

The Role of Causal Schemas in Inductive Reasoning

Ralf Mayrhofer (rmayrho@uni.goettingen.de)

Jonas Nagel (jnagel1@uni-goettingen.de)

Michael R. Waldmann (michael.waldmann@bio.uni-goettingen.de)

Department of Psychology, University of Göttingen

Gosslerstr. 14, 37073 Göttingen, Germany

Abstract

Inductive reasoning allows us to go beyond the target hypothesis and capitalize on prior knowledge. Past research has shown that both the similarity of categories and specific knowledge about causal relations affect inductive plausibility. We go one step further and focus on the role of abstract causal schemas about main effects and interactions. Two experiments show that both prior assumptions about abstract causal schemas and the similarity of the corresponding causal effects affect inductive judgments. Reasoners have different prior beliefs about the likelihood of main-effect versus interactive schemas, and rationally combine these prior beliefs with new evidence in a way that can be modeled as Bayesian belief updating.

Keywords: inductive reasoning; causal schemas; causal interactions; Bayesian inference

Introduction

Inductive reasoning occurs in various contexts. In associative learning we infer a general regularity from a set of learning trials. In causal learning we use a sample of observations and go beyond the information given to induce general causal laws. Inductions not only occur at the level of learning exemplars but can also relate prior knowledge about hypotheses to new hypotheses. For example, knowing that cats have hearts allows us to give an informed guess about the probable validity of the hypothesis that wolves have hearts, as well. The interconnectedness of our knowledge is a powerful tool to form informed guesses about new hypotheses.

Although inductive inferences based on exemplars have for a long time been studied in learning, inductions between hypotheses is a fairly recent research goal (see Feeney & Heit, 2007, for an overview). Many early studies have focused on the similarity of categories as the basis for inductive inference.

However, similarity between categories is not the only factor influencing inductive inferences. Based on a variant of causal-model theory, Rehder (2007) has proposed a theory which treats inductive generalizations as causal reasoning. According to this theory we assess the likelihood that a novel feature applies to a new category on the basis of our beliefs about the causal relations that connect that feature to the category. For example, subjects tend to infer that a category has a novel feature if they believe that this feature is caused by or is the cause of a characteristic feature of the category.

Whereas previous research on inductive reasoning has focused on the similarity of categories (i.e., feature overlap) or the causal relation connecting a novel feature to the categories, our current research explores the role of more abstract and complex causal schemas in inductive reasoning about hypotheses.

Causal Schemas in Learning

Thus far, causal schemas have mainly been studied in the context of learning, not inductive reasoning. We will briefly review this research to derive hypotheses for inductive reasoning. Within the causal learning literature it has typically been assumed that the default assumption about multiple causes is that they combine additively. For example, Cheng (1997) has postulated a noisy-OR schema as the default for generative causes according to which multiple causes independently generate an effect.

Although additive integration of multiple causes may be the default, causes may also interact (Kelley, 1972; Novick & Cheng, 2004). The majority of research about additive integration and interactions has been conducted within the associative learning literature. Popular examples of an interaction are positive and negative patterning, in which the effect cannot be predicted on the basis of an additive integration of the individual causes. Positive patterning (PP) refers to learning a situation in which two cues (e.g. A and B), when presented individually, are not paired with the outcome (A- and B- trials), but when presented together they are paired with the outcome (AB+ trials). For example, two drugs may individually not cause a side effect but only in combination. The corresponding additive cue combination (A-, B- => AB-) we henceforth call negative main effect (ME-). In contrast, negative patterning (NP) refers to a situation in which cues A and B, when presented alone, are paired with the outcome (A+ and B+ trials), but when presented together are not paired with the outcome (AB- trials). An example of this interaction are two drugs which each cause a side effect, but cancel out each other's effect when presented together. The complimentary main effect (A+, B+ => AB+) we will call positive main effect (ME+).

Shanks and Darby (1998) found that people can learn both of these interactions (PP and NP) concurrently, and can form the appropriate abstract schematic representations. Moreover, Shanks and Darby demonstrated that people transfer these schemas to new cues which have not previously been shown together. For example, participants

that underwent NP training with cues A and B, and were then shown C+ and D+ trials, could infer that the novel compound CD would not be followed by the outcome (see also Lucas & Griffiths, 2010).

Kemp et al. (2007) have proposed a Bayesian model which explains Shanks and Darby's (1998) data. The model learns causal schemas by monitoring the co-occurrences of cues and outcomes, and groups together cues that behave in a similar fashion. In the NP case this model groups together cues that co-occur with the outcome in isolation, but do not co-occur with the outcome when paired with another cue of the same kind. Importantly, the model can use these cue groupings to generate predictions about novel cue-combinations at test, and thus solve the [C+, D+, CD?] test cases.

Causal schemas differ in learning difficulty in a way consistent with the assumptions in the causal literature. Studies of patterning have shown that learning about patterning schemas is a difficult task and proceeds much slower than when organisms are confronted with main effects (Kehoe, 1988).

Causal Schemas in Inductive Reasoning

Previous research has shown that people are capable of abstracting causal schemas from learning data and transferring them to new situations. However, very little is known about how causal schemas affect inductive plausibility when knowledge is presented as a set of individual facts and hypotheses. Based on the findings in the learning literature and our predecessor study (O. Griffiths et al., 2009), we expect that reasoners bring to bear different priors about causal schemas. We expect them to consider main-effect relations more likely than interactions, especially disordinal ones as in the NP case. Different priors for schema knowledge should therefore constitute one important, hitherto neglected factor influencing inductive plausibility between facts and hypotheses. In particular, a simple application of Bayes' rule predicts that a new instance of an unlikely interaction should have a larger impact on inductive beliefs than a new instance of a schema that is already considered common (e.g., main effects).

A second factor we will explore in the present research concerns the question whether the similarity of the schemas influences induction. Since causal schema hypotheses are abstract not only with respect to the involved cues but also with respect to other properties of the underlying causal relation, similarity between patterning or main effect instances is obviously determined by at least two factors: the similarity of the involved cues in corresponding roles (thus the similarity between the pair of cues A and B and the pair of cues C and D) as well as the similarity of the effects which are generated by A and B on the one hand and C and D on the other hand. In the present experiments we will keep the similarity between cues constant across conditions but we will manipulate the similarity of the effect, a factor that has been neglected in previous research. Moreover, similarity will be investigated in the context of both

confirming and disconfirming evidence. Our main hypothesis is that both confirming and disconfirming evidence should more strongly increase or decrease the prior belief, respectively, if the similarity between the effects mentioned in the facts and the hypotheses is high rather than low.

Schema-based Priors and Belief Updating

O. Griffiths et al. (2009) proposed a simple Bayesian account of schema-based belief updating which models how people update their belief in some schema hypothesis H_i (i.e., H_{PP} , H_{NP} , H_{ME+} or H_{ME-} , respectively) given a confirming instance D via Bayes' rule:

$$P(H_i|D_i) = \frac{P(D_i|H_i)P(H_i)}{P(D_i)}$$

The posterior belief in H_i depends upon the likelihood of the confirming instance D_i given H_i being true, and the participants' prior belief in H_i . An example would be an inference about the hypothesis $(A+, B+) \Rightarrow (AB-)$ when it is already known that the conclusion $(C+, D+) \Rightarrow (CD-)$ is true for novel cues C and D from the same domain as A and B.

Assuming that people consider patterning schemas to be less plausible than main-effect schemas, Griffiths et al. (2009) used this Bayesian belief updating to derive the following predictions: Beliefs regarding patterning schemas will be change more profoundly in response to the observation of a confirming instance than beliefs regarding main-effect schemas. After updating, however, plausibility ratings for patterning schemas should still not exceed those for main-effects.

In an experiment Griffiths et al. (2009) tested these predictions. 32 participants were presented with a series of eight fictitious scenarios describing causal relationships between a number of cues and an effect in several different content domains (e.g., physics or biology). Each of the eight trials consisted of two subscenarios (see Table 1). Subscenario 1 contained three statements: The first two statements, the premises, were labeled as facts (*Fact A* and *Fact B*), and the participants were instructed to treat them as true facts. Each of these premises described one of two cues (A or B) that either did or did not cause an outcome. The third statement, labeled *Conclusion*, was a causal statement about the compound AB that again either did or did not cause the same effect. The distribution of presence or absence of the effect in the three statements determined the cue interaction type of the trial (see Table 1). Participants were then requested to indicate to what extent they believed the conclusion statement to be true as well. Given that the cues and their combinations were novel, these responses were taken as indicators of prior beliefs in the plausibility of the corresponding schema. Afterwards, subjects were presented with the second subscenario, in which they received confirming evidence in the form of three further premises (*Facts 1–3*). These premises described two different cues from the same domain (C and D) and their

compound CD causing or not causing the same effect as in the first subscenario. Moreover, the presented schema was the same. Then the participants were once more asked to indicate the extent to which they believed the conclusion statement about the compound AB to be true, this time in consideration of the additional evidence they had received about C and D.

Table 1: Design of the Experiment by Griffiths et al. (2009)

Subsc.	Statement	Cue Interaction			
		PP	ME-	NP	ME+
1	Fact A	A-	A-	A+	A+
	Fact B	B-	B-	B+	B+
	Conclusion	AB+	AB-	AB-	AB+
2	Fact 1	C-	C-	C+	C+
	Fact 2	D-	D-	D+	D+
	Fact 3	CD+	CD-	CD-	CD+
	Fact A	A-	A-	A+	A+
	Fact B	B-	B-	B+	B+
	Conclusion	AB+	AB-	AB-	AB+

Note. Letters A – D represent causes. Symbols + and – indicate statements in which the cause either produced the effect or did not, respectively. The dashed line separates premises from conclusions. Subsc. = Subscenario. Adapted from Griffiths et al. (2009).

The results of this study were in line with the predictions of Bayesian updating. The baseline ratings in subscenario 1 indicated that participants assigned a higher prior probability to main effects than to patterning interactions. The increase in the ratings between subscenarios 1 and 2 was higher for patterning interactions than for main effects, while the mean ratings in subscenario 2 remained higher for main effects than for patterning interactions even after updating.

Bayesian Belief Updating and Similarity

For the sake of simplicity, Griffiths et al. (2009) assumed that the likelihood $P(D_i|H_i)$ of a confirming instance D_i is represented by some fixed number larger than 0.5 (i.e., the instance is informative) but less than 1 (i.e., the inference from the instance to the hypothesis, which is formulated with respect to another pair of cues, is tainted with uncertainty), and that the likelihood is a function of the similarity between the confirming instance D_i and the instance addressed by the hypothesis H_i .

Since the similarity of instances was not manipulated by Griffiths et al. (2009), it remained open whether this factor influences judgments. Making the similarity component explicit and extending the model to disconfirming instances the proposal of Griffiths et al. (2009) can be revised as:

$$P(D_i|H_i) = \begin{cases} \frac{S(D_i, H_i)}{2} & D_i \text{ confirms } H_i \\ 1 - \frac{S(D_i, H_i)}{2} & D_i \text{ disconfirms } H_i, \end{cases}$$

with $S(D_i, H_i) \in [0,1]$ representing some monotone similarity measure expressing the subjective similarity between the instance D_i and the instance addressed by H_i . Thus, the more similar D_i and H_i are, the stronger the predicted belief update should be (either in the positive direction in the confirming case, or in negative direction in the disconfirming case). Disconfirming evidence is defined as evidence that confirms the contrast hypothesis. For example, $(C+, D+) \Rightarrow (CD-)$ confirms the negative-patterning hypothesis and disconfirms the positive main-effect hypothesis.

Experiment 1

This experiment aims at replicating and extending the results from the experiment by Griffiths et al. (2009). We make a first attempt to manipulate the similarity between the different hypotheses from the same domain. As laid out above, a decrease in similarity between the confirming instance and the instance about which the hypothesis is formulated should decrease the tendency to generalize from the confirmatory evidence to the conclusion statement in question. In the present experiment we will use identical cues in the two subscenarios (i.e., $A=C$, $B=D$) but vary the similarity of the effects. Thus, Bayesian belief updating predicts a stronger increase of inductive confidence from subscenario 1 to subscenario 2 in the high similarity condition as opposed to the low similarity condition.

Method

Participants 48 University of Göttingen undergraduates participated in a series of various unrelated computer-based experiments in our lab either for course credit or for €7/h.

Design The design was closely matched to the experiment by Griffiths et al. (2009). We manipulated two independent variables in the scenarios presented to participants. The first factor was the type of causal schema and had four levels: ME-, ME+, NP, and PP. The second factor was the similarity between the to-be-judged conclusion and the confirmatory evidence. We manipulated whether the two scenarios used the same causal effect or different effects from the same domain.

Each participant responded to a total of 16 scenarios. The scenarios were randomly assigned to the experimental conditions separately for each participant. We used a complete 4 (Cue Interaction: ME-, ME+, NP, PP) $\times$ 2 (Similarity: same effect vs. different effect) repeated-measurement ANOVA design. Each subject thus received two trials in each experimental condition. The trial order was randomly determined for each individual participant.

Materials and Procedure Participants completed the experiment individually on desktop computers. The experiment began with an instruction section which informed participants about the course of their task and briefly tested whether they thoroughly understood it.

Afterwards, participants were presented with 16 fictitious scenarios from different content domains (physics, chemistry, biology, medicine) that were constructed to cover

a broad range of settings. Fictitious cues and effects were used to make sure that participants could not rely on specific prior causal knowledge when making their inferences. Each scenario again consisted of two subscenarios. Subscenario 1 was set up exactly as in Griffiths et al. (2009) (see upper part of Table 1). After having read the two premises (*Facts A–B*) and the *Conclusion*, participants indicated the extent to which they believed the conclusion statement to be true (this rating will be labeled *Rating 1* from now on). For this task participants were provided with an 11-point scale, ranging from “definitely false” at the left-hand end (0) to “definitely true” at the right-hand end (10).

After having provided this rating, participants proceeded to subscenario 2. Its set-up was similar to that in Griffiths et al. (2009) in that three additional facts from the same domain were introduced (*Facts 1–3*; see lower part of Table 1). These facts always constituted confirming evidence for the hypothesis that the conclusion is true. The three statements from subscenario 1 were repeated below the three new statements. Apart from that, we made two important changes in the present experiment regarding the materials of subscenario 2 in order to manipulate the similarity between the to-be-judged conclusion and the confirmatory evidence. First, the confirming evidence consisted of statements about the *same* cues as in the statements in subscenario 1 (i.e., A & B), so that overall similarity was increased compared to the material in Griffiths et al. (2009). Second, we manipulated the similarity between the effect caused by these cues in subscenario 1 and in the new statements of subscenario 2. In half of the trials, cues A and B caused (or did not cause) the same effect in both the to-be-judged conclusion and the provided confirmatory evidence. In the other half, the effect differed between both sets of statements. This means that in all same-effect conditions, in subscenario 2 *Facts 1–2* were identical to *Facts A–B*, and *Fact 3* was identical to the to-be-judged conclusion. Logically, all participants should have indicated certainty about the truth of the conclusion in this condition, since it was already stated as true in the premises. Table 2 shows the material of subscenario 2 in a sample trial.

Table 2. Sample of Subsc. in an NP/Different Effect trial

<i>Fact 1:</i>	Aering Heptosulfin with methane causes the Heptosulfin to become crystalline.
<i>Fact 2:</i>	Aering Heptosulfin with butane causes the Heptosulfin to become crystalline.
<i>Fact 3:</i>	Aering Heptosulfin with a mixture of methane and butane does not cause the Heptosulfin to become crystalline.
<i>Fact A:</i>	Aering Heptosulfin with methane causes the Heptosulfin to become isomorph.
<i>Fact B:</i>	Aering Heptosulfin with butane causes the Heptosulfin to become isomorph.
<i>Conclusion:</i>	Aering Heptosulfin with a mixture of methane and butane does not cause the Heptosulfin to become isomorph.

Following the presentation of the statements, participants were asked once again to rate the *Conclusion*, using the same scale as in subscenario 1. This rating, which was given after confirming evidence had been presented, is from here on referred to as *Rating 2*. Participants then proceeded directly to the next scenario. This process was repeated until all 16 scenarios were complete. The computer program ensured that participants were not able to return to any previous questions.

Results

The results of Experiment 1 are summarized in Figure 1. First, different assumptions about the prior probability of main effects vs. patterning interactions are evident in the much higher mean ratings in *Rating 1* in ME- and ME+ trials compared to NP and PP trials ($F_{3,141}=177.3$, $p<.001$, $\eta_p^2=.79$). Thus, again main-effect schemas were assumed to be more likely than interaction schemas. Second, belief change (increase from *Rating 1* to *Rating 2* within conditions) was influenced by the Cue Interaction factor ($F_{3,141}=28.18$, $p<.001$, $\eta_p^2=.37$), by the Similarity factor ($F_{3,47}=82.39$, $p<.001$, $\eta_p^2=.64$), and by the interaction between both factors ($F_{3,141}=15.18$, $p<.001$, $\eta_p^2=.24$). That is, after receiving positive evidence, participants tended to increase their confidence in the conclusion more in the cases exhibiting patterning-interactions than in the cases exhibiting main-effects. The dependence of belief updates on prior knowledge is predicted by basic Bayesian belief updating and replicates the findings of Griffiths et al. (2009). The belief change is also larger when the cues cause the same effect in both instances rather than a different effect. Furthermore, the interaction indicates that the difference in belief change between same-effect and different-effect conditions was much more pronounced in patterning-interactions than in main-effect trials (planned contrast¹: $F_{1,47}=22.05$, $p<.001$).

Experiment 2

The main goal of Experiment 2 was to investigate how effect similarity and type of evidence (i.e., confirming vs. disconfirming evidence) interact with the type of causal schemas. In Experiment 1 we have already shown that the more similar the instances are, the more confident the participants are in the truth of the hypothesis. Bayesian belief updating predicts that the opposite is expected if disconfirming evidence is presented. To test this prediction, we included disconfirming evidence in half of the trials. In contrast to Experiment 1, we increased the dissimilarity of the cues between the subscenarios to test whether the similarity of the effect event also influences inductive ratings when the cues are more dissimilar.

¹ $(M_{\text{Same/PP\&NP}} - M_{\text{Diff/PP\&NP}}) - (M_{\text{Same/ME}} - M_{\text{Diff/ME}}) = 4.15$

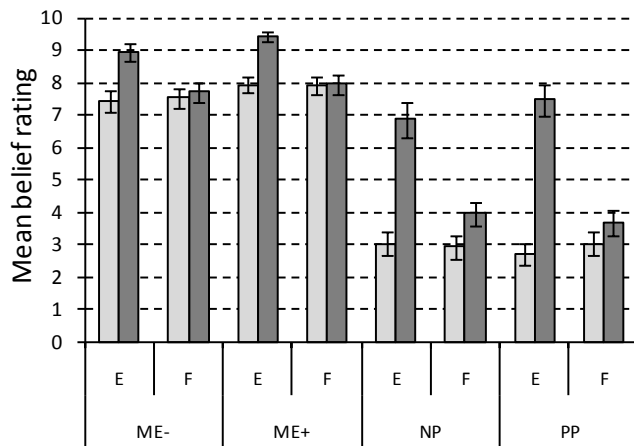


Figure 1: Means of belief ratings of conclusions, and standard errors for pattering-interaction and main-effect schemas (ME-, ME+, NP, PP) plotted against effect similarity (E: same effect, F: different effect). Light grey bars (left hand side) indicate ratings before the confirming instance was shown (*Rating 1*), dark grey bars (right hand side) show ratings after the confirming instance was presented (*Rating 2*).

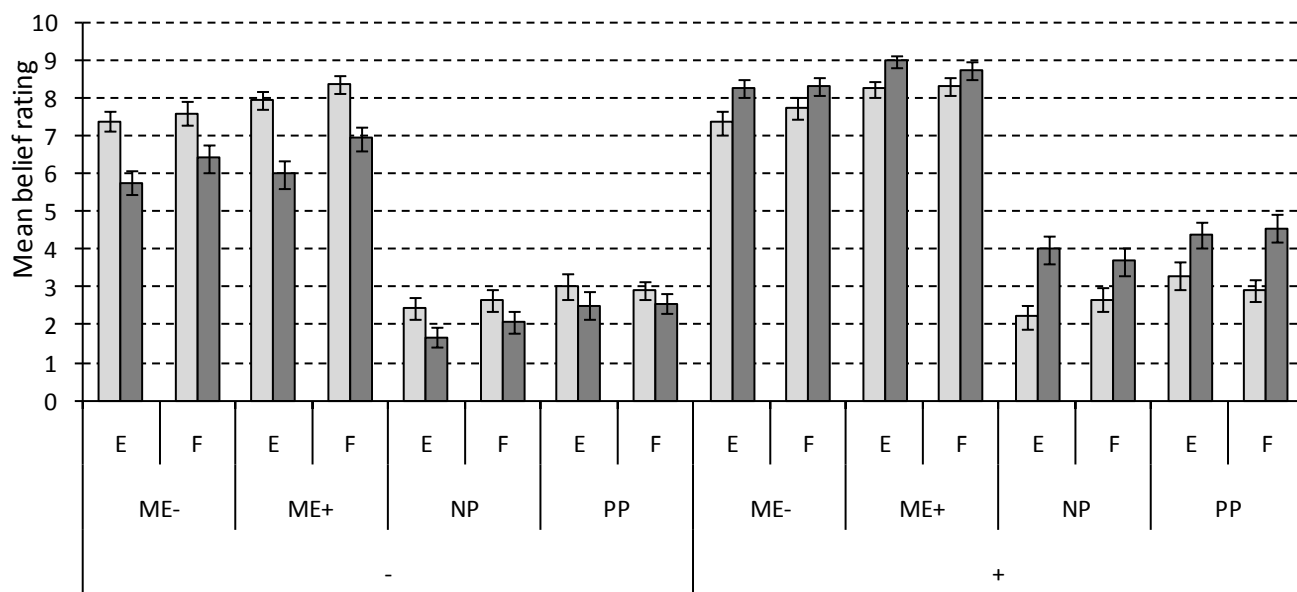


Figure 2: Mean belief ratings in conclusion statements, and standard errors for pattering-interaction and main-effect schemas (ME-, ME+, NP, PP) plotted against effect similarity (E: same effect, F: different effect) and evidence type ("−": disconfirming evidence, "+": confirming evidence). Light grey bars (left hand side) indicate ratings before the (dis)confirming instance was shown (*Rating 1*), dark grey bars (right hand side) show ratings after the (dis)confirming instance was presented (*Rating 2*).

Method

Participants A different sample of 48 University of Göttingen undergraduates participated.

Design We manipulated the same two independent variables as in Experiment 1. Additionally, we varied whether subscenario 2 contained confirming vs. disconfirming evidence for the to-be-judged conclusion statement. This

yielded a complete 4 (Cue Interaction: ME-, ME+, NP, PP) $\times$ 2 (Similarity: same effect vs. different effect) $\times$ 2 (Evidence: confirming vs. disconfirming) repeated-measurement ANOVA design. Each participant again responded to 16 trials, one from each condition.

Materials and Procedure The procedure and materials corresponded to Experiment 1, apart from two changes. First, we manipulated an Evidence factor. In half of the trials, the additional evidence in subscenario 2 was disconfirming rather than confirming evidence. That is, *Facts 1–3* did not instantiate the same causal schema presented in subscenario 1, but rather its complement. If the statements in subscenario 1 formed a positive (negative) main effect, the additional evidence in subscenario 2 was confirming for negative (positive) patterning, and vice versa. The second change was that *Facts 1–3* were no longer about the same cues as *Facts A–B* and *Conclusion* (i.e., A and B), but about different cues from the same domain (C and D). That is, in all same-effect conditions, both instances differed with regards to the cues, whereas they differed with regard to both the cues and the effect in the different-effect conditions. We thus introduced a constant level of dissimilarity in all conditions on the cue level so that in none of the cases the to-be-judged

conclusion was identical to one of the premises.

Results

The results are summarized in Figure 2. First, different assumptions about the prior probability of main effects vs. patterning interactions are evident in the much higher mean ratings in *Rating 1* in ME- and ME+ trials compared to NP

and PP trials ($F_{3,141}=251.6$, $p<.001$, $\eta_p^2=.84$). Thus, main effects were again assumed to be more likely than interactions.

Second, belief change (difference between *Rating 1* and *Rating 2* within conditions) was influenced by Cue Interaction ($F_{3,141}=14.10$, $p<.001$, $\eta_p^2=.23$) and type of additional evidence ($F_{3,47}=124.42$, $p<.001$). Planned contrasts revealed that in the case of confirming evidence, the increase in belief was stronger for patterning interactions than for main effects ($F_{1,47}=11.04$; $p<.01$); in the case of disconfirming evidence, the decrease in belief was stronger for main effects than for patterning interactions ($F_{1,47}=40.06$; $p<.001$), as predicted in (ii).

Finally, there was no main effect of the similarity factor on belief change ($F_{3,47}<1$, $p=.45$). This, however, was to be expected: While the confidence in the hypothesis should increase more after confirming evidence about the same effect than about a different effect, it should also *decrease* more after disconfirming evidence about the same effect than about a different effect (thus, on the level of the similarity factor both effects cancel out each other). This prediction, in turn, is reflected in the significant Evidence $\times$ Similarity interaction ($F_{3,47}=82.39$, $p<.05$, $\eta_p^2=0.09$) which is driven by the predicted specific differences (planned contrast²: $F_{1,47}=4.43$, $p<.05$). These results confirm prediction (iii).

General Discussion

We have presented two experiments which replicate and extend a previous study testing a rational model of belief updating (Griffiths et al., 2009). We showed again that people lacking specific causal knowledge may use knowledge about abstract causal schemas in inductive reasoning. Moreover, we found again that people find interactions less plausible than main effects, while, in line with Bayesian updating, evidence about a case of an interaction increases confidence more than evidence about main effects, which stays at a relatively high level. In the present study we elaborated our model to accommodate variations of similarity and cases of confirming versus disconfirming evidence.

The present research suggests a number of directions for future research. In the present experiments we have shown that the similarity of effect events influences inductive reasoning with both confirmatory and disconfirmatory evidence. It would be interesting to additionally explore the role of the similarity of the cues (A-D), which was only varied across experiments in the present paper. We expect that both the similarity of the cues and of the effect will equally contribute to similarity-related effects.

The present research used extremely abstract materials and a subset of possible interaction types. It might be interesting to look at differences between different causal schemas when more domain knowledge is allowed (see Waldmann, 2007, for other domain related schemas).

Finally, in the Introduction we have separated learning tasks from inductive reasoning tasks, but combinations are conceivable. Previous knowledge need not be stated as facts but can be presented in the form of statistical evidence (e.g., learning trials). It would certainly be interesting to develop a model of inductive reasoning that integrates prior beliefs about abstract and specific causal relations, similarity, and different types of evidence.

Acknowledgments

This research was supported by a research grant of the Deutsche Forschungsgemeinschaft (DFG Wa 621/20).

References

- Feeney, A., & Heit, E. (2007). *Inductive reasoning: Experimental, developmental, and computational approaches*. New York, NY: Cambridge University Press.
- Griffiths, O., Mayrhofer, R., Nagel, J., & Waldmann, M. R. (2009). Causal schema-based inductive reasoning. In N. Taatgen, H. van Rijn, L. Schomaker & J. Nerbonne (Eds.), *Proceedings of the Thirty-First Annual Conference of the Cognitive Science Society* (pp. 691-696). Austin, TX: Cognitive Science Society.
- Kehoe, E. J. (1988). A layered network model of associative learning: Learning to learn a configuration. *Psychological Review*, 95, 411-433.
- Kelley, H. H. (1972). *Causal schemata and the attribution process*. New York: General Learning Press.
- Kemp, C., Goodman, N. D., & Tenenbaum, J. B. (2007). Learning causal schemas. *Proceedings of the Twenty-Ninth Annual Conference of the Cognitive Science Society* (pp. 64-70). Austin, TX: Cognitive Science Society.
- Lucas, C. G., & Griffiths, T. L. (2010). Learning the form of causal relationships using hierarchical Bayesian models. *Cognitive Science*, 34, 113-147.
- Novick, L. R., & Cheng, P. W. (2004). Assessing interactive causal influence. *Psychological Review*, 111, 455-485.
- Rehder, B. (2007). Property generalization as causal reasoning. In A. Feeney & E. Heit (Eds.), *Inductive reasoning: Experimental, developmental, and computational approaches*. New York: Cambridge University Press.
- Shanks, D. R., & Darby, R. J. (1998). Feature- and rule-based generalization in human associative learning. *Journal of Experimental Psychology: Animal Behavior Processes*, 24, 405-415.
- Waldmann, M. R. (2007). Combining versus analyzing multiple causes: How domain assumptions and task context affect integration rules. *Cognitive Science*, 31, 233-256.

² ($M_{\text{Same/Conf}} - M_{\text{Diff/Conf}}$) + ($M_{\text{Diff/Disconf}} - M_{\text{Same/Disconf}}$) = 1.92