

Metareasoning: Thinking about Thinking

Michael T. Cox (ed.), Anita Raja (ed.)

<https://doi.org/10.7551/mitpress/9780262014809.001.0001>

Published online: 22 August 2013 **Published in print:** 18 March 2011

9780262295284

Print ISBN: 9780262014809

Online ISBN:

Search in this book

CHAPTER

4 Learning Expertise with Bounded Rationality and Self-Awareness

Susan L. Epstein, Smiljana Petrovic

<https://doi.org/10.7551/mitpress/9780262014809.003.0004>

Published: March 2011

Abstract

This chapter describes how an architecture enhanced by metareasoning assesses expertise, manages large bodies of heuristics, and learns to think more effectively. FORR (FOr the Right Reasons) is a learning and problem-solving architecture that models the development of expertise with bounded rationality and self-awareness. Metareasoning is essential in FORR's ability to learn to solve problems within a given class. As it learns to manage heuristics, FORR uses metareasoning to decide when to abandon an unpromising learning attempt, when to stop learning, how to select heuristics during learning, and how to prioritize heuristics. FORR also reasons about its performance on previous problems, its previous decisions, and the relative discriminatory power of its heuristics. Thinking about thinking makes FORR-based applications more incisive and more successful.

Keywords: [metareasoning](#), [FORR](#), [expertise](#), [bounded rationality](#), [self-awareness](#), [problem solving architecture](#)

Subject: [Artificial Intelligence](#)

When people provide only rudimentary knowledge about a problem domain, a *self-adaptive* system must learn its own expertise, that is, become better and faster at its task than the rest of us (D'Andrade, 1991). Essential to that development are *bounded rationality* (imposed resource limitations) and *self-awareness*, the ability to monitor one's own problem-solving behavior and reasoning. From this perspective, *metareasoning* examines and reflects upon the components of decision making and how they are controlled, rather than merely upon the decisions themselves. *Learning with metareasoning* occurs when the system seeks to improve its performance by the application of metareasoning to modify its behavior. A system with this capacity could, for example, decide to ignore some of its experience, to prefer some procedures to others, or even to stop learning because it is satisfied with what it has accomplished. Under the aegis of an architecture that supports learning

obstruction that lies directly between the robot and the goal. At most one plan is active at a time. The tier-1 Advisor *Enforcer* appropriately supports the execution of an active plan or terminates it.

- *Most good reasons are fallible heuristics.* These tier-3 Advisors express preferences for choices numerically, as *strengths*. For example, Ariadne's *Giant Step* assigns greater strengths to longer steps. The object level in figure 4.1 coordinates all three tiers of Advisors to make a decision. It progresses to the next tier only if the current one does not produce a decision.
- *Good reasons make different contributions on different problem classes.* FORR evaluates each tier-3 Advisor on each new problem class. The more constructive are *class-appropriate*; the others are *class-inappropriate*. ACE's heuristics, for example, are provided in *dual pairs* (one to maximize and the other to minimize the same metric); ACE sorts out whether either is appropriate. Lines 1 and 2 in table 4.1 show that duals may perform differently on different problem classes.
- *A combination of good reasons offers substantial benefits.* There is evidence for this both in people's reliance on multiple heuristics (Biswas et al., 1995; Crowley & Siegler, 1993; Ratterman & Epstein, 1995; Schraagen, 1993) and in programs that integrate multiple rationales to their advantage (Keim et al., 1999). Table 4.1 (lines 3–5) shows how a pair of heuristics, one for variable selection and the other for value selection, may outperform individual heuristics (lines 1–2).
- *Good combinations of reasons vary with the problem class.* Table 4.1 (lines 3 and 5) shows that pairs of Advisors successful on one problem class may do less well on another. ↴
- *A program can learn to become an expert.* In the *learning phase*, FORR addresses a sequence of problems, and all of figure 4.1 is active. Metareasoning examines the data from the object level after each learning-phase problem, and may reformulate its strategy. Then, in the *testing phase*, FORR turns learning off and evaluates its performance on a second set of problems. No further directives from the metareasoning module occur during testing. Because decisions may be nondeterministic, FORR evaluates performance over a set of such runs (an *experiment*). Relatively few learning problems should suffice. For example, ACE learns to solve each problem class well after learning on only 30 problems. Table 4.1 (line 6) shows how the mixture of heuristics that ACE learned for each of them outperformed both individual heuristics and pairs of them.
- *Multiple representations enhance reasoning.* Each Advisor has access to every representation of events and states from the ground level and the object level. These representations are shared and computed at the object level only on demand. For example, ACE has dozens of CSP descriptions, including many that identify different kinds of subproblems and relationships among variables.

Bounded rationality sets an arbitrary performance standard. One might, for example, learn only from solutions achieved within bounded resources. This presumes that a fast solution or one that visits fewer search states is better. Setting such bounds is problematic, however. Under bounds too low for typical problems in the class, only the easiest are solved, solutions are fewer, and the learned behavior is untrustworthy because it arose from problems where most any decision would do. If bounds are set too high, however, long but successful searches provide traces of not-so-expert behavior. For example, figure 4.2 shows the impact on ACE of a limit on *search nodes* (partial assignments) while it learns to solve the geometric problems. A run was *successful* if it solved at least 80 percent of its 50 testing problems within the node limit. Observe how higher node limits during learning produced fewer successful runs and incurred a higher search cost than did lower limits.

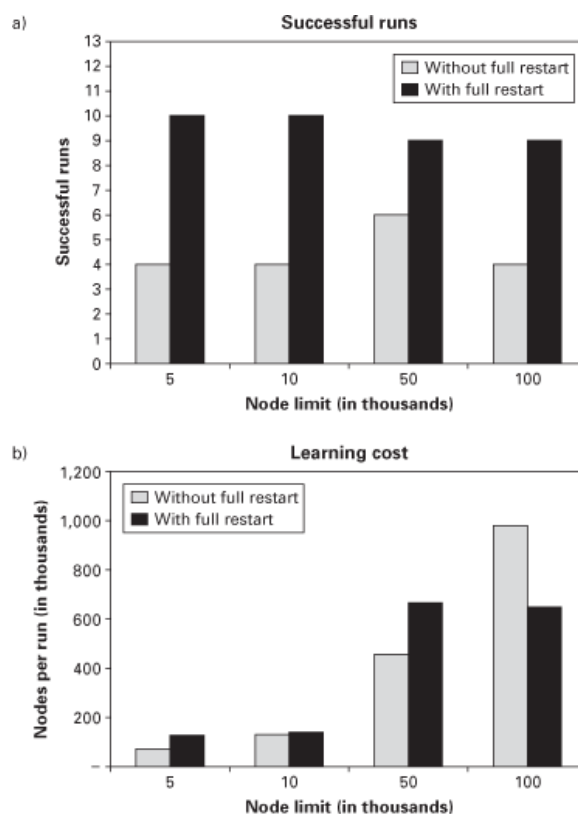


Figure 4.2 The impact of resource bounds on 10 runs for the geometric problems. (a) More runs in an experiment are successful with full restart (darker bars) and (b) fewer nodes are expanded.

Metareasoning can monitor overall skill and use it to control the learner. A coarse but significant measure of success is the frequency of solved problems during learning. FORR introduces *full restart*, the ability to reinitialize all heuristics' weights to some small value and begin again, on different problems from the same class (Petrovic & Epstein, 2006a). (This is different from restart on a problem, which tries to diversify search for its solution [Gomes & Sellman, 2004]). When problems frequently go unsolved during a learning phase, full restart begins the entire phase again. This abandons a run on an unusually difficult sample of problems, or one where solutions to very easy problems introduce misinformation. The resource limit is crucial here, as figure 4.2 indicates. As one would expect, higher resource bounds incur a higher learning cost, but under a high node limit, full restart considerably modulates that effect.

Modeling Expertise

Modeling expertise requires some standard: an oracle, perhaps, or an expert opponent. Such guidance, although important, may not offer enough variety to develop a robust

learner. Experiments with Hoyle in table 4.2, for example, found competition against a *perfect player* (a program that makes only optimal moves) too narrow (Epstein, 1994b). An expert player should succeed against opponents of any strength. After training against a perfect player, Hoyle lost testing contests to opponents with far less prowess—it was repeatedly flummoxed by their errors. Although some percentage of random decisions in an otherwise flawless opponent drives the learner’s experience outside perfect play, such moves lack good rationales and are therefore often of poor quality—not the kind of decisions Hoyle should learn to make. If Hoyle learned only against itself, it did not always develop sufficiently strong expertise either. The most effective approach, *lesson and practice training*, had Hoyle alternately play two contests against a perfect player and then practice in seven contests against itself.

Table 4.2 Skill after learning against different playing expertise (Epstein, 1995). For competition against opponents of different strengths at this draw game, lesson and practice training was better preparation than training against a perfect player. Hoyle played 100 testing contests against each of 4 opponents: one moved perfectly, the others had some percentage of random moves among otherwise perfect play: 10 percent random (expert), 70 percent random (novice), and 100 percent (random).

Outcomes	Percentage of contests played correctly	
	After learning against a perfect player	After lesson and practice training
Wins against an expert	12%	18%
Wins against a novice	59%	63%
Wins against random play	80%	85%
Wins or draws against perfect play	100%	100%
Wins or draws against an expert	93%	98%
Wins or draws against a novice	80%	97%
Wins or draws against random play	88%	100%

Without an external standard of expertise, a system can take traces of its own successes as a model. Neither Ariadne nor ACE has an external model; each learns alone and can only judge the correctness of its actions from its ability to solve problems. When the robot finds the goal, metareasoning excises any closed loops from a trace of the robot’s path, and takes the remainder as a model. Similarly, when ACE solves a problem, metareasoning excises any decisions that have no alternative or are subsequently retracted and takes the remaining ones as a model. Nonetheless, neither a loop-free path nor a retraction-free search is likely to be ideal—there may have been a better way to solve the problem.

RSWL judges correctness with relative support (Petrovic & Epstein, 2006b). Under RSWL, reinforcements are directly proportional to relative support, but penalties are also inversely proportional to the number of choices. Two variations on RSWL modify rewards and penalties based on different estimates of training instance difficulty. RSWL- κ uses *constrainedness* (a measure of a CSP class difficulty (Gent, Prosser, & Walsh, 1996)) to estimate the difficulty of each search state dynamically. RSWL-d is computationally less expensive; it uses search tree depth and assumes that decisions are more difficult at the top of the search tree. RSWL-d is based on the premise that, with respect to a given algorithm, every CSP has a *backdoor*, a set of variables after whose consistent assignment with the constraints search becomes extremely easy (Williams, Gomes, & Selman, 2003).

Learning on a sequence of problems applies knowledge from one successful search to subsequent searches. The first success increases the weights of those Advisors that contributed to its decisions. It may, however, be quite expensive to solve any first problem at all, particularly when the problems are hard, there are many Advisors, and they disagree with one another. On occasion a run fails because the learner has solved no problems at all, and therefore changed no weights. Otherwise, only the easiest problems in a class may be solved, which produces relatively few training instances, all from relatively easy situations. FORR's metareasoning therefore includes the ability to work with *random subsets* of Advisors (Petrovic & Epstein, 2008). For each problem in the learning phase, this method chooses a fresh subset of tier-3 Advisors to consult; if search solves that problem, FORR learns weights from it only for that subset. The expectation is that eventually the subset chosen for some problem will be dominated by class-appropriate heuristics, the problem will be solved, the class-appropriate weights in the subset will increase, and whenever any of those Advisors appears in a subsequent subset it will be more likely to influence search, making further successes more likely. Subset size is crucial, however: it must be large enough to select Advisors frequently, yet small enough to speed processing and to give an otherwise minority voice the opportunity to dominate decisions. Table 4.3 shows how performance improves using random subsets: It reduces the number of problems addressed during learning and speeds computation time on each decision. It even functions well when we deliberately skew the initial Advisor pool with many more class-inappropriate than class-appropriate Advisors.

p. 53

Table 4.3 Random subsets of 30 percent improve ACE's learning performance on the geometric problems with full restart under a 5,000-node limit, with no statistically significant change in testing performance. An early failure is an unsolved problem before any solved one.

	Learning	
	Without random subsets	With random subsets
Number of learning problems	44.8	36.1
Number of learning failures	13.7	6.4
Number of early (learning) failures	7.1	0.9
Number of successful runs (of 10)	10	10
Seconds per learning decision	0.0161	0.0106
Seconds per learning run	1651.6	593.6
Average number of nodes in testing	192.7	195.9
Percentage of solved testing problems	98.6%	96.4%

Finally, a self-aware system can monitor the rate at which its performance improves and stop learning once it is no longer learning anything new. FORR's metareasoning includes such *learning to stability* (Epstein, Freuder, & Wallace, 2005). Under this option, FORR monitors its Advisors' weights across a recent time window (e.g., the last 20 problems) and terminates a learning phase when the standard deviation of changes in the Advisors' weights across the window are less than ϵ . Learning to stability assumes that stable weights will remain stable; our experience over far longer learning phases confirms this. Making a learner more responsive to its own learning experience this way has proved successful in all three domains: there is no change in performance, only a reduction in the resources that would have been devoted to learning after the system found no way to improve its weights further. For example, Hoyle recognizes that it has learned all it can on simple games after 12 contests; more difficult games require as many as 120.

Learning to Reason Less

Metareasoning can monitor the traces from individual problems to reduce computation in a variety of ways, thereby restructuring the reasoning process itself. Unless errors can prove fatal, more thinking is not always better. Of course, such economy should be evaluated against the risks error presents and the cost required to recover from it.

p. 54 Ideally, decision making uses only the best Advisors and emphasizes those that offer better advice. The metalevel can instruct the object level to omit from computation any Advisors that produce no advice during learning (e.g., Material in tic-tac-toe). Moreover, once weights are learned, FORR uses *benchmark Advisors*, which produce random advice, to identify class-appropriate Advisors. A benchmark Advisor does not vote, but it does receive a learned weight. After the learning phase, metareasoning eliminates from participation any Advisor whose weight is lower than its benchmark's. Moreover, any representation referenced only by eliminated Advisors will no longer be computed. Filtering the Advisors in these ways speeds decisions after learning, without decreasing performance. In ACE, for example, the average testing decision is accelerated by about 30 percent.

Although a tier-3 Advisor with a particularly high weight could be promoted to tier 1, our experience in all three domains has found this a dangerous practice. Instead, FORR partitions tier-3 Advisors retained after weight learning into groups of uneven size, based on their weights. Under this *prioritization*, tier-3 Advisors with the highest weights vote first; only under a tie do subsequent groups of Advisors have an opportunity to comment on the best tied actions. Fewer resources are consumed this way, since ties are relatively rare after the first group or two. With too many groups, prioritization effectively produces a ranked list. (Ranking underperforms a weighted mixture in all three of our domains, for every problem class we have investigated.) Partitions of three to seven subsets produce the best results, but the number of subsets depends on the problem class. We have rarely experienced more than a 10 percent speedup with prioritization. Although fewer Advisors make more mistakes, most representations are still computed, particularly the computationally intensive ones on which class-appropriate Advisors often rely.

Metareasoning can identify portions of the solution process where different behavior is warranted. In some domains, the last part of problem solving is more formulaic. Game players, for example, may have an endgame library and play there by lookup. ACE estimates the last part of CSP search as that after the maximum search depth at which it has experienced a wipeout within the problem class. Below this depth and only during testing, ACE's tier-1 Advisor *Pusher* consults the single highest-weighted tier-3 variable-selection Advisor as if it were in tier 1, bypassing tiers 2 and tier 3 entirely (Epstein, Freuder, & Wallace, 2005). Pushing generally reduces computation time by about 8 percent. ACE does not, however, push value selection. Experiments indicate that one can think less about where to search after the backdoor but not about the values to assign there.

Fast and frugal reasoning is a form of human metareasoning that favors recognized choices and then breaks ties among a random pair of them with one heuristic (Gigerenzer, Todd, & Group, 1999). For ACE, a recognized choice is one made earlier in search (and subsequently retracted) on the same problem. On problems where retractions are common, reusing prior decisions with the highest weighted Advisor to break ties accelerated decision time, despite increased errors (Epstein & Ligorio, 2004).

Conclusions

On difficult problems, errors in a model of expertise may be inevitable, and training instances from the same model may vary in their quality and significance. Nonetheless, a self-aware system can recognize its own prowess or lack thereof, and respond accordingly. Moreover, as we have shown here, a self-aware system can evaluate and reorganize its components to improve its performance.

Constraint solving is a paradigm for many kinds of AI problems. ACE learns to solve problems in many difficult classes, where problems stymie off-the-shelf solvers that cannot monitor and modify their own behavior. Hoyle and Ariadne each learn how to search a single space, a game tree or a maze, from which all the problems in a class are drawn. ACE learns about how to search a *set* of spaces, all of which are supposedly alike, a considerably more difficult task.

Metareasoning is essential in FORR's ability to learn to solve problems within a given class. As it learns to manage heuristics, FORR uses metareasoning to decide when to abandon an unpromising learning attempt, when to stop learning, how to select heuristics during learning, and how to prioritize heuristics. FORR also reasons about its performance on previous problems, its previous decisions, and the relative discriminatory power of its heuristics. Thinking about thinking makes FORR-based applications more incisive and more successful.

Acknowledgments

This work was supported in part by the National Science Foundation under IIS-0811437 and IIS-0739122. ACE is an ongoing joint project with Eugene Freuder and Richard Wallace of The Cork Constraint Computation Centre.

References

Aardal, K. I., van Hoesel, S. P. M., Koster, A. M. C. A., Mannino, C., & Sassano, A. (2003). Models and solution techniques for frequency assignment problems. *4OR: A Quarterly Journal of Operations Research*, 1(4), 261–317.

[Google Scholar](#) [WorldCat](#)

Biswas, G., Goldman, S., Fisher, D., Bhuva, B., & Glewwe, G. (1995). Assessing design activity in complex CMOS circuit design. In P. Nichols, S. Chipman, & R. Brennan (Eds.), *Cognitively diagnostic assessment* (pp. 167–188). Hillsdale, NJ: Lawrence Erlbaum.

Borrett, J. E., Tsang, E. P. K., & Walsh, N. R. (1996). Adaptive constraint satisfaction: The quickest first principle. In W. Wahlster (Ed.), *Proceedings of the 12th European Conference on Artificial Intelligence* (pp. 160–164). Chichester: Wiley. ↵

Crowley, K., & Siegler, R. S. (1993). Flexible strategy use in young children's tic-tac-toe. *Cognitive Science*, 17(4), 531–561. [10.1207/s15516709cog1704_3](https://doi.org/10.1207/s15516709cog1704_3)

[Google Scholar](#) [WorldCat](#) [Crossref](#)

D'Andrade, R. G. (1991). Culturally based reasoning. In A. Gellatly & D. Rogers (Eds.), *Cognition and social worlds* (pp. 795–830). Oxford: Clarendon.

Epstein, S. L. (1994a). For the right reasons: The FORR architecture for learning in a skill domain. *Cognitive Science*, 18(3), 479–511. [10.1207/s15516709cog1803_4](https://doi.org/10.1207/s15516709cog1803_4)

Google Scholar WorldCat Crossref

Epstein, S. L. (1994b). Toward an ideal trainer. *Machine Learning*, 15(3), 251–277. [10.1007/BF00993346](https://doi.org/10.1007/BF00993346)

[Google Scholar](#) [WorldCat](#) [Crossref](#)

Epstein, S. L. (1995). Learning in the right places. *Journal of the Learning Sciences*, 4(3), 281–319. [10.1207/s15327809jls0403_2](https://doi.org/10.1207/s15327809jls0403_2)

[Google Scholar](#) [WorldCat](#) [Crossref](#)

Epstein, S. L. (1998). Pragmatic navigation: Reactivity, heuristics, and search. *Artificial Intelligence*, 100(1–2), 275–322. [10.1016/S0004-3702\(97\)00083-0](https://doi.org/10.1016/S0004-3702(97)00083-0)

Google Scholar WorldCat Crossref

Epstein, S. L. (2001). Learning to play expertly: A tutorial on Hoyle. In J. Fürnkranz & M. Kubat (Eds.), *Machines that learn to play games* (pp. 153–178). Huntington, NY: Nova Science.

Epstein, S. L., Freuder, E. C., & Wallace, R. J. (2005). Learning to support constraint programmers. *Computational Intelligence*, 21(4), 337–371. [10.1111/j.1467-8640.2005.00277.x](https://doi.org/10.1111/j.1467-8640.2005.00277.x)

[Google Scholar](#) [WorldCat](#) [Crossref](#)

Epstein, S. L., & Ligorio, T. (2004). Fast and frugal reasoning enhances a solver for really hard problems. In K. Forbus, D. Gentner & T. Regier (Eds.), *Proceedings of the 26th Annual Conference of the Cognitive Science Society* (pp. 351–356). Mahwah, NJ: Lawrence Erlbaum.

Fukunaga, A. S. (2002). Automated discovery of composite SAT variable-selection heuristics. In *Proceedings of the Eighteenth National Conference on Artificial Intelligence* (pp. 641–648). Menlo Park, CA: AAAI Press.

[Google Scholar](#)
[Google Preview](#)
[WorldCat](#)
[COPAC](#)

Gagliolo, M., & Schmidhuber, J. (2006). Dynamic algorithm portfolios. In *Proceedings of the Ninth International Symposium on Artificial Intelligence and Mathematics*. Fort Lauderdale, FL. Available at <http://anytime.cs.umass.edu/aimath06>.

WorldCat

Gent, I. P., Prosser, P., Walsh, T. (1996). The constrainedness of search. *AAAI/IAAI*, 1, 246–252.

Google Scholar WorldCat

