

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Causal Learning through Repeated Decision Making

Permalink

<https://escholarship.org/uc/item/8r88h20k>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 30(30)

ISSN

1069-7977

Authors

Hagmayer, York
Meder, Bjoern

Publication Date

2008

Peer reviewed

Causal Learning through Repeated Decision Making

York Hagmayer (york.hagmayer@bio.uni-goettingen.de)

Björn Meder (bmeder@uni-goettingen.de)

Department of Psychology, University of Göttingen, Gosslerstr. 14, 37073 Göttingen, Germany

Abstract

Many decisions refer to actions that have a causal impact on other events. Such actions allow for mere learning of expected values, but also for causal learning about the structure of the decision context. Whereas most theories of decision making neglect causal knowledge, causal learning theories emphasize the importance of causal beliefs and assume that people represent decision problems in terms of their causal structure. In three studies we investigated the representations people acquire when repeatedly making decisions to maximize a certain payoff. Our results show that (i) initial causal hypotheses guide the interpretation of decision feedback, (ii) consequences of interventions are used to revise existing causal beliefs, (iii) decision makers use the experienced feedback to induce a causal model of the choice situation, which (iv) enables them to adapt their choices to changes of the decision problem.

Keywords: Decision making, causal models, learning

Introduction

When playing a computer game to build up a new civilization we are repeatedly confronted with the same decisions (e.g., build a new settlement, construct roads, hire soldiers, etc.). By making our choices we may learn different things. First, we could learn which actions increase the value of our civilization the most (e.g., temples are more valuable than houses). However, we can also learn about the causal structure of the system acted upon. For example, we may learn that a sewage system both increases the value of the civilization and prevents later epidemics, which in turn preserves the civilization's value. Particularly this kind of knowledge would later enable us to decide which actions to take when new options arrive (e.g., build a thermal bath) or the situation changes (e.g., a crisis due to an epidemic).

As the example shows, outcomes of different options, their expected values *and* the causal structure of the system may be learnt from repeated decisions. Surprisingly, most theories of decision making do not address the role of causal relations or the decision makers' causal beliefs. Consider likelihood $\times$ value theories, such as expected utility theory, which is still the dominant theoretical framework for consequentialist decision theories. At the heart of these theories is the distinction between decision-makers' beliefs about the situation (options and their possible outcomes), its uncertainties (probabilities, expectations), and the evaluation of the potential outcomes (values, utilities). However, often not probabilities in general, but only probabilities that reflect causal relations (and not merely spurious associations) are relevant for making good decisions, because only causal relations allow it to actively influence a desired outcome. Although this distinction is crucial, no causal learning is assumed by these theories. Similarly, associative learning theories presuppose that relations among actions and out-

comes reflect instrumental (i.e., causal) relations and are not due to other unknown factors. These models capture relations among different outcomes variables, but – again – no distinction between causal and non-causal (i.e., merely statistical) relations is made. Current models of repeated decision making (e.g. Erev & Barron, 2005) agree with associative learning models in assuming that learners do only acquire limited causal knowledge (e.g., action-payoff contingencies). None of these approaches assumes that people learn about the causal texture of the decision problem.

By contrast, causal learning theories propose that people take causal structure explicitly into account (e.g., Gopnik & Schulz, 2007; Waldmann, Hagmayer, & Blaisdell, 2006). Research has shown that, for example, causal beliefs guide the interpretation of covariational information, and that people are able to learn about causal structure through a number of cues (cf. Lagnado, Waldmann, Hagmayer, & Sloman, 2007). Particularly relevant for decision making is also the finding that people use their causal knowledge to derive predictions for interventions not performed yet (Meder, Hagmayer & Waldmann, 2008; Sloman & Lagnado, 2005; Waldmann & Hagmayer, 2005).

The *causal model theory of choice* (Sloman & Hagmayer, 2006) extends causal learning theories to decision making. The basic idea is that people use the available information to induce a causal model of the decision problem and the choice situation. A causal model of the decision problem encompasses knowledge about the structure of the system targeted by the intervention and the interrelatedness of actions, outcomes, and payoffs (cf. Figure 1). Such models enable decision makers to simulate the causal consequences of the available courses of action thereby ensuring that decisions are based on causal and not merely statistical relations. Previous research has shown that people indeed use causal models when making simple one-shot decisions (Hagmayer & Sloman, 2005). However, in these studies participants only made decisions based on hypothetical scenarios without actually engaging in decision making. Thus far, it has not been investigated whether people learn about causal structure and use causal knowledge when making repeated decisions.

The Question of Learning and Representation

The question pursued here is what people learn when repeatedly deciding about interventions on a causal system in order to optimize their payoffs. Figure 1a depicts a causal system, which we used in the studies reported here. It comprises three alternative options (1, 2, 3), three outcome variables (A, B, C), and a final effect variable, which represents the decision maker's payoff. As can be seen from the model, the effect is not directly influenced by any of the available options but only via the intermediate variables A,

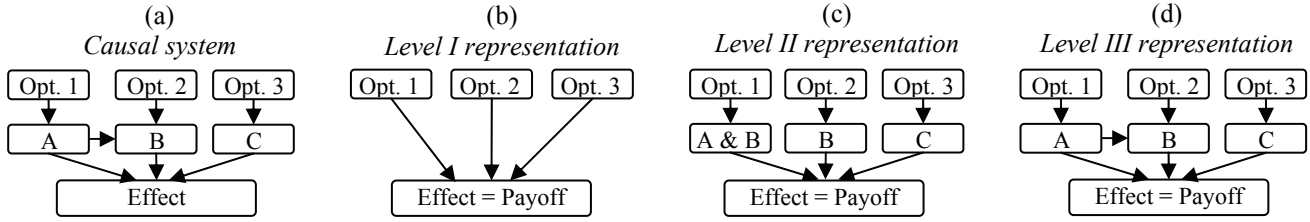


Figure 1: Possible representations of a decision problem.

B, and C. Of particular importance is that the intermediate variables are also interrelated. Variable A not only influences the effect variable but also exerts a causal impact on variable B. By repeatedly choosing among the options people can observe which outcome variables are affected by the chosen action and how these intermediate variables relate to the payoff variable. In this case, option 1 generates variable A and, in turn, B, while options 2 and 3 will generate only a single outcome variable (B and C, respectively). Thus, the experienced feedback comprises information not only about the relation among options and payoffs but also about relations of options, outcomes variables, and payoffs. However, competing theories of decision making and learning differ with respect to what kind of representation they assume to result from the experienced feedback. We here group these representations into three types and discuss them in ascending order of complexity (Levels I to III).

Level I representations include only relations between options (i.e., actions) and payoffs (Figure 1b). Associative learning of relations among actions and payoffs results in such representations. An exemplary approach from the literature on animal learning is *habit learning* (cf. Niv, Joel & Dayan, 2006). The characteristic feature of these representations is that the outcomes connected with each action are not represented separately. Thus, according to these approaches decisions are only driven by the expected values of the available options.

Level II representations encompass options, outcome variables, and the payoffs associated with these variables (Figure 1c). Expected utility theories are the classical example from theories of JDM; a similar approach from the animal learning literature is *goal-directed learning* (cf. Niv et al., 2006). The main difference between knowledge structures of level I and level II is that the latter differentiate the outcomes resulting from an action and the value connected with these outcomes. Therefore they can accommodate motivational shifts or revaluations of outcomes by altering values while preserving the actual option-outcome relations. However, these accounts are not sensitive to the potential causal interrelatedness of the intermediate outcome variables.

Level III (causal model) representations comprise options, their causal relations to outcome variables, the causal relations among these variables, and their relations to payoffs. Figure 1d depicts such a model. Causal model theories (e.g., Waldmann et al., 2006) assume that learners acquire these models through learning. While some theories assume that the structure of the model is derived from cues like temporal delays and interventions (Lagnado, Waldmann, Hagmayer, & Sloman, 2007), others believe that

people primarily use statistical properties of the observed data (e.g., Gopnik et al., 2004; Griffiths & Tenenbaum, 2005). Virtually all causal model theories assume that the parameters of the model (i.e., the strength of the causal relations, base rates of causes, integration rules) are estimated on the basis of the observed covariations. While level II representations are sensitive to the correlation of outcome variables, only causal models distinguish between causal and merely statistical relations, which is often crucial for making good decisions. A formal theory of causal models in learning, reasoning, and decision making is offered by causal Bayes nets theories (e.g., Spirtes, Glymour, & Scheines, 1993; Pearl, 2000).

In summary, a number of different approaches have been proposed to describe learning during repeated decision making. All of them predict that eventually people will shift their choices to the option that maximizes their payoffs. However, different assumptions are made about the acquired representations. Associative learning models assume either the acquisition of level I or level II representations. We, by contrast, propose that people also tend to represent the task with respect to their causal texture and, for example, do not merely encode action-outcome contingencies. More specifically, we hypothesize that (i) decisions are influenced by existing causal beliefs (Exp. 1), (ii) decision makers use consequentialist outcomes to refine their causal hypotheses (Exp. 1), (iii) people will spontaneously induce a causal model representation of the decision problem if no prior knowledge about the causal system is available (Exp. 2a,b), and (iv) will use this knowledge to adapt their choices to changes of the decision problem (Exp. 1 and 2a,b).

Experiment 1

The first goal was to examine how repeated decision making is guided by peoples' causal beliefs. We therefore manipulated decision makers' causal hypotheses about the decision problem while keeping the consequences of the available actions identical across conditions. Thus, potential differences between conditions cannot be attributed to differential learning input. The second goal was to investigate whether participants would use the outcomes of their decisions to revise the initially presented causal model, which was not fully compatible with the experienced feedback. The experiment consisted of two repeated decision making phases, with the second being the test phase. In both phases participants were requested to maximize their payoffs by repeatedly choosing from a set of options. In contrast to studies on causal induction participants were never asked to learn

about causal structure.

Participants and Design. Participants were 36 undergraduates from the University of Göttingen who were randomly assigned to the two causal model conditions.

Material and Procedure. A biological scenario was employed in which three genes (A, B, C) of mice could be activated by injecting them with different types of so-called ‘messenger-RNA’ (U, V, W). The genes, in turn, influenced the level of a growth hormone (the payoff variable). Options U, V, and W as well as the genes A, B, and C were binary variables (active vs. inactive); the value of the payoff ranged from 0-100 points.

Prior to the decision making phase participants were suggested one of two causal structures (see upper row of Figure 2). According to the *causal chain model*, actions U, V, and W affect only one of the genes A, B, or C. However, the genes are interrelated and an activation of gene A also triggers an activation of B ($W \rightarrow A \rightarrow B$). Allegedly, there is also a causal chain relating U to B via C ($U \rightarrow C \rightarrow B$). By contrast, in the *common cause model* (CC) condition participants were instructed that W directly affects genes A and B ($A \leftarrow W \rightarrow B$) and U directly affects B and C ($B \leftarrow U \rightarrow C$). However, both causal models were only partially correct; the true underlying causal models are depicted in the middle row of Figure 2. Contrary to the presented models, ‘Do U’ neither directly nor indirectly affected gene B; thus this intervention generated C only. The table in Figure 2 specifies the feedback decision-makers experience, which was identical across conditions. Choosing ‘Do W’ results in an activation of both A and B and a payoff of +80, ‘Do V’ and ‘Do U’ activate variables B and C, respectively, yielding payoffs of +40 and +60. Choosing not to take any of the options on a particular trial (“no int”) leaves the genes inactive and generates no payoff.

In the initial *Repeated Decision Making (RDM) Phase* participants were requested to maximize their payoffs (i.e., the hormone level) by repeatedly choosing one of the three interventions. Each of the 30 trials referred to a new mouse whose genes were inactive prior to any intervention. Feedback was provided as outlined above. In line with all theoretical accounts we expected participants to quickly switch their preference to option W. In the *Test Phase* a new set of options was introduced. Subjects were informed that instead of interventions U, V, and W they now have only ‘A-RNA’ and ‘C-RNA’ available, which are known to deterministically activate genes A and C, respectively. Then decision-makers were instructed to maximize the hormone level of ten new mice by injecting them with either A- or C-RNA. In this phase, no feedback was provided. Thus, the consequences of the new interventions could not be observed but needed to be inferred from previous learning experiences. Here the two causal models make divergent predictions which interventions people should choose to maximize their payoffs (cf. Figure 2). In the causal-chain condition, due to the causal link $A \rightarrow B$ ‘Do A’ should be assumed to activate A and B. Therefore, the expected payoff is +80 while generating C would only yield a payoff of +60. Thus, decision-

makers should opt for ‘Do A’. By contrast, the common-cause model implies that generating A would only yield a payoff of +40 since A and B are not directly causally related. Accordingly, decision-makers should opt for ‘Do C’. Thus, despite identical previous learning experiences participants should show differential preferences in the test phase because of their causal model assumptions.

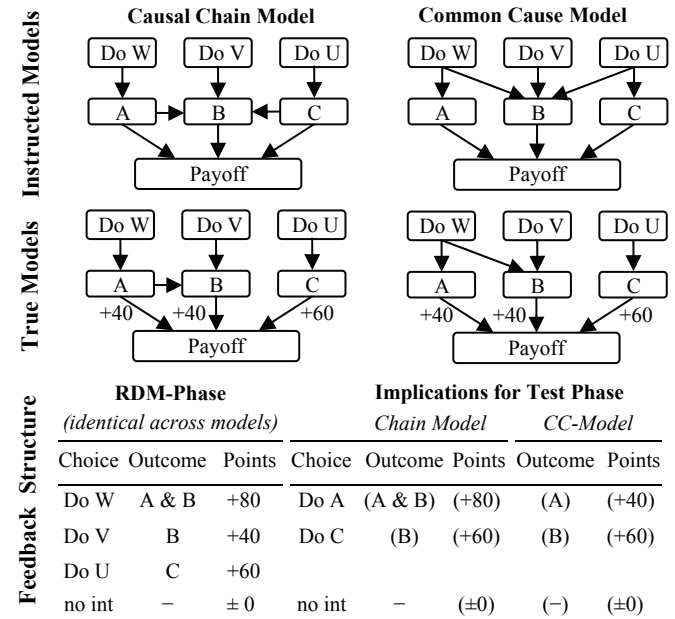


Figure 2: Causal models and feedback structure of Experiment 1. Numbers in parentheses indicate causally expected values (not observable).

Subsequent to the test phase participants were also requested to estimate the expected payoffs for all options during the RDM and the test phase (‘Do U’, ‘Do V’, ‘Do W’, ‘Do A’, ‘Do C’). Thus, they had to give an estimate of how many points they would gain from each possible intervention. Finally, participants were provided with a graph like the one depicted in Figure 2, but without any arrows (i.e., only the variables were depicted). Their task was to express their causal hypotheses by drawing all causal relations the assumed to hold between options, outcome variables, and payoff. This additional test was employed to directly tap onto learners representations and to examine whether learners revised the initially instructed model (i.e., realized that there is no causal link $C \rightarrow B$ and $U \rightarrow B$, respectively).

Results and Discussion. Table 1 depicts participants’ choices for the decision making phase (30 trials) and test phase (10 trials). In the first phase participants exhibited a clear preference for ‘Do W’ regardless of condition. Statistical analyses (t-tests for ‘Do U’, ‘Do V’, ‘Do W’, and ‘no int’) revealed no differences between conditions (all $p > .19$). By contrast, a clear difference between conditions was obtained for the test phase. In line with our predictions participants chose ‘Do A’ significantly more often when assuming a chain model than when being initially presented with a common cause model, $t(34) = 3.16$, $p < .01$. Conversely, the mean of ‘Do C’ choices was higher in the com-

mon cause condition, $t(34) = 3.15, p < .01$.

Table 1. Mean number (SE) of choices in Experiment 1.

Model	RDM-Phase				Test Phase		
	<i>Do W</i>	<i>Do V</i>	<i>Do U</i>	<i>no int</i>	<i>Do A</i>	<i>Do C</i>	<i>no int</i>
Chain Model	20.6 (1.65)	3.2 (.62)	4.5 (.72)	1.7 (.43)	7.8 (.65)	2.1 (.63)	0.1 (.06)
CC-Model	19.1 (1.70)	4.1 (.70)	5.5 (.94)	1.3 (.28)	4.1 (.91)	5.8 (.92)	0.1 (.06)

These choices are also consistent with participants' expected payoffs (Table 2). As with the choices, no differences resulted for the first phase but only for the test phase. In accordance with the respective models participants derived higher ratings for 'Do A' in the causal chain than in the common cause condition, $t(34) = 5.96, p < .001$. Obviously they figured out that 'Do A' would still activate B in the chain but that this was not the case in the common cause condition. Finally, the causal models drawn by the participants revealed that most of them also revised the initial model. In the common cause condition and chain condition 67% and 94%, respectively, correctly stated that 'Do U' only affects C but not B.

Table 2. Means (SE) of expected payoffs in Exp. 1.

Model	RDM-Phase			Test Phase	
	<i>Do W</i>	<i>Do V</i>	<i>Do U</i>	<i>Do A</i>	<i>Do C</i>
Chain Model	78.9 (1.11)	43.3 (1.81)	55.6 (2.58)	70.0 (3.70)	54.4 (3.81)
CC-Model	80.0 (.00)	43.3 (2.43)	54.4 (2.17)	42.4 (2.59)	46.3 (4.02)

Overall, the findings indicate that many participants derived causal representations of the decision task, and these causal hypotheses affected participants' choices in the test phase. Neither level I nor level II representations originating from associative learning can account for the obtained differences because neither the learning input nor participants' choices during the initial decision making phase differed across conditions.

Experiment 2a

In Experiment 1 decision-makers were presented with hypotheses about causal structure prior to the decision making phase. The main goal of Experiment 2 was to examine whether decision-makers would spontaneously induce a causal model representation of the choice task. This would demonstrate that people not only use existing causal beliefs but also strive to infer the causal texture of the choice task. Therefore, no initial models were suggested but learners were only provided with intervention options which would enable them to infer the underlying causal structure from the experienced feedback. Two causal models were used. The first was identical to the true causal model shown in Figure 2 left hand side (i.e., chain model). However, the available interventions were now labeled according to which variable they affected, that is, 'Do A', 'Do B', and 'Do C', respectively. These specific interventions enable learners to infer a

causal model. For example, if the activation of A also results in an activation of B this implies that A causes B and therefore 'Do A' indirectly affects B by way of A. In the second causal model the causal link between A and B was removed (i.e., there were no causal relations among the intermediate variables), although the payoffs were held constant between conditions (cf. Table 3).

Participants and Design. 36 undergraduates from the University of Göttingen participated. They were randomly assigned to the two causal model conditions.

Materials and Procedure. The same materials and procedure as in Experiment 1 were used. In contrast to the first study participants were not presented with a causal model prior to the repeated decision making phase. Again participants were asked to maximize payoffs. In all conditions choosing to intervene on A yields the highest payoff (+80), followed by options C (+60) and B (+40) (Table 3). However, the way variable A generates the outcome differed across conditions. Whereas in Model 1 ($A \rightarrow B$) variable A affected the payoff both directly and indirectly via B with each variable yielding a +40 point payoff, in Model 2 ($A \mid B$) variable A is unconnected to B and directly leads to a payoff of +80.

Table 3. Feedback structure of Experiment 2a.

Choice	Model 1 $A \rightarrow B$				Model 2 $A \mid B$			
	RDM-Phase		Test Phase		RDM-Phase		Test Phase	
Do A	A & B	+80	(A)	(+40)	A	+80	(A)	(+80)
Do B	B	+40	—	—	B	+40	—	—
Do C	C	+60	(C)	(+60)	C	+60	(C)	(+60)
no int	—	± 0	(—)	(± 0)	—	0	(—)	(± 0)

After completing the RDM-phase, in which participants made 30 decisions and received feedback about the resulting consequences, they proceeded to the test phase, in which we presented them with a change of the causal system by confronting participants with mice not possessing gene B. Thus, variable B and the associated option 'Do B' were removed from the causal system. Again participants had to make ten decisions in a row without receiving feedback. The removal of B has diverging implications for the two conditions. Given Model 2 ($A \mid B$) the payoffs of the remaining options 'Do A' and 'Do C' are not affected; therefore intervening in A is still the best choice. By contrast, in Model 1 ($A \rightarrow B$) the removal of B eliminates the causal link from A to B and therefore decreases the *causally* expected value for option A from +80 to +40 (cf. Table 3) making 'Do C' the better choice. In other words, causal learning during decision making should lead to differences in the test phase. If participants merely learn to associate options and payoffs no differences should result and people should stick with their previous preference. Finally, participants were also asked to estimate the expected payoffs of all options in both phases and to express their assumptions about the causal system.

Results. As can be seen from Table 4 people had a clear preference for interventions in A in the first phase regardless of the underlying model. Again there was no difference be-

tween conditions (all $p > .40$). By contrast, in the test phase there was a significant difference between conditions for both interventions, ‘Do A’: $t(34) = 2.81, p < .01$; ‘Do C’: $t(34) = 3.87, p < .01$. These results show clearly that participants considered causal structure. This finding is corroborated by participants’ expected payoffs (Table 5). As for the choices, there was no difference between conditions in the first phase. However, when variable B was removed from the system, subjects’ estimates for the causal impact of A differed across conditions: consistent with the (causally) expected values ‘Do A’ received lower ratings in the $A \rightarrow B$ condition than in the $A \mid B$ condition, $t(33) = 4.59, p < .01$.

Table 4. Mean number (SE) of choices Experiment 2a.

Model	RDM-Phase				Test Phase		
	<i>Do A</i>	<i>Do B</i>	<i>Do C</i>	<i>no int</i>	<i>Do A</i>	<i>Do C</i>	<i>no int</i>
M1	20.7	3.3	3.9	2.1	4.5	5.4	0.1
$A \rightarrow B$	(1.57)	(.52)	(.64)	(.59)	(1.03)	(1.04)	(.01)
M2	20.1	3.9	3.1	2.9	8.1	1.1	0.8
$A \mid B$	(1.84)	(.77)	(.52)	(.93)	(.73)	(.42)	(.53)

Table 5. Means (SE) of expected payoffs in Exp. 2a.

Model	RDM- Phase			Test Phase	
	<i>Do A</i>	<i>Do B</i>	<i>Do C</i>	<i>Do A</i>	<i>Do C</i>
M1	82.9	37.1	60.0	54.7	57.1
$A \rightarrow B$	(2.86)	(1.94)	(4.19)	(5.15)	(2.68)
M2	80.0	38.6	58.6	78.9	56.7
$A \mid B$	(0.00)	(1.54)	(1.54)	(1.18)	(1.91)

Finally, we analyzed whether participants explicitly indicated the presence of the causal link $A \rightarrow B$ in their drawings of the causal model. While many theories of causal learning suggest that people will mentally build causal models, it has been rarely tried to elicit these representations directly. This test intends to directly tap into learners’ representations of the decision context. The results showed a clear difference between conditions; 44% assumed a causal link $A \rightarrow B$ given Model 1 ($A \rightarrow B$) but only 6% given Model 2 ($A \mid B$). However, the fact that only 44% of the learners in the $A \rightarrow B$ condition detected the link also indicates that not all learners spontaneously induced a causal model representation of the choice task. Further analyses revealed a strong concurrence between peoples’ causal beliefs and their choices and estimates: participants who detected the causal relation clearly switched preferences from ‘Do A’ to ‘Do C’. These participants also estimated a reduced payoff of ‘Do A’ in the test phase. By contrast, those who did not indicate the presence of the link $A \rightarrow B$ stuck to their previous preferences and estimates. This finding suggests a high concordance between learners’ choices and expectations and their causal hypotheses.

Experiment 2b

Experiment 2b aimed to replicate and extend the findings of Experiment 2a using the same causal structures but different payoffs. Otherwise the two experiments were identical. Table 6 outlines the feedback structure of Experiment 2b.

In Model 1 ($A \rightarrow B$) the actually negative influence of A (-40) is compensated by the strong positive impact of B (+80). As a consequence, choosing to intervene on A still yields a positive payoff (+40). As before in the alternative condition (Model 2 $A \mid B$) decision makers experience identical payoffs for the options but in this model each intervention affects only one of the variables. For the transfer test, again, variable B was removed from the causal model.

Table 6. Feedback structure of Experiment 2b.

Choice	Model 1 $A \rightarrow B$				Model 2 $A \mid B$			
	RDM-Phase		Test Phase		RDM-Phase		Test Phase	
Do A	A & B	+40	(A)	(-40)	A	+40	(A)	(+40)
Do B	B	+80	—	—	B	+80	—	—
Do C	C	+20	(C)	(+20)	C	+20	(C)	(+20)
no int	—	± 0	(-)	(± 0)	—	± 0	(-)	(± 0)

Results. Table 7 depicts participants’ choices for the two repeated decision making phases. Again no differences between conditions emerged in phase 1 (all $p > .23$); regardless of condition participants exhibited a clear preference for ‘Do B’. By contrast, a clear difference between conditions was obtained for participants’ choices between conditions after the removal of variable B, ‘Do A’: $t(28) = 3.13, p < .01$ and ‘Do C’: $t(28) = 3.06, p < .01$. Thus, with Model 1 ($A \rightarrow B$) underlying the task decision makers preferred interventions in C whereas for Model 2 ($A \mid B$) people predominantly chose ‘Do A’.

Table 7. Mean number (SE) of choices in Experiment 2b.

Model	RDM- Phase				Test Phase		
	<i>Do A</i>	<i>Do B</i>	<i>Do C</i>	<i>no int</i>	<i>Do A</i>	<i>Do C</i>	<i>no int</i>
M1	3.7	22.3	3.0	1.0	3.6	6.2	0.2
$A \rightarrow B$	(.96)	(1.72)	(.72)	(.22)	(1.06)	(1.11)	(.14)
M2	3.3	22.1	3.1	1.6	7.93	1.9	0.2
$A \mid B$	(.79)	(1.81)	(.78)	(.43)	(.89)	(.87)	(.11)

The expected payoffs provide further evidence for subjects’ sensitivity to the underlying causal structure (cf. Table 8). In accordance with the causally expected values a difference was obtained between conditions for participants’ estimates of ‘Do A’ in the test phase, $t(28) = 3.31, p < .01$. Again only half of the participants (47%) correctly indicated the causal link $A \rightarrow B$ in their drawings of the model. Nobody did so given Model 2. Only participants detecting the causal relation in Model 1 clearly ascribed a higher expected payoff to ‘Do C’ than ‘Do A’. Overall, these results corroborate the findings of Experiment 2a.

Table 8. Means (SE) of expected payoffs in Exp. 2b.

Model	RDM- Phase			Test Phase	
	<i>Do A</i>	<i>Do B</i>	<i>Do C</i>	<i>Do A</i>	<i>Do C</i>
M1	57.3	72.0	44.0	28.3	41.3
$A \rightarrow B$	(3.30)	(5.79)	(2.89)	(8.07)	(1.33)
M2	56.0	80.0	44.0	56.0	42.0
$A \mid B$	(2.14)	(.00)	(2.14)	(2.14)	(2.00)

General Discussion

The goal of the studies presented here was to examine the representations decision makers acquire and use when repeatedly choosing among interventions on a causal system. Experiment 1 demonstrated how pre-existing causal beliefs guide decision making, and how the consequences of decisions are evaluated relative to peoples' beliefs. Identical learning experiences resulted in different choices, different expected values for the same options, and different causal model assumptions. This finding challenges associative and other data-driven theories of decision making. Experiments 2a,b showed that many participants spontaneously induced a causal model of the decision problem although they never had been asked to do so. The most important finding throughout the experiments is that people used their causal hypotheses to adapt their choices to changes of the decision problem. For example, in Experiment 2a a chain model assumption made them switch away from the option with the highest expected value when the worst option ('Do B') was removed. Conventional decision theories require new data input to infer the outcomes of interventions not performed yet or to evaluate modifications of the causal system acted upon. By contrast, a causal model representation of the decision problem enables people to flexibly adapt their choices to these changes.

The results of Experiment 2, however, also show that not all decision makers spontaneously induced causal model representations. The models drawn by participants in Experiment 2 as well as their expected values indicate that roughly half of the participants acquired causal level III representations. The decisions and judgments of other participants seemed to be mainly driven by the previously acquired option-payoff contingencies, that is, level I representations. Removing an outcome variable did not influence their preferences and predicted values. This neglect might be rooted in explicitly instructing participants to maximize their payoffs. Thus, it may not have occurred to those participants that acquiring a causal representation may turn out to be useful later on. What about level II representations? The findings of Experiment 1 cannot be explained by level II representations because they do not allow for differentiating between the two causal models, both of which were compatible with the observed feedback. Participants' choices and expected values in Experiment 2, however, are compatible with level II representations. But, only causal representations can explain why only participants that correctly detected the causal link in Model 1 switched their preferences in the test phase.

To sum up, the current findings point out an interesting interplay of causal learning and repeated decision making. While learning to maximize their payoffs people seem also to generate hypotheses about the causal texture of the choice task, which, in turn, affected later decisions. Though learners sometimes had problems to discover all causal relations, the results show a strong concordance between peoples' causal beliefs, their expected payoffs, and their decisions. In future studies we plan to systematically investigate how decision makers' goals (e.g., payoff maximization vs. struc-

ture learning) influence their learning and decision making processes. Studying the tradeoff between causal learning and maximizing payoffs promises to yield interesting new insights into how people make repeated decisions.

References

- Erev, I., & Barron, G. (2005). On adaptation, maximization, and reinforcement learning among cognitive strategies. *Psychological Review*, 112, 912-931.
- Gopnik, A., & Schulz, L. (2007) (Eds.). *Causal learning: Psychology, philosophy, and computation*. Oxford: Oxford University Press.
- Gopnik, A., Glymour, C., Sobel, D. M., Schulz, L. E., Kushnir, T., & Danks, D. (2004). A theory of causal learning in children: Causal maps and Bayes nets. *Psychological Review*, 111, 3-32.
- Griffiths, T. L., & Tenenbaum, J. B. (2005). Structure and strength in causal induction. *Cognitive Psychology*, 51, 354-384.
- Hagmayer, Y., & Sloman, S. A. (2005). A causal model theory of choice. In B. G. Bara, L. Barsalou, & M. Bucciarelli (Eds.), *Proceedings of the Twenty-Seventh Annual Conference of the Cognitive Science Society* (pp. 881-886). Mahwah, NJ: Erlbaum.
- Lagnado, D. A., Waldmann, M. R., Hagmayer, Y., & Sloman, S. A. (2007). Beyond covariation: Cues to causal structure. In A. Gopnik & L. Schulz (Eds.), *Causal learning: Psychology, philosophy, and computation* (pp. 154-172). Oxford: Oxford University Press.
- Meder, B., Hagmayer, Y., & Waldmann, M. R. (2008). Inferring interventional predictions from observational learning data. *Psychonomic Bulletin & Review*, 15, 75-80.
- Niv, Y., Joel, D., & Dayan, P. (2006). A normative perspective on motivation. *Trends in Cognitive Science*, 8, 375-381.
- Pearl, J. (2000) *Causality*. Cambridge, MA: Cambridge University Press.
- Sloman, S. A., & Hagmayer, Y. (2006). The causal psychology of choice. *Trends in Cognitive Science*, 10, 407-412.
- Sloman, S. A., & Lagnado, D. A. (2005). Do we "do"? *Cognitive Science*, 29, 5-39.
- Spirtes, P., Glymour, C., & Scheines, P. (1993). *Causation, prediction, and search*. New York: Springer.
- Waldmann, M. R., Hagmayer, Y., & Blaisdell, A. P. (2006). Beyond the information given: Causal models in learning and reasoning. *Current Directions in Psychological Science*, 15, 307-311.
- Waldmann, M. R., & Hagmayer, Y. (2005). Seeing versus doing: Two modes of accessing causal knowledge. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 31, 216-227.