. The outcome value of the retrieved chunk is the expected outcome for the trial. The second step is to generate the decision that most likely leads to that outcome given the context. Note that this process is not guaranteed to generate optimal decisions, and indeed people do not. Rather, it represents a parsimonious way to leverage our memory of past decisions in this paradigm that still provides functional behavior. A significant theoretical achievement of our approach is that it unifies control models and choice models in a single decision-making paradigm.

After feedback is received, the model learns a resource allocation decision chunk that associates the leading category probability, the quantity of resources assigned to the leading category, and the actual outcome of the trial (i.e., the resource allocation score for that trial). Additionally, up to two counterfactual chunks are committed to declarative memory. The counterfactuals represent what would have happened if a winner-take-all resource assignment had been applied, and what would have happened if a pure probability-matched resource assignment (i.e., using the same values as the final probabilities) had been applied. The actual nature of the counterfactual assignments is not important; what is essential is to give the model a broad enough set of experience representing not only the choices made but also those that could have been made. The use of a counterfactual strategy to generate a diversity of outcomes, experienced or imagined, can be seen as a very general and effective metacognitive heuristic.

The advantage of this approach is that the model is not forced to choose between a discrete set of strategies such as winner-take-all or probability matching; rather, various strategies can emerge from instance-based learning. By priming the model with the winner-take-all and probability matching strategies (essentially the boundary conditions), it is possible for the model to learn any strategy in between them, such as a tendency to more heavily weigh the leading candidate, or even suboptimal strategies such as choosing 25% for each of the four categories (assuring a score of 25 on the trial) if the model receives enough negative feedback (i.e., poor scores) so as to encourage risk aversion. Instance-based learning can thus be seen in this instance as a highly flexible metacognitive strategy from which a number of more limited, hardwired strategies can emerge.

7. Discussion

So far, we have argued that cognitive architectures aid in the study of intuitive decision-making by providing a falsifiable theory for the study of mechanisms, processes, and representations involved in decision-making. By using a cognitive architecture, one is adopting constraints involved in managing the flow of knowledge and processes involved in these knowledge operations. Architectures expand our ability to go beyond ‘just-so’ explanations to describe the underlying processes and knowledge leading up to decisions. They also provide more flexibility beyond the constraints of expertise-based systems when operating outside of very constrained and/or very stable environments. In many cases, models based in a cognitive architecture can perform tasks and provide testable predictions that are too complex to analyze with traditional experimental methods due to the combinatorics of possible decisions.

The next step in the development of cognitive architectures should be mechanisms to support the generalizability of models and reduce degrees of freedom. Some preliminary thrusts include: the integration of neurally plausible associative learning to drive implicit statistical learning of regularities within the environment (Thomson & Lebiere, 2013), the development of expectation-driven cognition to cue episodic memory formation (Kurup et al., 2012), and more generally the development of strategy selection (or metacognitive awareness) within the architecture to guide the selection of features used in the representation and retrieval of instances (Lebiere et al., 2009; Marewski & Link, 2014; Marewski & Schooler, 2011; Reitter, Juvina, Stocco, & Lebiere, 2010; Reitter, 2010).

Ideally, the strategies and heuristics implemented in the architecture should be selected (if not created) by the model itself rather than provided by the modeler. The model (driven by the architecture) should be responsible for the selection and evolution of strategies. To get started, however, several general procedures are needed to bootstrap learning until sufficient knowledge is learned, at which point processes implicated in generating expertise should lead to interesting emergent behaviors (and novel predictions) within the model. The question of which minimal set of procedures best captures human performance is an empirical one, and one that needs to be a center of focus. The adoption of general frameworks such as instance-based learning and the adoption of a common set of heuristics across tasks appear to be the next step in the right direction.

Conflict of Interest Statement

The authors certify that they have NO affiliations with or involvement in any organization or entity with any financial interest (such as honoraria; educational grants; participation in speakers' bureaus; membership, employment, consultancies, stock ownership, or other equity interest; and expert testimony or patent-licensing arrangements), or non-financial interest (such as personal or professional relationships, affiliations, knowledge or beliefs) in the subject matter or materials discussed in this manuscript.

References

- Anderson, J. R. (1982). Acquisition of cognitive skill. *Psychological Review*, 89, 369–406.
- Anderson, J. (1990). *The Adaptive Character of Thought*. Hillsdale, NJ: Erlbaum Associates.
- Anderson, J. R. (2007). *How can the human mind occur in the physical universe?* Oxford, UK: Oxford University Press.
- Anderson, J. R., Bothell, D., Byrne, M. D., Douglass, S., Lebiere, C., & Qin, Y. (2004). An integrated theory of mind. *Psychological Review*, 11(4), 1036–1060.
- Anderson, J. R., & Lebiere, C. J. (Eds.). (1998). *The atomic components of thought*. Mahwah, NJ: Psychology Press.
- Anderson, J. R., & Schooler, L. J. (1991). Reflections of the environment in memory. *Psychological Science*, 2, 396–408.
- Berger, J. O. (1985). *Statistical decision theory and Bayesian analysis*. New York, NY: Springer.
- Bowers, J. S., & Davis, C. J. (2012). Bayesian just-so stories in psychology and neuroscience. *Psychological Bulletin*, 138(3), 389–414.
- Brunswick, E. (1957). Scope and aspects of the cognitive problem. In H. Gruber, K. R. Hammond, & R. Jessor (Eds.), *Contemporary approaches to cognition* (pp. 5–31). Cambridge: Harvard University Press.
- Bussemeyer, J. R., Pothos, E. M., Franco, R., & Trueblood, J. S. (2011). A quantum theoretical explanation for probability judgment errors. *Psychological Review*, 118(2), 193.
- Card, S. K., Moran, T. P., & Newell, A. (1983). *The psychology of human-computer interaction*. Hillsdale, NJ: Lawrence Erlbaum Associates.
- Chase, W. G., & Simon, H. A. (1973). Perception in chess. *Cognitive Psychology*, 4(1), 55–81.
- Cooper, R. P. (2007). The role of falsification in the development of cognitive architectures: Insights from a Lakatosian analysis. *Cognitive Science*, 31(3), 509–533.
- Dimov, C. M., Marewski, J. N., & Schooler, L. J. (2013). Constraining ACT-R models of decision strategies: An experimental paradigm. In M. Knauff, M. Pauen, N. Sebanz, & I. Wachsmuth (Eds.), *Proceedings of the 35th annual conference of the Cognitive Science Society* (pp. 2201–2206). Austin, TX: Cognitive Science Society.
- Dougherty, M. R., Franco-Watkins, A. M., & Thomas, R. (2008). Psychological plausibility of the theory of probabilistic mental models and the fast and frugal heuristics. *Psychological Review*, 115(1), 199–211.
- Erickson, M. A., & Kruschke, J. K. (1998). Rules and exemplars in category learning. *Journal of Experimental Psychology: General*, 127, 107–140.
- Heuer, R. J., & Pherson, R. H. (2010). *Structured analytic techniques for intelligence analysis*. Thousand Oaks, CA: CQ Press.
- Gagliardi, F. (2011). Instance-based classifiers applied to medical databases: Diagnosis and knowledge extraction. *Artificial Intelligence in Medicine*, 52(3), 123–139.
- Gibson, J. J. (1977). The Theory of Affordances. In R. Shaw, & J. Bransford (Eds.), *Perceiving, Acting, and Knowing. Toward an Ecological Psychology* (pp. 67–82). Hillsdale, NJ: Lawrence Erlbaum Associates.
- Gigerenzer, G., Hoffrage, U., & Goldstein, D. G. (2008). Fast and frugal heuristics are plausible models of cognition: Reply to Dougherty, Franco-Watkins, and Thomas (2008). *Psychological Review*, 115(1), 230–239.
- Gonzalez, C., & Lebiere, C. (2005). Instance-based cognitive models of decision making. In D. Zizzo, & A. Courakis (Eds.), *Transfer of knowledge in economic decision making*. New York: New York: Palgrave MacMillan.
- Gonzalez, C., & Lebiere, C. (2013). Cognitive architectures combine formal and heuristic approaches. *Behavioral and Brain Sciences*, 36(3), 285–286.
- Gonzalez, C., Lerch, J. F., & Lebiere, C. (2003). Instance-based learning in dynamic decision making. *Cognitive Science*, 27, 591–635.
- Gigerenzer, G., & Selten, R. (Eds.). (2002). *Bounded rationality: The adaptive toolbox*. Cambridge, MA: The MIT Press.
- Gigerenzer, G., Todd, P. M., & the ABC Research Group. (1999). *Simple heuristics that make us smart*. New York, NY: Oxford University Press.
- Griffiths, T. L., Chater, N., Norris, D., & Pouget, A. (2012). How the Bayesians got their beliefs (and what those beliefs actually are): Comment on Bowers and Davis (2012). *Psychological Bulletin*, 138(3), 415–422.
- Hogarth, R. M. (2001). *Educating intuition*. Chicago, IL: University of Chicago Press.
- Kahneman, D., & Frederick, S. (2005). A model of heuristic judgment. In K. J. Holyoak, & R. G. Morrison (Eds.), *The Cambridge handbook of thinking and reasoning* (pp. 267–293). New York: Cambridge University Press.
- Kahneman, D., & Klein, G. (2009). Conditions for intuitive expertise: A failure to disagree. *American Psychologist*, 64(6), 515.
- Kahneman, D., & Tversky, A. (1979). Prospect theory: An analysis of decision under risk. *Econometrica: Journal of the Econometric Society*, 47, 263–291.
- Kahneman, D., & Tversky, A. (1996). On the reality of cognitive illusions. *Psychological Review*, 103(3), 582–591.
- Kennedy, W., & Patterson, R. (2012). Modeling intuitive decision making in ACT-R. In *Proceedings of the eleventh international conference on computational modeling* (pp. 1–6).
- Klein, G. A. (1998). *Sources of power: How people make decisions*. Cambridge, MA: MIT Press.
- Klein, G. (2009). *Streetlights and shadows: Searching for the keys to adaptive decision making*. Cambridge, MA: The MIT Press.
- Klein, G. A., Calderwood, R., & Clinton-Cirocco, A. (1986). Rapid decision making on the fireground. In *Proceedings of the human factors and ergonomics society 30th annual meeting*, Vol. 1 (pp. 576–580).
- Kruglanski, A. W., & Gigerenzer, G. (2011). Intuitive and deliberate judgments are based on common principles. *Psychological Review*, 118(1), 97–109.
- Kurup, U., Lebiere, C., Stentz, A., & Hebert, M. (2012). Using expectations to drive cognitive behavior. In *Proceedings of the Twenty-Sixth AAAI Conference on Artificial Intelligence*. Ontario, Canada: AAAI Press.
- Lebiere, C., & Anderson, J. R. (2011). Cognitive constraints on decision making under uncertainty. *Frontiers in Cognition: Psychology*, 2, 305–309.
- Lebiere, C., & Best, B. J. (2009). From Microcognition to Macro cognition: Architectural support for asymmetric adversarial behavior. *Journal of Cognitive Engineering and Decision Making*, 3(2), 176–193.
- Lebiere, C., Gonzalez, C., & Warwick, W. (2009). Convergence and constraints revealed in a qualitative model comparison. *Journal of Cognitive Engineering and Decision Making*, 3(2), 131–155.

- Lebiere, C., Gonzalez, C., & Martin, M. (2007). Instance-based decision making model of repeated binary choice. In *Proceedings of the 8th international conference on cognitive modeling* Ann Arbor, MI, USA.
- Lebiere, C., Gray, R., Salvucci, D., & West, R. (2003). Choice and learning under uncertainty: A case study in baseball batting. In *Proceedings of the 25th annual meeting of the Cognitive Science Society* (pp. 704–709).
- Lebiere, C., Pirolli, P., Thomson, R., Paik, J., Rutledge-Taylor, M., Staszewski, J., et al. (2013). A functional model of sensemaking in a neurocognitive architecture. *Computational Intelligence and Neuroscience*, <http://dx.doi.org/10.1155/2013/921695>. Special Issue on Neurocognitive models of Sense-Making. We are article ID 921695
- Lebiere, C., & Wallach, D. (2001). Sequence learning in the ACT-R cognitive architecture: Empirical analysis of a hybrid model. In *Sequence learning*. Berlin/Heidelberg: Springer.
- Lebiere, C., Wallach, D., & Taatgen, N. (1998). Implicit and explicit learning in ACT-R. In *Second European conference on cognitive modelling*.
- Lebiere, C. (1999). The dynamics of cognitive arithmetic. *Kognitionswissenschaft [Journal of the German Cognitive Science Society]* Special issue on cognitive modelling and cognitive architectures, D. Wallach & H. A. Simon (Eds.), 8 (1), 5–19.
- Lewandowsky, S. (1993). The rewards and hazards of computer simulations. *Psychological Science*, 4, 236–243.
- Marewski, J. N., Gaissmaier, W., Schooler, L. K., Goldstein, D., & Gigerenzer, G. (2010). From recognition to decisions: Extending and testing recognition-based models for multialternative inference. *Psychonomic Bulletin & Review*, 3, 287–309.
- Marewski, J. N., & Link, D. (2014). Strategy selection: An introduction to the modeling challenge. *WIREs Cognitive Science*, 5, 39–59.
- Marewski, J. N., & Mehlhorn, K. (2011). Using the ACT-R architecture to specify 39 quantitative process models of decision-making. *Judgment and Decision Making*, 6(6), 439–519.
- Marewski, J. N., & Schooler, L. J. (2011). Cognitive niches: An ecological model of strategy selection. *Psychological Review*, 118(3), 393.
- Marr, D. (1982). *Vision: A computational approach*. San Francisco, CA: Freeman & Company.
- Newell, A. (1990). *Unified theories of cognition*. Cambridge, MA: Harvard University Press.
- Newell, A., & Simon, H. A. (1972). *Human problem solving* (Vol. 14) Englewood Cliffs, NJ: Prentice-Hall.
- Reitter, D. (2010). Metacognition and multiple strategies in a cognitive model of online control. *Journal of Artificial General Intelligence*, 2(2), 20–37.
- Reitter, D., Juvina, I., Stocco, A., & Lebiere, C. (2010). Resistance is futile: Winning lemonade market share through metacognitive reasoning in a three-agent cooperative game. In *Proceedings of the behavior representation in modeling and simulations (BRIMS 2010) conference* Charleston, SC.
- Roberts, S., & Pashler, H. (2000). How persuasive is a good fit? A comment on theory testing. *Psychological Review*, 107, 358–367.
- Schultheis, H. (2009). Computation and explanatory power of cognitive architectures: The case of ACT-R. In *Proceedings of the 9th international conference of cognitive modeling*.
- Shallice, T., & Cooper, R. (2011). *The organisation of mind*. Oxford UK: Oxford University Press.
- Simon, H. A. (1990). *Reason in human affairs*. Stanford CA: Stanford University Press.
- Simon, H. A. (1992). What is an “explanation” of behavior? *Psychological Science*, 3(3), 150–161.
- Simon, H. A. (1995). Rationality in political behavior. *Political Psychology*, 16, 45–61.
- Stanovich, K. E., & West, R. F. (1999). Discrepancies between normative and descriptive models of decision making and the understanding/acceptance principle. *Cognitive Psychology*, 38, 349–385.
- Stewart, T. C., West, R., & Lebiere, C. (2009). Applying cognitive architectures to decision-making: How cognitive theory and the equivalence measure triumphed in the Technion Prediction Tournament. In N. A. Taatgen, & H. van Rijn (Eds.), *Proceedings of the 31st annual conference of the Cognitive Science Society*. Austin, TX: Cognitive Science Society.
- Taatgen, N. A., & Anderson, J. R. (2008). Constraints in cognitive architectures. In *Cambridge handbook of computational psychology*. Cambridge, NY: Cambridge University Press.
- Taatgen, N., & Anderson, J. R. (2010). The past, present, and future of cognitive architectures. *Topics in Cognitive Science*, 2(4), 693–704.
- Taatgen, N. A., Lebiere, C., & Anderson, J. R. (2006). Modeling paradigms in ACT-R. In R. Sun (Ed.), *Cognition and multi-agent interaction: From cognitive modeling to social simulation* (pp. 29–52). Cambridge: University Press.
- Thomson, R., & Lebiere, C. (2013). Constraining Bayesian inference with cognitive architectures: An updated associative learning mechanism in ACT-R. In M. Knauff, M. Pauen, N. Sebanz, & I. Wachsmuth (Eds.), *Proceedings of the 35th annual conference of the Cognitive Science Society*. Austin, TX: Cognitive Science Society.
- Todd, P. M., & Gigerenzer, G. (2003). Bounding rationality to the world. *Journal of Economic Psychology*, 24, 143–165.
- Wallach, D., & Lebiere, C. (2003a). Conscious and unconscious knowledge: Mapping to the symbolic and subsymbolic levels of a hybrid architecture. In L. Jimenez (Ed.), *Attention and implicit learning*. Amsterdam, Netherlands: John Benjamin Publishing Company.
- Wallach, D., & Lebiere, C. (2003b). Implicit and explicit learning in a unified architecture of cognition. *Advances in Consciousness Research*, 48, 215–252.
- Weick, K. (1995). *Sensemaking in organizations*. Thousand Oaks, CA: Sage.
- West, R. L., & Lebiere, C. (2001). Simple games as dynamic, coupled systems: Randomness and other emergent properties. *Journal of Cognitive Systems Research*, 1(4), 221–239.