

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Causal (In)efficiency: Breaking Markov Violations Through Structural Uncertainty in Causal Chains

Permalink

<https://escholarship.org/uc/item/7hv4f6wp>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 48(0)

Authors

Marchant, Nicolas

Garrido, Diego Ignacio

Chaigneau, Sergio E.

Publication Date

2026

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

Causal (In)efficiency: Breaking Markov Violations Through Structural Uncertainty in Causal Chains

Nicolás Marchant (nicolas.marchant@pucv.cl)
Pontificia Universidad Católica de Valparaíso, Chile

Diego Garrido (diegogarridocerpa@gmail.com) and Sergio E. Chaigneau (sergio.chaigneau@uai.cl)
Centro de Neurociencia Social y Cognitiva, Universidad Adolfo Ibáñez, Chile

Abstract

One of the most persistent findings in causal cognition is that individuals frequently deviate from normative Bayesian reasoning, most notably through Markov violations, where they incorporate conditionally independent information into their judgments. Recent research has sought rational explanations for these deviations, ranging from structural uncertainty to memory sampling limitations. In this study, we evaluated four rational models using a novel experimental paradigm that introduced structural uncertainty into a causal chain by manipulating the functional properties of the causal links. We found that participants systematically violated the Markov assumption when causal links shared the same mechanism. Computational modeling identified the Bayesian Uncertainty Model (BUM) as the most plausible explanation for these results, with sampling-based methods like the Mutation Sampler appearing as close contenders. We discuss these findings in light of the need to rethink the definition of rationality in Bayesian causal reasoning.

Keywords: Causal reasoning; Causal Bayes Networks; Markov condition; Computational modeling

Introduction

One of the most striking and persistent cognitive inefficiencies in causal reasoning literature is that people constantly deviate from the normative standards prescribed by Causal Bayes Nets (CBNs; Rehder, 2024; Rottman & Hastie, 2016; Sloman & Lagnado, 2015; Walsh & Sloman, 2008). CBNs became prominent in cognitive science because they offer an axiomatic formalization of inference judgments through the application of probability calculus and graph theory (Griffiths & Tenenbaum, 2005; Hagmayer, 2016; Pearl, 2000; Rottman & Hastie, 2014). However, although CBNs provide a normative account in several domains of cognition, the fact that people do not behave according to these normative judgments poses a major challenge for CBN-based theories.

In causal reasoning, a well-documented deviation from normativity is the Markov assumption violation (violation of the screening-off rule). Suppose we were asked to reason about a causal chain with three variables: $X_1 \rightarrow Y \rightarrow X_2$. Then, we are asked to infer whether variable X_2 is present. The Markov condition holds that if the state of Y is known, then the state of X_1 should be irrelevant for estimating the likelihood of X_2 . However, people often consider the state of X_1 in their judgment, violating the Markov assumption and deviating from the predicted normative response (Kolvoort et al., 2024; Mayrhofer & Waldmann, 2015; Rehder, 2014, 2018; Rehder & Burnett, 2005; Rehder & Waldmann, 2017).

In recent years, the experimental effort in the literature has been to find some nuance regarding the conditions under which the Markov condition (screening-off rule) is violated. Park and Sloman (2013) showed that Markov violations arise more often when the causal mechanism explanation is the same (i.e., in a chain structure, all causal links operate using a chemical reaction), but people tend to respect the Markov assumption when the causal mechanisms are different (i.e., some links use a chemical reaction and others use a physical mechanism). In another study, Mayrhofer and Waldmann (2015) demonstrated that Markov violations appeared only when the cause was treated as an agent with the disposition to influence a patient (i.e., an effect variable, according to dispositional theories; Wolff & Barbey, 2015). For example, in a causal chain the violation occurred when the judgment was predictive (i.e., from a cause to an effect; sender condition), but it disappeared when the judgment was diagnostic (from an effect to a cause; reader condition). In Rehder and Waldmann (2017) they also showed that Markov violations disappeared in common-cause and common-effect structures when participants were presented only with data information without a clear depiction or graphical representation of the causal structure. When the data was accompanied by a graphical description of the model, or even when the description was presented alone, people did not adhere to the screening-off rule.

In this study, we aim to further investigate why individuals frequently violate the Markov assumption in causal reasoning. We hypothesize that individual-level uncertainty regarding the causal structure significantly influences causal judgments. To introduce structural uncertainty into the causal model, we manipulated the causal link mechanisms by introducing red links that tend to function on hot days and blue links that tend to function on cold days. Additionally, we instructed participants that the current weather was unknown, so they could infer the weather through the colors of the causal links. Our experimental design involved two experimental conditions: a congruent condition, in which both links in a causal chain were of the same color ($X_1 \rightarrow Y \rightarrow X_2$; $X_1 \rightarrow Y \rightarrow X_2$), and an incongruent condition, in which both links were of different colors ($X_1 \rightarrow Y \rightarrow X_2$; $X_1 \rightarrow Y \rightarrow X_2$). To capture individual response variability, we implemented combined inference trials with complete and incomplete information about the causal variables for all nodes in a causal chain.

Furthermore, we fitted our empirical data to four rational models of causal reasoning to elucidate the cognitive mechanisms underlying these axiomatic deviations from CBNs.

Rational accounts of Causal Reasoning

In recent years, several candidate computational models have emerged to account for Markov assumption violations (Chaigneau et al., 2026; Davis & Rehder, 2020; Kolvoort et al., 2023; Marchant et al., 2024; Mayrhofer & Waldmann, 2015; Rehder, 2018, 2024, 2025). In this section, we describe three state-of-the-art rational models—the Bayesian Uncertainty Model (BUM; Chaigneau et al., 2026; Marchant et al., 2023a, 2023b, 2024), Mutation Sampler (MS; Davis & Rehder, 2020; Rehder, 2024, 2025), and Bayesian Mutation Sampler (BMS; Kolvoort et al., 2023, 2025)—that derive normative predictions for causal inference judgments through Noisy-Or function integration. We also introduce the standard normative model (i.e., Causal Generative Model; Pearl, 2000; Rehder, 2017). We do not include non-normative models based on heuristics or hidden causal nodes in the present work (see Rottman & Hastie, 2016, for a brief review of those models).

Normative Model

The Causal Generative Model is widely considered the axiomatic and normative standard for ideal Bayesian reasoning (Pearl, 2000; Rehder, 2017). It assumes that the presence of a cause increases the likelihood of generating its direct effect. When causes act independently, the Noisy-Or function is defined as:

$$Pr(E = 1 | c_1 \dots c_n) = 1 - (1 - b) \prod (1 - m)^{c_i} \quad (1)$$

Here, the probability of a given variable E is determined by the presence of its potential causes $c_1 \dots c_n$ (where $c_i = 1$ if the i -th cause is present, and 0 otherwise). The parameter m denotes the causal strength, and b is the base-rate probability of E in the absence of any causes. If a variable has no direct parents in the causal model, its probability is given by the prior c . In a causal chain ($X_1 \rightarrow Y \rightarrow X_2$), the joint probability is factored as the product of the conditional probabilities: $Pr(X_1)Pr(Y | X_1)Pr(X_2 | Y)$. Critically, the normative model always adheres to the Markov assumption, serving as a benchmark to measure the magnitude of descriptive violations.

Bayesian Uncertainty Model

The Bayesian Uncertainty Model (BUM; Chaigneau et al., 2026; Marchant et al., 2023a, 2023b, 2024) proposes that people deviate from normative Bayesian reasoning and exhibit Markov violations because they are uncertain about what is the true causal structure. Previous accounts have explored this idea in contexts where people are uncertain about the causal structure they are reasoning with, often by assuming invisible nodes or causal relations not present in the explicit model (Mayrhofer & Waldmann, 2015; Meder et al.,

2014; Park & Sloman, 2013). However, BUM assumes that perfect confidence in the given explicit model is impossible; thus, a prior parameter must be introduced to the Noisy-Or integration function. To model this uncertainty, the BUM assumes an alternative counterfactual model in which the causal relations are broken (i.e., a null model). Consequently, BUM assumes that participants weight the evidence according to its likelihood under the given model (H) and the null-model (H'). Joint probabilities are then computed by integrating the distributions $Pr(H | E)$ and $Pr(H' | E)$ through Bayesian Model Averaging (BMA; Raftery et al., 1997), as defined by:

$$P_{BUM}(E) = wP_H(E) + (1 - w)P_{H'}(E) \quad (2)$$

where w functions as the prior belief in the given model H , while $1 - w$ is the prior for the null-model H' . Additionally, $P_H(E)$ is computed following Eq. 1, implementing the same causal parameters (i.e., c , m , and b). In contrast, $P_{H'}(E)$ is computed through the Noisy-Or function but with a different parameterization: here, the causal relation is set to 0 ($m = 0$) to represent a null causal model, and the base rates are represented by a free parameter q specific to H' .

Mutation Sampler Model

The Mutation Sampler Model (MS; Davis & Rehder, 2020; Rehder, 2024, 2025) computes causal inferences by sampling over the state space of a causal system using Markov Chain Monte Carlo (MCMC) methods, such as the Metropolis-Hastings algorithm. However, given that humans likely operate with bounded rationality and lack access to the entire contents of semantic memory, the model assumes that a reduced and biased selection of samples is taken. This limited sampling represents how semantic memory functions in the human mind and offers an explanation for why people deviate from normative reasoning.

According to Davis and Rehder (2020), the model relies on several key assumptions: 1) Sampling always begins at the causal system's prototype—states that are qualitatively consistent with the system's causal structure (e.g., where all variables are present or all absent: $X_1 = 1, Y = 1, X_2 = 1$ or $X_1 = 0, Y = 0, X_2 = 0$, respectively). 2) On each sampling step, one variable is chosen from the current state (q) and mutated (i.e., $0 \leftrightarrow 1$) to generate a proposal state (q'). The acceptance decision is made through the Metropolis-Hastings $a(q' | q)$ rule, defined by:

$$a(q' | q) = \min(1, \frac{\pi(q')}{\pi(q)}) \quad (3)$$

Where $\pi(q)$ is the joint probability of the causal system being in state q . This rule is integral to the MS as it prioritizes high-probability system states while still allowing lower-probability states to be visited. 3) People's sampling capacity is limited to a few samples, represented by the parameter λ . As λ increases, the samples approximate the normative joint distribution; conversely, a small λ results in normative violations, such as violations of the Markov condition.

Consequently, conditional probability queries are computed by averaging the visited states.

Bayesian Mutation Sampler Model

The Bayesian Mutation Sampler (BMS; Kolvoort et al., 2023, 2025) is an extension of the MS model designed to address its inability to fully capture the distributional properties of human judgments. Previously, Kolvoort et al. (2023) showed that the distribution of MS inferences often exhibits “spikes” at 0, 50, and 100%; in contrast, human subjects avoid extreme responses, showing a preference toward the center of the scale (i.e., 50%), a phenomenon often referred to as conservatism. To address these deficiencies, the BMS integrates a β prior into inference queries, moderating the raw sample frequencies. Thus, the BMS operates similarly to the MS by generating random samples using the Metropolis-Hastings algorithm. However, to compute conditional probability queries it introduces the β prior parameter, thus when $\beta > 1$, the prior assigns more probability mass to values near 0.5.

The BMS can account for Markov violations through the same sampling mechanisms as the MS (i.e., biased starting points and limited chain length λ , representing constrained cognitive resources). Consequently, the BMS produces more moderate probability judgments, especially when the sample size is small, resulting in a bias toward conservatism.

Behavioral experiment

Participants

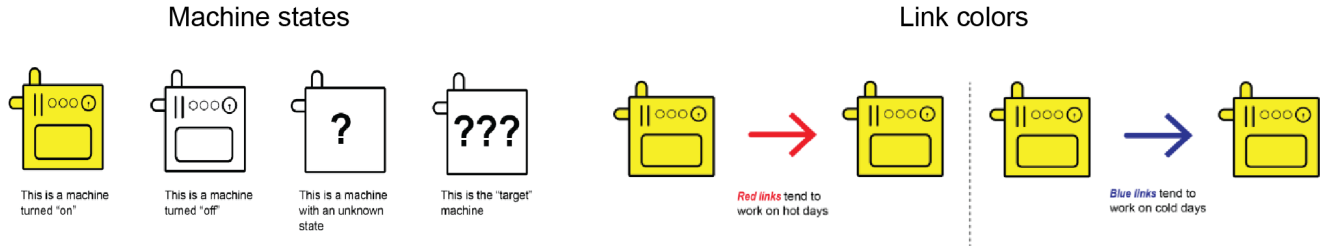
A total of 78 participants (38 females), aged between 20 and 47 years ($M_{age} = 28.19$, $SD_{age} = 6.52$), agreed and provided informed consent to participate in the online experiment. Participants were recruited through the Prolific platform and received monetary compensation (approximately £2.70). We used Prolific’s built-in pre-screening system to ensure that participants were fluent in English (as all instructions and materials were in English), were between 18 and 50 years of age, and had an approval rate on Prolific between 90% and 100%.

Design and materials

We implemented a mixed between-within subjects design: 2 (Condition: Congruent; Incongruent) $\times$ 2 (Color Block Order: first, second) $\times$ 8 (Inference Query), where the last two factors were repeated measures.

Regarding the procedure, we employed pictorial stimuli, in contrast to the common practice in the causal reasoning literature of using verbal-based scenarios in specific knowledge domains. Here, we employed schematic machines (see Fig. 1) that could take only three states: ON (yellow machine), OFF (white machine), or unknown (machine with a ‘?’; See Fig. 1A). These machines were ordered in a straight line following a causal chain structure ($X_1 \rightarrow Y \rightarrow X_2$). Note that all

A.



B.

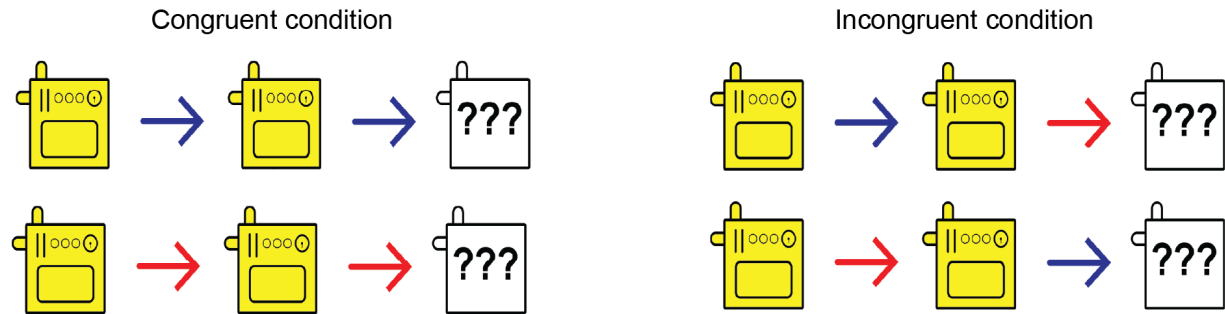


Figure 1: Experimental setup and schematic machine stimuli. Panel A illustrates all possible discrete states of the machines (yellow = ON; white = OFF; ? = unknown; ??? = TARGET) and describes the causal links: red links function on hot days and blue links function on cold days. Note that the instructions specify that the current weather is unknown. Panel B provides examples of trials in both experimental conditions: congruent (links of the same color) and incongruent (links of different colors). This example depicts a predictive judgment where X_2 is the target and the parent nodes are ON (i.e., $Pr(X_2 = 1 | Y = 1, X_1 = 1)$).

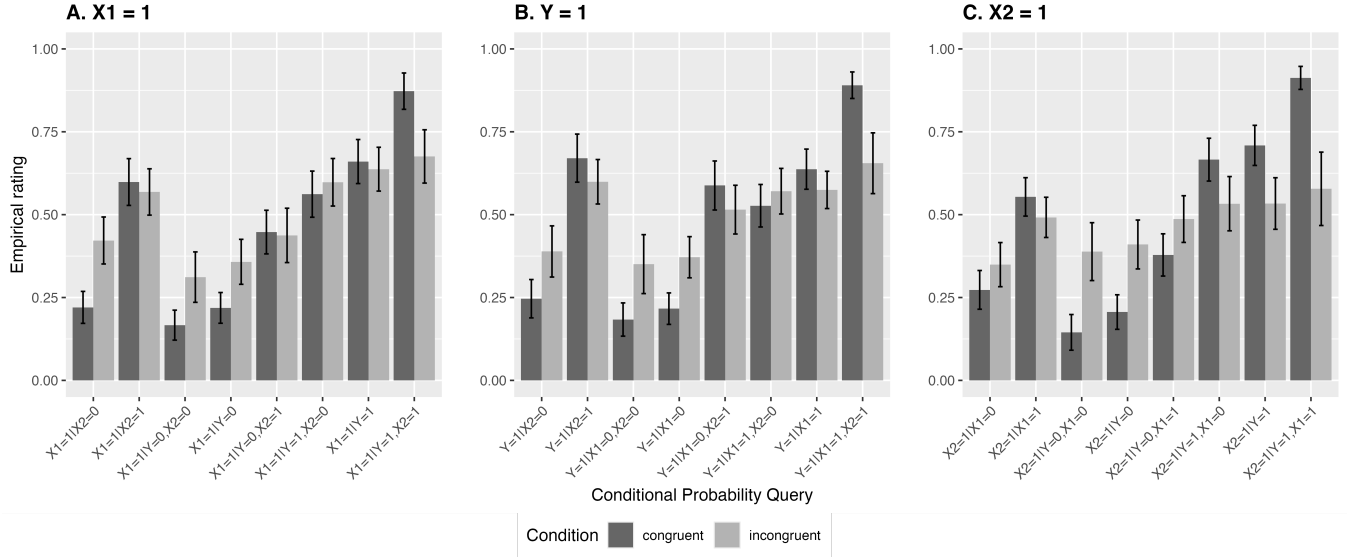


Figure 2: Mean inference judgments across conditions. Panels A, B, and C show diagnostic (X_1 ON), middle-term (Y ON), and predictive (X_2 ON) inferences, respectively. Darker bars indicate the congruent condition; lighter bars indicate the incongruent condition. Error bars represent 95% confidence intervals.

machines had the same size and shape and were created using Adobe Illustrator¹

As participants entered the study, they were asked to read general instructions depicting the states of the machines; they were told that their task was to infer the probability of the machine marked with three question marks (i.e., ???) being ON, based on the information presented on the screen. Participants used a rating scale from 0 to 100 to submit their best guess. Each participant completed 48 trials in total, consisting of 24 trials per color block. We taught participants that the machines were related to each other by arrows that worked differentially according to the weather. Thus, red arrows ($\rightarrow$) tended to work on hot days, while blue arrows ($\rightarrow$) tended to work on cold days. However, in each trial, we stated that the weather for the current machine system was unknown.

Unbeknownst to participants, they were randomly assigned to one of the two between-subjects conditions. In the congruent condition, participants saw an entire machine system with the same colors (i.e., all blue or all red) distributed in two color blocks. In contrast, in the incongruent condition, participants saw a machine system with different colors (i.e., red-blue or blue-red; See Fig. 1B). The presentation of color blocks was randomly assigned across participants.

Behavioral Results

We conducted a repeated measures ANOVA on inference judgments across all conditions. Because the order effects in both conditions (incongruent: hot-cold and cold-hot; con-

gruent: hot-hot and cold-cold) yielded non-significant results (incongruent $p = .60$; congruent $p = .29$), we collapsed the data across the order of presentation. Consequently, as each participant judged each inference query twice, subsequent analyses were performed on the average of each pair. The final model was a 2 (Condition: Incongruent, Congruent) $\times$ 24 (Inference Query; see Figure 2) design, where the latter factor was the repeated measure. The ANOVA revealed a significant interaction between condition and inference query, $F(1, 5.59) = 11.21, MSE = 0.17, p < .001, \eta_p^2 = .13$; a significant main effect of inference query, $F(1, 5.59) = 58.27, MSE = 0.17, p < .001, \eta_p^2 = .43$; and a non-significant main effect of condition ($p = .46$).

To better measure the magnitude of Markov violations, we computed a difference score (Rehder, 2018) for both conditions in predictive (inferring $X_2 = 1$) and diagnostic (inferring $X_1 = 1$) inferences in a causal chain. This difference score for predictive inferences was computed by subtracting $Pr(X_2 = 1 | Y = 1, X_1 = 0)$ from $Pr(X_2 = 1 | Y = 1, X_1 = 1)$. Similarly, the diagnostic difference score was computed by subtracting $Pr(X_1 = 1 | Y = 1, X_2 = 0)$ from $Pr(X_1 = 1 | Y = 1, X_2 = 1)$.

To assess the magnitude of Markov violations through difference scores, we conducted a 2 (Condition: Incongruent, Congruent) $\times$ 2 (Markov Direction: Diagnostic ($X_1 = 1$), Predictive ($X_2 = 1$)) mixed ANOVA. The results revealed a non-significant interaction ($p = .67$) and a non-significant main effect of Markov direction ($p = .18$). However, we found a significant main effect of condition, $F(1, 76) = 20.86, MSE = 0.09, p < .001, \eta_p^2 = .22$. As shown in Fig. 3, regardless of the direction of the Markov violation (i.e., diagnostic or predictive), the magnitude of the violation was consistently greater

¹The dataset involved in our analysis was collected as part of a larger project in collaboration with Tadeq Quillien, whose results will be reported elsewhere.

in the congruent condition compared to the incongruent condition. This result contrasts with those from Mayrhofer and Waldmann (2015), in which Markov violations disappeared in diagnostic inferences, and further supports the hypothesis that structural uncertainty induced by causal link information may drive deviations on causal reasoning regardless of the type of inference being performed.

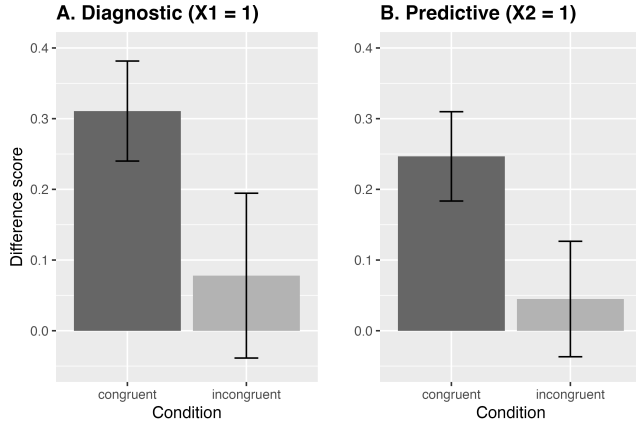


Figure 3: Difference scores for inferences of interest. Panel A displays difference scores for diagnostic inferences (where X_1 is ON), and Panel B displays difference scores for predictive inferences (where X_2 is ON). These scores quantify the magnitude of Markov violations by calculating the difference between the two critical conditional probability queries in each direction.

Computational modelling

Model fitting procedure

To assess the computational mechanisms underlying the occurrence of Markov violations in human inference judgments, we compared the four theoretical approaches introduced previously. As mentioned in the introduction, all of these are rational models, as they represent generalizations of Bayesian inference. We fitted all four computational models to individual participant data using a grid search followed by optimization via the *dfoptim* package in R, employing the Nelder-Mead algorithm (*nmkb* function with a bounded parameter space). All models were fitted by minimizing the Mean Squared Error (*MSE*) between model predictions and participants' empirical responses. We assigned different grid values for initialization:

For the Noisy-Or function parameters (c , m , and b)², we used a grid from 0.05 through 0.95, in steps of 0.15, constrained between 0 and 1. For the BUM, two free parameters were added, w and q , which followed the same grid space.

²Note that all rational models derive predictions by implementing the Noisy-Or function to compute joint probabilities; thus, the same grid was used across all models. However, each parameter of the Noisy-Or is fitted separately for each model.

For the MS and BMS, we added a λ parameter representing the chain length, with grid values set at 2, 4, 8, 12, 16, and 32, constrained between 2 and 64. The BMS adds an extra parameter, the β prior, with grid values set at 0.053, 0.250, 0.538, 1, 1.86, 4, and 19, constrained between 0 and 24. Critically, because the MS and BMS are stochastic models—meaning that each run generates a randomly generated sample yielding different conditional probabilities through the Metropolis-Hastings algorithm—we fitted the MS and BMS by computing the expected joint distribution of the samples (i.e., the ideal distribution achievable if an infinite number of samples were taken and averaged). This expected joint distribution is constructed via a transition matrix where each entry represents an MCMC transition probability between mutated states. Given that each state in the conditional probability task has exactly three nodes, each node has a mutation probability of exactly $1/3$. This expected joint distribution is then iteratively multiplied λ times, starting from the prototype states. In summary, this transition matrix approach allows model prediction probabilities to be computed deterministically rather than stochastically.

Model comparison was conducted using the Akaike Information Criterion (AIC), as it penalizes for the number of free parameters, following the equation:

$$AIC = n * \log(SSE/n) + 2(k + 1) \quad (4)$$

where $n = 24$ (reflecting the averaging of the two repeated trials per inference query to make computational modeling more approachable; no statistical differences were found between color block orders), *SSE* is the Sum of Squared Errors computed via *MSE* as $SSE = MSE * n$, and k is the number of free parameters. For each participant, we identified the best-fitting model as the one with the lowest AIC value.

Model fitting results

Table 1 shows the average AIC results and the percentage of subjects best fitted in each condition. As can be seen in the congruent condition (links of the same color), the BUM achieves the lowest AIC values ($AIC = -92.1$) and best fits 43.6% of subjects, followed closely by the MS. Notably, the congruent condition is where Markov violations arise, as shown in Figure 3. In the incongruent condition (links of different colors), the BUM again achieves the lowest average AIC ($AIC = -81.2$); however, the normative model fits the largest proportion of subjects (56.4%), followed by the MS. Interestingly, the BMS achieves poor fits across both conditions. This is intriguing because, in the incongruent condition, participants pulled their inference judgments toward 50% of the scale and generally did not exhibit Markov violations, as seen in Figure 2.

Table 2 presents the median best-fitting parameters for each model across conditions. As previously described, each model computes joint probabilities via the Noisy-Or function, extending this foundation based on its specific axiomatic assumptions. Specifically, the BUM incorporates the w and q

Table 1: Results of model fitting by AIC and % of best fitted subjects in both conditions

Condition	Normative		BUM		MS		BMS	
	AIC	%S	AIC	%S	AIC	%S	AIC	%S
Congruent	-82.2	15.4	-92.1	43.6	-90.0	38.5	-86.0	2.6
Incongruent	-78.7	56.4	-81.2	17.9	-81.0	25.6	-78.3	0

parameters; the MS adds the λ parameter; and the BMS adds both λ and β . These parameter results are addressed in the Discussion section.

Table 2: Median best-fitting parameter values

model	condition	c	m	b	w	q	λ	β
Normative	Congruent	.45	.64	.23				
	Incongruent	.50	.36	.40				
BUM	Congruent	.41	.51	.13	.58	.47		
	Incongruent	.51	.39	.36	.42	.48		
MS	Congruent	.37	.28	.30			3.41	
	Incongruent	.49	.11	.45			16.5	
BMS	Congruent	.35	.50	.20			4	.05
	Incongruent	.50	.35	.35			12	.54

Discussion

In this work, we investigated the conditions under which individuals deviate from normative Bayesian reasoning in causal inference, focusing on violations of the Markov assumption (Park & Sloman, 2013; Rehder, 2018; Rehder & Burnett, 2005; Rottman & Hastie, 2016). Using a novel paradigm that introduced structural uncertainty through causal link mechanisms (i.e., red links that tend to work on hot days, and blue links that tend to work on cold days), we found that Markov violations were consistently larger in the congruent condition, regardless of whether the inference was predictive or diagnostic. We interpret these behavioral results as indicating that the congruent condition introduces structural uncertainty via the covariational information about the shared mechanism conveyed by causal link colors (e.g., if two links are red, my belief that the weather is hot increases). As a reviewer suggested, this can be examined by inspecting the inference queries in Figure 2. Consider, for example, the predictive inference in Figure 2 panel C, where both links are red: $X_1 \rightarrow Y \rightarrow X_2$, in where $P(X_2 = 1 | Y = 1, X_1 = 1)$ is greater than $P(X_2 = 1 | Y = 1, X_1 = 0)$, because observing $X_1 = 1$ instantiates the belief that the weather is hot, raising the expectation that the target machine is also on ($X_2 = 1$). However, when $X_1 = 0$, uncertainty about causal link functioning is introduced: something else may have caused $Y = 1$ even though the link mechanism is nominally the same. This creates ambiguity about whether the link color operated as intended given a hot day, thereby reducing the probability that $Y = 1$ will cause $X_2 = 1$ and pulling the estimated judgment toward 50%.

In contrast, participants in the incongruent condition were

insensitive to the covariational information about causal link mechanisms, leaving them uncertain about whether the weather was cold or hot (i.e., both weather conditions were approximately equally likely). This uncertainty pulled judgments toward the 50% mark on the estimated inference probability, which may reflect a conservative response strategy among these participants.

To elucidate the cognitive drivers on causal reasoning normative deviations, we implemented computational modeling. The modeling results suggests that these deviations can be captured by a set of different rational strategies. Overall, the Bayesian Uncertainty Model (BUM) and the Mutation Sampler (MS) emerged as the best-performing models across conditions (see Table 1). The BUM explains violations by assuming structural uncertainty of the given causal model, whereas the MS attributes deviations to constraints on cognitive resources during sampling. Importantly, while both approaches provide a good descriptive account at the level of model fit, they make distinct commitments at the parameter level, which has consequences for how these mechanisms should be interpreted across experimental conditions.

An inspection of the parameter estimates in Table 2 raises interpretive issues. In both the MS and BMS, the fitted λ parameter is larger in the incongruent than in the congruent condition. If the MS and BMS's λ is interpreted as the length of the memory search chain, larger values imply more extensive sampling and, consequently, judgments closer to normative Bayesian inference. However, there is no clear theoretical reason to expect participants to engage in more exhaustive memory search precisely when the causal structure is less coherent.

This pattern therefore challenges a straightforward resource-based interpretation of λ in the present task. In contrast, the BUM parameters exhibit a more transparent and psychologically interpretable pattern: the weight w is clearly lower in the incongruent condition, indicating reduced confidence in the instructed causal model when link colors conflict, while the q parameter remains stable across conditions, suggesting that the alternative null-model is represented similarly regardless of congruency.

Future directions emerge from this study. In particular, it would be valuable to generalize the BUM and MS within a Hierarchical Bayesian Framework; it is possible that people represent different causal relations within the same model, e.g., using simple links in the incongruent versus the complete structure in the congruent condition. Overall, this contribution expands the discussion on causal cognition, advocating that the Markov violation is not a cognitive failure, but a crucial piece of the puzzle in rethinking human rationality.

Acknowledgment

This work received financial support to the first author from ANID Fondecyt (Fondo Nacional de Desarrollo Científico Tecnológico, Chile) grant 11250063.

References

- Chaigneau, S. E., Marchant, N., & Rehder, B. (2026). Breaking the chains of independence: A Bayesian uncertainty model of normative violations in human causal probabilistic reasoning. *New Ideas in Psychology*, 81, 101231.
- Davis, Z. J., & Rehder, B. (2020). A process model of causal reasoning. *Cognitive Science*, 44(5), e12839.
- Griffiths, T. L., & Tenenbaum, J. B. (2005). Structure and strength in causal induction. *Cognitive psychology*, 51(4), 334–384.
- Hagmayer, Y. (2016). Causal bayes nets as psychological theories of causal reasoning: Evidence from psychological research. *Synthese*, 193, 1107–1126.
- Kolvoort, I., Davis, Z., Rehder, B., & van Maanen, L. (2025). Models of variability in probabilistic causal judgments. *Computational Brain & Behavior*, 1–27.
- Kolvoort, I., Fisher, E. L., van Rooij, R., Schulz, K., & van Maanen, L. (2024). Probabilistic causal reasoning under time pressure. *Plos One*, 19(4), e0297011.
- Kolvoort, I., Temme, N., & van Maanen, L. (2023). The bayesian mutation sampler explains distributions of causal judgments. *Open Mind*, 1–32.
- Marchant, N., Puebla, G., Quillien, T., & Chaigneau, S. (2024). Rationally uncertain: Investigating deviations from explaining away and screening off in causal reasoning. *Proceedings of the Annual Meeting of the Cognitive Science Society*, 46.
- Marchant, N., Quillien, T., & Chaigneau, S. E. (2023a). A context-dependent bayesian account for causal-based categorization. *Cognitive Science*, 47(1), e13240.
- Marchant, N., Quillien, T., & Chaigneau, S. E. (2023b). Uncertainty can explain apparent mistakes in causal reasoning. *Proceedings of the 45th Annual Meeting of the Cognitive Science Society*, 45.
- Mayrhofer, R., & Waldmann, M. R. (2015). Agents and causes: Dispositional intuitions as a guide to causal structure. *Cognitive science*, 39(1), 65–95.
- Meder, B., Mayrhofer, R., & Waldmann, M. R. (2014). Structure induction in diagnostic causal reasoning. *Psychological Review*, 121(3), 277.
- Park, J., & Sloman, S. A. (2013). Mechanistic beliefs determine adherence to the Markov property in causal reasoning. *Cognitive Psychology*, 67(4), 186–216.
- Pearl, J. (2000). *Causality: Models, reasoning and inference*. Cambridge University Press.
- Raftery, A. E., Madigan, D., & Hoeting, J. A. (1997). Bayesian model averaging for linear regression models. *Journal of the American Statistical Association*, 92(437), 179–191.
- Rehder, B. (2014). Independence and dependence in human causal reasoning. *Cognitive Psychology*, 72, 54–107.
- Rehder, B. (2017). Concepts as causal models: Induction. *The Oxford handbook of causal reasoning*, 377.
- Rehder, B. (2018). Beyond markov: Accounting for independence violations in causal reasoning. *Cognitive Psychology*, 103, 42–84.
- Rehder, B. (2024). Extending a rational process model of causal reasoning: Assessing Markov violations and explaining away with inhibitory causal relations. *Journal of Experimental Psychology: Learning, Memory, and Cognition*, 50(9), 1463.
- Rehder, B. (2025). A magic act in causal reasoning: Making Markov violations disappear. *Entropy*, 27(6), 548.
- Rehder, B., & Burnett, R. C. (2005). Feature inference and the causal structure of categories. *Cognitive Psychology*, 50(3), 264–314.
- Rehder, B., & Waldmann, M. R. (2017). Failures of explaining away and screening off in described versus experienced causal learning scenarios. *Memory & Cognition*, 45(2), 245–260.
- Rottman, B. M., & Hastie, R. (2014). Reasoning about causal relationships: Inferences on causal networks. *Psychological bulletin*, 140(1), 109.
- Rottman, B. M., & Hastie, R. (2016). Do people reason rationally about causally related events? markov violations, weak inferences, and failures of explaining away. *Cognitive psychology*, 87, 88–134.
- Sloman, S. A., & Lagnado, D. (2015). Causality in thought. *Annual review of psychology*, 66, 223–247.
- Walsh, C., & Sloman, S. (2008). Updating beliefs with causal models: Violations of screening off. In M. A. Gluck, J. R. Anderson, & S. M. Kosslyn (Eds.), *Memory and mind: A festschrift for Gordon H. Bower* (pp. 345–358). Lawrence Erlbaum Associates Publishers.
- Wolff, P., & Barbey, A. K. (2015). Causal reasoning with forces. *Frontiers in human neuroscience*, 9, 1.