

ROCK: A Causal Inference Framework for Reasoning about Commonsense Causality

Anonymous ACL submission

Abstract

Commonsense causality reasoning (CCR) aims at identifying plausible causes and effects in natural language descriptions that *deemed reasonable by an average person*. Although being of great academic and practical interest, this problem is still shadowed by the lack of a well-posed theoretical framework; existing work usually relies on various notions of correlation and is susceptible to *confounding co-occurrences*. This paper articulates the central question of CCR and develops a novel framework, ROCK, to Reason O(A)bout Commonsense K(C)ausality based on classical causal inference principles. ROCK leverages temporal signals as incidental supervision, and makes use of *temporal propensities* that are analogous to propensity scores (Rosenbaum and Rubin, 1983) for balancing confounding effects. We implement a modular zero-shot pipeline, which is effective and demonstrates good potential for CCR on various datasets.

1 Introduction

Commonsense causality reasoning (CCR) has been an important yet non-trivial task in natural language processing (NLP) communities that exerts broad industrial and societal impacts (Kuipers, 1984; Gordon et al., 2012; Mostafazadeh et al., 2020; Sap et al., 2020). We articulate this task as

reasoning about plausible causes and effects of semantic meanings conveyed by natural languages that are accessible by an average reasonable person.

This definition naturally excludes questions that are beyond commonsense knowledge, such as those scientific in nature (e.g., does a surgery procedure reduce mortality?). Instead, it accommodates causal queries within the reach of an ordinary reasonable person. For example, in Figure 1, is Alice’s “entering a restaurant” (E_1) a plausible cause for her “ordering a pizza” (E_2)? Although the precedence

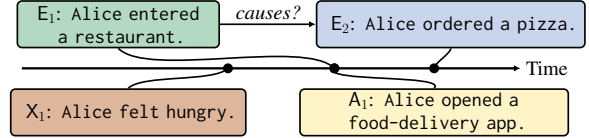


Figure 1: **An Example of CCR: does E_1 cause E_2 ?** The temporal order $E_1 \prec E_2$ does not necessitate causation due to confounding co-occurrences (e.g., X_1). Since when *conditioning on X_1* , a *comparable* intervention A_1 of E_1 also precedes E_2 , the effect from E_1 to E_2 shrinks.

from E_1 to E_2 is logical, it might be less a “cause” compared with Alice’s “feeling hungry” (X_1).

We know that temporality certainly informs causation, but how to account for confounding co-occurrences such as X_1 ? Motivated by causal principles (Section 2), we formulate CCR as estimating the *change* in the likelihood of E_2 ’s occurrence due to intervening E_1 (denoted by $\neg E_1$):

$$\Delta = \mathbb{P}(E_1 \prec E_2) - \mathbb{P}(\neg E_1 \prec E_2) \quad (1)$$

where $\mathbb{P}(\cdot)$ can be estimated by language models (LMs) e.g., via mask language modeling (Section 4). If the occurrence of E_1 and $\neg E_1$ is purely random, Δ directly estimated informs CCR; however, due to confounding co-occurrences (X_1), one need to *balance* the covariates (events that precede E_1) to eliminate potential spurious correlations. In light of the capabilities and limitations of LMs, we propose *temporal propensity*, a surrogate propensity score for this purpose (Section 3). Indeed, we show in Section 5 that *although temporality is essential for CCR, it is vulnerable to spurious correlations without being properly balanced*.

Contributions. We articulate CCR from a completely new perspective using causal inference principles, and our contributions include (i) a novel causal framework for CCR that is the first of its kind to the best of our knowledge; (ii) mitigating confounding co-occurrences by matching temporal propensities; (iii) a modular pipeline for zero-shot CCR with demonstrated effectiveness.

2 Background

The problem of reasoning causal relationships, and differentiating them from innocuous associations has been contemplated and studied extensively in human populations research spanning clinical trials, epidemiology, political and social sciences, economics, and many more (Fisher, 1958; Cochran and Chambers, 1965; Rosenbaum, 2002; Imbens and Rubin, 2015), among which causal practitioners usually base their studies on the potential outcomes framework (also known as the Rubin causal model, see Neyman, 1923; Rubin, 1974; Holland, 1986) and the structural equation model (Pearl, 1995; Heckman, 2005; Peters et al., 2017) in observational studies; Granger causality (Granger, 1969) for time series data, to name a few.

With the recent celebrated empirical success of language models especially transformers (Devlin et al., 2019; Radford et al., 2019), there is an increasing interest in the NLP community on drawing causal inference using textual data, where the major focus treats textual data as either covariates or study units (Keith et al., 2020; Feder et al., 2021) on which causal queries are formed (e.g., does taking a medicine affects recovery, which are recorded in textual medical records?). On the other hand, CCR with natural language descriptions struggles to fit in a causal inference framework: *textual data in this case are just vehicles conveying semantic meanings, not to be taken at the face value*, hence it is difficult to draw the parallel between classical causal inference that requires a clear definition of study units, treatments, and outcomes.

2.1 Existing Approaches

Existing works related to CCR are usually grouped under the umbrella term of commonsense reasoning (Rashkin et al., 2018; Ning et al., 2019a; Sap et al., 2020) or event detection (O’Gorman et al., 2016). Some of the notable progresses usually come from leveraging explicit causal cues/links (tokens such as “due to”) and use conditional probabilities to measure “causality” (Chang and Choi, 2004; Do et al., 2011; Luo et al., 2016); leveraging deep LMs via augmenting training datasets, designing training procedures and/or loss functions (Sap et al., 2019; Shwartz et al., 2020; Tamborrino et al., 2020; Zhang et al., 2021; Staliunaite et al., 2021).

There are several works that share similarities in perspective with ours, yet different in various ways: Granger causality, which measures associa-

tion, is used by Kang et al. (2017) to detect event causes-and-effects; Bhattacharjya et al. (2020) studies events as point-processes, in a way arguably closer to association; Gerstenberg et al. (2021) uses a simulation model to reason physical causation. To the best of our knowledge, there is very little literature, if not none, on adopting a causal perspective in solving CCR.

2.2 Challenges of CCR

Many existing CCR methods (mostly supervised) are based on ingenious designs and creative LM engineering, theoretical justifications, however, are sometimes desirable, as only then do we know how general they can be. Indeed, recent studies reveal that several supervised models may have exploited certain artifacts in datasets to ace the evaluations (Kavumba et al., 2019; Han and Wang, 2021).

This dilemma of constructing a well-founded theoretical framework versus engineering models to achieve excellent empirical performances is not surprising, perhaps, given that the challenges of CCR from causal perspectives are not trivial at all: what is the study unit, treatment, and outcome in this case? What does it mean to “intervene”, or “manipulate” the treatment? Is treatment *stable*, or is it desirable to consider multiple versions of it?

2.3 Principles of the ROCK Framework

In this paper, we attempt to these questions using, among several causal principles, the following two that are intuitive and directly appeal to human nature (see e.g., Russell, 1912; Bunge, 1979):

- Precedence does not imply causation.
- Causation implies precedence.

Here the first principle warns us of *post-hoc* fallacies; the second informs us that the events must be compared with those that are *in pari materia* (Mill, 1851; Hill, 1965), or having *balanced* covariates (also called “pretreatments,” which we mean events that occur prior to E_1 , cf. Rosenbaum, 1989). Our CCR formulation in terms of temporality has several benefits: (i) the intrinsic temporality of causal principles characterizes its central role in CCR; (ii) temporal signals bring about incidental supervision (Roth, 2017; Ning et al., 2019a); (iii) although being a non-trivial question *per se*, reasoning temporally has witnessed decent progresses lately, making it more accessible than directly detecting causal signals (Ning et al., 2017, 2018, 2019b; Zhou et al., 2020; Vashishtha et al., 2020).

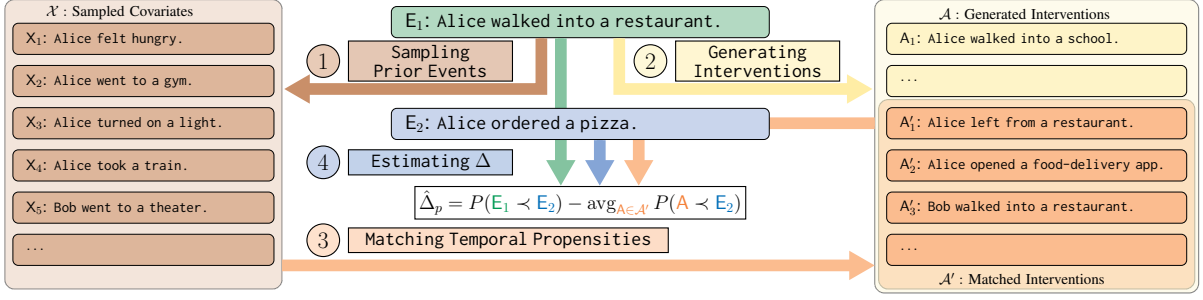


Figure 2: **Illustration of the ROCK framework.** Does E_1 cause E_2 ? To answer this query, ① the event sampler samples a set of covariates $\mathcal{X}$ of events X_k that occur preceding E_1 . ② The intervention generator generates a set $\mathcal{A}$ of interventions A_k on E_1 . ③ A subset $\mathcal{A}' \subset \mathcal{A}$ of interventions is selected whose temporal propensities $q(x; A)$ is close to that of E_1 , $q(x; E_1)$ (Equation (5)). ④ The temporal predictor uses $\mathcal{A}'$ to estimate Δ .

3 The ROCK Framework

Notations. We use sans-serif letters for events, uppercase serif letters *indicators* of whether the corresponding event occurs¹, and lowercase serif letters the realizations of those indicators. For example, in Figures 1 and 2, E_1 : “Alice walked into a restaurant.” $E_1 = \mathbb{1}\{E_1 \text{ occurs}\}$ and $e_{1,i} = \mathbb{1}\{E_1 \text{ occurs to the } i\text{-th study unit}\}$ ². We view the occurrence of events as point processes $E(t)$ on $t \in \{0, 1\}$ (e.g., present vs past). We write $E_1 \succ E_2$ (resp. $\prec$) to mean E_1 occurs succeeding (resp. preceding) E_2 . With a slight abuse of notation, we write $\mathbb{P}(E_1 \prec E_2) = \mathbb{P}(E_1(0), E_2(1))$ and $\mathbb{P}(E_2|E_1) = \mathbb{P}(E_2(1)|E_1(0))$. We write P for estimates of $\mathbb{P}$, and omit measure-theoretic details³.

Overview of the ROCK framework. We set the stage in this section and discuss implementation details in Section 4. Given E_1 and E_2 , as shown in Figure 2, ROCK samples the covariates set $\mathcal{X}$ and interventions set $\mathcal{A}$, from which a matched subset $\mathcal{A}'$ is selected via temporal propensities (Section 3.4). The score Δ is then estimated by Equation (5).

3.1 The Central Question of CCR

Given E_1 and E_2 , as aforementioned in Section 1, we articulate CCR as to estimate, the change of temporal likelihood *had* E_1 been *intervened*:

$$\Delta = \mathbb{P}(E_1 \prec E_2) - \mathbb{P}(\neg E_1 \prec E_2) \quad (2)$$

which assumes values in $[-1, 1]$ and is analogous to the *average treatment effect* in human populations research. Should we know the occurrence of E_1

is purely random, direct estimation is fair enough; however, since these probabilities are eventually estimated from training data, if there are confounding events X_k that always co-occurs with E_1 in the training data, they will bias this estimation. To this end, it is necessary to first clear out several key notions associated with this causal query, and then properly define the intervention $\neg E_1$.

3.2 Units, Covariates, and Interventions

One major challenge of framing a causal query for CCR is the ambiguity of the underlying mechanism. Unlike human populations research, where experiments and study units are obvious to define, it is not immediately clear what they are when faced with semantic meanings of languages. Yet, we can draw parallel again between semantic meanings and human subjects via the following thinking experiment: suppose each human-being keeps a journal/script detailing the complete timeline of her experiences since her conception, then we can treat each individual as a study unit where the temporal relations of events can be inferred from the journal. The treatment will then be the indicator $E_1 = \mathbb{1}\{E_1 \text{ occurs}\}$, the outcome $E_2 = \mathbb{1}\{E_2 \text{ occurs}\}$, and the covariates all events preceding E_1 . Hence a hypothetical observational study can be based on observations of a subpopulation of such journals, from which we form observed treatments $e_{1,1}, e_{1,2}, \dots, e_{1,k}$, the associated outcomes $e_{2,1}, e_{2,2}, \dots, e_{2,k}$, and covariates $x_1, x_2, \dots, x_k$ (as vectors of indicators).

This identification naturally complies with the temporal nature of covariates (Rubin, 2005), since by definition they are *pretreatment* that take place *before* the treatment. We shall now address the issue of intervention (manipulation). Generally speaking, events are complex, and therefore inter-

¹By “occurs,” we mean “is observed.” We treat them interchangeable in the rest of our paper.

²Defined among other concepts in Section 3.2.

³Let $\mathcal{E}$ be the set of commonsense events we consider, the probability space we are working on is $(\mathcal{E} \times \mathcal{E}, \sigma(\mathcal{E} \times \mathcal{E}), \mathbb{P})$.

vention in this case would be better interpreted in a broader sense than one particular type of manipulation such as negation. For example, with E_1 being “Alice walked into a restaurant,” suppose hypothetically, before E_1 , Alice did not walk into a restaurant ($\neg E_1^1$), we can thus compare $\mathbb{P}(E_1 \prec E_2)$ with $\mathbb{P}(\neg E_1^1 \prec E_2)$ to reason to what extent some event E_2 can be viewed as the effect due to E_1 . However, this is not the complete picture: Alice may have walked into somewhere else such as a bar; she may have, instead of walked into, but left a restaurant; instead of Alice, perhaps it was Bob who walked into a restaurant. The temporal information between these events and E_2 are also likely to inform causation between E_1 and E_2 , and they are no less interventions than negation. As such we interpret intervention in our framework in a broader sense, not necessarily only negation or the entailment of negations, but *any events that leads to plausible states of counterfactuality*. We will denote all possible interventions of E_1 as $\mathcal{A}$.

Remark. The well-celebrated *stable unit treatment value assumption* (SUTVA, Rubin, 1980) requires that for each unit there is only one version of the non-treatment. We can view the outcome of our framework in Equation (1) to be the temporal probability threshold at a certain level, $\mathbb{1}\{\mathbb{P}(\cdot \prec E_2) > 1 - \delta\}$, thus complying SUTVA.

3.3 Balancing Covariates

Direct estimation of Δ in Equation (1) is feasible only in an ideal world where those probabilities are estimated from randomized controlled trials (RCTs) such that the treatment (E_1) is assigned completely at random to study units. Due to confounding co-occurrences, events that precede E_1 need to be properly balanced (Mill, 1851; Rosenbaum and Rubin, 1983; Pearl and Mackenzie, 2018). Taking again as an example E_1 : “Alice walked into a restaurant,” and E_2 : “Alice ordered a pizza.” Suppose hypothetically, Alice’s twin sister Alicia, who has the exactly same life experiences up to the point when E_1 took place, but opted not walking into a restaurant, but opened a food delivery app on her phone ($\neg E_1$). Then we can reason that the cause-and-effect relationship from E_1 to E_2 is perhaps not large. On the other hand, if we know another irrelevant person, say Bob, undergone $\neg E_1$ and then E_2 , then perhaps we are not ready to give that conclusion since we do not know if Bob and Alice are comparable at the first place.

This example illustrates the importance of adjusting or balancing pretreatments. Hence we rewrite Equation (1) as conditional expectations among study units that are comparable, i.e.,

$$\mathbb{E}_{\mathbf{x}} [\mathbb{P}(E_1 \prec E_2 | \mathbf{x}) - \mathbb{P}(\neg E_1 \prec E_2 | \mathbf{x})], \quad (3)$$

where $\mathbf{x}$ is the vector of covariates.

Remark. (i) We should define $\mathbf{x}$ as events preceding E_1 , but *not* E_2 , which will potentially introduce posttreatment biases (Rosenbaum, 1984): if an X' that occurs between E_1 and E_2 is adjusted, Δ thus estimated quantifies the effect from E_1 to E_2 *without* passing through X' . (ii) Although $\mathbf{x}$ should be those that are correlated with E_1 , adjusting for un-correlated effects does not introduce biases.

3.4 Matching Temporal Propensities

There are several techniques for balancing covariates such as sub-classification, matched sampling, and covariance adjustment, and via structural equations (Cochran and Chambers, 1965; Pearl, 1995). Rosenbaum and Rubin (1983) showed that such adjustment can be done using any balancing score, with the coarsest one being the propensity $p(\mathbf{x}) = \mathbb{P}(E_1(1) = 1 | \mathbf{x}(0))$, the probability of the occurrence of E_1 (at time 1) conditioning on the covariates being $\mathbf{x}$ (at time 0).

To properly identify what events constitute the covariate set is essential for our CCR framework. In the best scenario, it should include the real cause(s), which is, however, exactly what CCR solves. To circumvent this circular dependency, we use an LM to sample a large number of events preceding E_1 , which should provide a reasonable covariate set. In this case, directly computing $p(\mathbf{x})$ is less feasible; hence we propose to use a surrogate which we call *temporal propensities*:

$$q(\mathbf{x}) = q(\mathbf{x}; E_1) = (\mathbb{P}(E_1(1) = 1 | \mathbf{x}))_{\mathbf{x} \in \mathbf{x}} \quad (4)$$

with each coordinate measures the conditional probability of the event E_1 given an event in $\mathbf{x}$. Therefore, for some fixed threshold ϵ and $p \in \{1, 2\}$, we have the following estimating equation for the L_p -balanced score:

$$\hat{\Delta}_p = f(E_1, E_2) - \frac{1}{|\mathcal{A}'|} \sum_{A \in \mathcal{A}'} f(A, E_2), \quad (5)$$

$$\mathcal{A}' := \left\{ A \in \mathcal{A} : \frac{1}{|\mathbf{x}|} \|q(\mathbf{x}; A) - q(\mathbf{x}; E_1)\|_p < \epsilon \right\}.$$

4 Implementation of ROCK

Having established a framework for CCR, we provide an exemplar implementation of ROCK in this section. Our purpose is not to ace in every CCR dataset, and we expect engineering efforts such as prompt design can bring further improvements.

The core tool we will use is (fine-tuned) pre-trained deep LMs. With the sheer amount of training data (e.g., over 800GB for the Pile dataset, [Gao et al. \(2020\)](#)), it is reasonable to assume that those models would imitate responses of an average reasonable person. On the other hand, it is hard for generation models (masked or open-ended) to parse information that are far from the mask tokens. It is thus more feasible to use LMs to sample summary statistics of the relationships between a pair of events, which is one of the main motivations for using temporal propensities (Equation (4)).

4.1 Components of ROCK

For practical purposes, we represent an event as a 3-tuple $(\text{ARG0}, \text{V}, \text{ARG1})$. ROCK takes two events E_1 and E_2 as inputs, and returns an estimate $\hat{\Delta}$ for Δ according to Equation (5). It contains four components (cf. Figure 2): an event sampler that samples a set $\mathcal{X}$ of events that are likely to occur preceding E_1 ; a temporal predictor whose output $f(X_1, X_2)$ given two input events X_1 and X_2 is an estimate of the temporal probability $\mathbb{P}(X_1 \prec X_2)$; an intervention generator that generates a set $\mathcal{A}$ of events that are considered as interventions of the event E_1 ; and finally a scorer that first forms the temporal propensity vectors $q(x; A) \in \mathbb{R}^{|\mathcal{X}|}$ for each sampled interventions $A \in \mathcal{A}$, then estimates Δ via Equation (5). We next discuss in greater details our implementation of this pipeline.

4.2 Implementation Details

Event Sampling. Given an event E_1 (e.g., E_1 : Alice walked into a restaurant.), we construct the prompt by adding “Before that,” to the sentence, forming “Alice walked into a restaurant. Before that, ” as the final prompt. We use the GPT-J model ([Wang and Komatsuzaki, 2021](#)), which is pretrained on the Pile dataset ([Gao et al., 2020](#)) for open-ended text generation where we set max length of returned sequences to be 30, temperature to be 0.9. We sample $n = 100$ events, cropping at the first stop token of the newly generated sentence to form $\mathcal{X}$.

Temporal Prediction. Given two events E_1 and E_2 , we use mask language modeling to predict their temporal relation by form the prompt $E_1 \langle \text{MASK} \rangle E_2$ and collect the score $s_a(E_1, E_2)$ and $s_b(E_1, E_2)$ for the tokens after and before. We then symmetrize the estimates to form $s(E_1, E_2) = \frac{1}{2}(s_a(E_1, E_2) + s_b(E_2, E_1))$. We can direct use $s(E_1, E_2)$ for $f(E_1, E_2)$; we discuss possible normalizations of this score in Section 5.

Temporality Fine-Tuning. Directly using a pre-trained LM as the temporal predictor is likely to suffer from low coverage, since the tokens before and after usually are not among the top- k most probable tokens. We can increase k but this does not heuristically justify if the outputted scores are meaningful. We thus use the New York Times (NYT) corpus which contains NYT articles from 1987 to 2007 ([Sandhaus, 2008](#)) to fine-tune an LM. Following the same procedure as [Zhou et al. \(2020\)](#), we perform semantic role labeling (SRL) using AllenNLP’s BERT SRL model ([Gardner et al., 2017](#)) to identify sentences with a temporal argument (ARG-TMP) that starts with a temporal connective tmp (either before or after). We then extract the verb and its two arguments (V, ARG0, ARG1) as well as of its temporal argument, thus forming an event pair (E_1, E_2, tmp) . We are able to extract 397, 174 such pairs and construct them into the fine-tuning dataset consisting of “ $E_1 \text{ tmp } E_2$ ” and “ $E_2 \neg \text{tmp } E_1$ ” for each extracted pair, where $\neg \text{tmp}$ is the reverse temporal connective (e.g., after if tmp is before). We then fine-tune a pretrained RoBERTa model (RoBERTa-BASE) using HuggingFace Transformers ([Wolf et al., 2020](#)) via mask language modeling with masking probability $p = 0.1$ for each token. We choose a batch size of 500 and a learning rate of 5×10^{-5} , and train the model to convergence, which was around 135, 000 iterations when the loss converges to 1.37 from 2.02.

Intervention Generator. Given event E_1 , the intervention generator generates a set $\mathcal{A}$ of events that are considered as interventions of the event A in the sense of Section 3.2, which includes manipulating ARG0, V, and ARG1 respectively. We achieve this goal by masking these components individually and filling in the masks using an LM. There are several existing works on generating interventions of sentences ([Feder et al., 2021](#)), and we select PolyJuice ([Wu et al., 2021](#)) in our pipeline due to its robustness. PolyJuice al-

lows conditional generation via control codes such as negation, lexical, resemantic, quantifier, insert, restructure, shuffle, and delete, each corresponds to different flavors on which the sentence is intervened. We drop the fluency-evaluation component of PolyJuice as they will be evaluated by the temporal predictor.

Score Estimation. Given the interventions $\mathcal{A}$ and the sampled covariates $\mathcal{X}$, we can use the temporal predictor to estimate $\mathbb{P}(X \prec A)$ for all $X \in \mathcal{X}$ and $A \in \mathcal{A}$. To obtain temporal propensities $q(x; A)$ for all interventions, we need to estimate $\mathbb{P}(A = 1|X)$ for each X and A . Since by our sampling method, X occurs preceding E_1 , there is an implicit conditioning on E_1 , we may thus set $P(X(0)) = f(X, E_1)$ and $P(X(0), A(1)) = f(X, A)$; we will discuss possible normalizations in Section 5.2. We then form temporal propensity vectors as

$$q(x; A) = \left(\frac{P(X(0))}{P(X(0), A(1))} \right)_{X \in \mathcal{X}}. \quad (6)$$

5 Empirical Studies

We put the ROCK framework into action, our findings reveal that *although temporality is essential for CCR, without balancing covariates, it is prone to spurious correlations.*

5.1 Setup and Details

Evaluation Datasets. We evaluate the ROCK framework on the Choice of Plausible Alternatives dataset (COPA, Gordon et al., 2012) and a self-constructed dataset of 153 instances using the first dimension (cause-and-effect) of GLUCOSE (GLUCOSE-D1, Mostafazadeh et al., 2020). Each instance in COPA consists of a premise, two plausible choices, and a question type asking which choice is the choice (or effect) of the premise. When asking for cause, we set the premise as E_1 , and two choices as E_2 respectively; otherwise we take the premise as E_2 and two choices as E_1 respectively. We choose the choice with the higher score. We evaluate the development set of 100 instances (COPA-DEV) and the test set of 500 instances (COPA-TEST). To construct GLUCOSE-D1, we take the test set and set the cause as premise, the effect and another candidate event as two choices then follow the same procedure.

Baseline Scores and Variants. To test the validity and the effectiveness of ROCK, we compare

with the adjusted score $\hat{\Delta}_p$ with several other reasonable scores that may be intuitive at first sight.

- L_1 -balanced score $\hat{\Delta}_1$: set $p = 1$ in (5).
- L_2 -balanced score $\hat{\Delta}_2$: set $p = 2$ in (5).
- Vanilla temporal score $\hat{\Delta}_{E_1} = \mathbb{P}(E_1 \prec E_2)$.
- Unadjusted score $\hat{\Delta}_{\mathcal{A}}$: set $\mathcal{A}' = \mathcal{A}$ in (5).
- Misspecified score $\hat{\Delta}_{\mathcal{X}}$: set $\mathcal{A}' = \mathcal{X}$ in (5).

Here the L_p -balanced scores are those balanced using temporal propensities with L_p norm in Equation (5); the vanilla temporal score is perhaps the most straightforward one, which treats temporal precedence as causation; the unadjusted score is obtained without balancing the covariates; the misspecified score mistakes the covariates for interventions. All these three have intuitive explanations but are either insufficient for CCR or prone to spurious correlations. Note that $\lim_{\epsilon \downarrow 0} \hat{\Delta}_p = \hat{\Delta}_{E_1}$ (when nothing is kept) and $\lim_{\epsilon \uparrow 1} \hat{\Delta}_p = \hat{\Delta}_{\mathcal{A}}$ (when everything is kept).

5.2 Design Choices and Normalizations

We discuss several design choices and normalizations that might stabilize estimation procedures. As shown in Section 5.4, although some of them may benefit certain datasets, the improvements are *marginal* compared with what temporal propensity matching brings.

Direct Matching (D). In (6), we may directly match the vectors of probabilities $(f(A, X))_{X \in \mathcal{X}}$.

Temporality Pre-Filtering (F). As the covariate sampler and temporal predictor are two different LMs, a sampled covariate might not be a preceding event judged by the temporal predictor. We may filter the covariates before matching temporal propensities such that $f(X, E_1) > f(X, E_1)$.

Score Normalization (S). In Section 4 we use $s(E_1, E_2)$ for $f(E_1, E_2)$, we can also normalize it and form $f(E_1, E_2)$ through

$$f(E_1, E_2) = \frac{s(E_1, E_2)}{s(E_1, E_2) + s(E_2, E_1) + s(E_1, N) + s(N, E_1)} \quad (7)$$

where N represents the null event when no additional information is given, set as an empty string.

Propensity Normalization (Q). In Equation (6), we can also normalize the estimates first before forming the q vectors via $P(X(0)) = f(X, E_1) / \sum_{X' \in \mathcal{X}} f(X', E_1)$ and $P(X(0), A(1)) = f(X, A) / \sum_{X' \in \mathcal{X}} f(X', A)$.

	Random Baseline	$\hat{\Delta}_1 \uparrow$ L_1 -Balanced	$\hat{\Delta}_2 \uparrow$ L_2 -Balanced	$\hat{\Delta}_{E_1} \uparrow$ Temporal	$\hat{\Delta}_A \uparrow$ Unbalanced	$\hat{\Delta}_X \uparrow$ Misspecified
COPA-DEV	0.5 ± 0.050	0.6900	0.7000	0.5800	0.5600	0.5300
COPA-TEST	0.5 ± 0.022	0.5640	0.5640	0.5200	0.5400	0.5240
GLUCOSE-D1	0.5 ± 0.040	0.6645	0.6968	0.5677	0.5742	0.6581
COPA-DEV (-T)	0.5 ± 0.050	0.6200	0.6300	0.5300	0.4800	0.5300
COPA-TEST (-T)	0.5 ± 0.022	0.5800	0.5740	0.4540	0.4600	0.4860
GLUCOSE-D1 (-T)	0.5 ± 0.040	0.6065	0.6194	0.5548	0.4387	0.3742

Table 1: **Best zero-shot results.** Shaded rows have temporal fine-tuning (T) disabled. (i) Estimators with temporal propensities balanced ($\hat{\Delta}_1$ and $\hat{\Delta}_2$) perform consistently better than the unbalanced and the temporal estimators. (ii) (2) In general, without temporality fine-tuning (“-T”, see Section 4), the performances degrade.

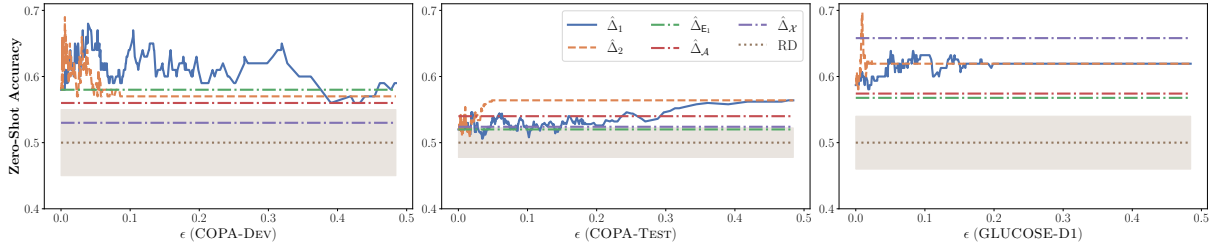


Figure 3: **Best zero-shot result vs ϵ .** Balanced estimators significantly outperform un-balanced and other variants for both COPA-DEV (left), COPA-TEST (middle) and GLUCOSE-D1 (right).

Co-occurrence Stabilization (C). The fine-tuned temporal predictor may sometimes still fail to cover the connectives. We can stabilize $\mathbb{P}(X \prec A)$ by setting it to $(P(A(0), X(1)) + P(X(0), A(1)))/2$.

Estimand Normalization (E). We can normalize the probability $\mathbb{P}(A \prec B)$ in the estimand Δ by dividing $(P(A(0), B(1)) + P(B(0), A(1)))$.

5.3 Results

5.3.1 A Concrete Example

We first examine a particular example when the vanilla temporal score $\hat{\Delta}_{E_1}$ fails but $\hat{\Delta}_1$ does not.

Example 1 (Did $E_1^{(1)}$ or $E_1^{(2)}$ cause E_2 ?).

$E_1^{(1)}$: I was preparing to wash my hands.

$E_1^{(2)}$: I was preparing to clean the bathroom.

E_2 : I put rubber gloves on.

$A_{15}^{(1)}$: I was preparing to wash my feet.

$A_5^{(2)}$: Kevin was preparing to clean the bathroom.

This is the 63-rd instance in COPA-DEV together a matched intervention (L_2 -balancing with optimal ϵ) for each choice. The unadjusted scores are $\hat{\Delta}_A(E_1^{(1)}, E_2) = 0.067894$ and $\hat{\Delta}_A(E_1^{(2)}, E_2) = 0.097539$ while the L_2 -balanced scores are $\hat{\Delta}_2(E_1^{(1)}, E_2) = -0.010274$ and $\hat{\Delta}_2(E_1^{(2)}, E_2) = 0.001311$. The balanced score

selects the correct choice ($E_1^{(2)}$) with higher confidence. More details are in given the Appendix.

5.3.2 Discussion

We show best zero-shot results over design choices (and over ϵ) in Figure 3 and Table 1. As ROCK tackles CCR from a completely new perspective, there are no real baselines to compare with; our goal is to demonstrate that *the causal inference motivated method, temporal propensity matching, mitigates spurious correlations* by comparing balanced scores with unbalanced ones. We think this perspective would also benefit the NLP community at large for solving CCR and other related tasks.

Comparison with existing methods. Among unsupervised baselines that we are able to reproduce, the self-talk method (Shwartz et al., 2020) achieves 66% on COPA-DEV without external knowledge and 69% when the CoMET-Net (Bosselut et al., 2019) that contains commonsense knowledge is used. Our L_2 -balanced score achieves 70% without using any external commonsense knowledge.

Temporal propensity matching is effective. In Table 1 (unshaded rows), we observe that balanced scores have generally better performances on all datasets compared with the temporal estimator and the unadjusted estimator, implying that (i) temporality is important for CCR, yet they are susceptible

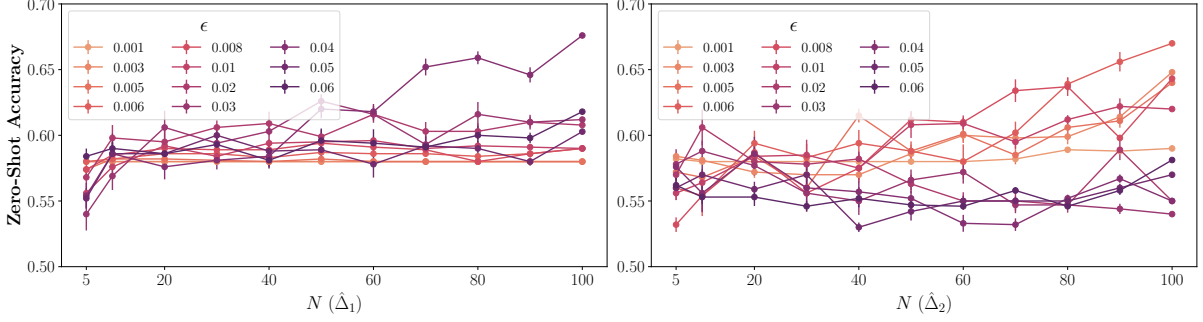


Figure 4: **Zero-shot result on COPA-DEV vs covariate set size $N = |\mathcal{X}|$ with 95%-confidence bands.** In general, using a larger N improves performances for both L_1 -balanced score ($\hat{\Delta}_1$, left) and L_2 -balanced score ($\hat{\Delta}_2$, right).

to spurious correlations; (ii) balancing covariates via matching temporal propensities is effective.

Rules-of-thumb for choosing ϵ . The parameter ϵ controls the threshold of covariates selection and p controls its geometry (see e.g., [Hastie et al., 2015](#)). Hinted by Figure 3, a general rule-of-thumb should be $\epsilon < 0.1$. Table A.1 shows optimal ϵ values when constrained to $[0, 0.1]$, where all are global optimal except for COPA-TEST under L_1 -balanced score (whose accuracy is 0.552). Hence we recommend setting ϵ to be reasonably small ϵ such as within $(0.01, 0.1)$ when $p = 1$ and relatively smaller such as $(0.005, 0.05)$ when $p = 2$. The optimal value depends on the implementation details of ROCK components and domains of CCR to be performed, yet these choices should result in a good start.

5.4 Ablation Studies

Temporality Fine-Tuning. Shaded rows in Table 1 show that when we use the pretrained RoBERTa-BASE without temporality fine-tuning (we increase k to 30), almost all estimators do not have decent performance. We conclude that (i) pretrained LMs usually have poor “temporal awareness,” and (ii) temporal fine-tuning helps LMs to extract temporal knowledge essential to CCR.

Covariate Set Size. Figure 4 depicts zero-shot results on COPA-TEST against the covariate set size $N = |\mathcal{X}|$ together with 95%-confidence bands. Here we only enable score normalizations (N) among all six normalizations. We observe that in general, increasing covariate set size improves performances if ϵ is reasonable: if ϵ is too small, added covariates may have little impacts while they may introduce more noises if ϵ is too large.

Normalizations. In Section 5.2 we discussed six possible normalizations. We report the best per-

	COPA-DEV		COPA-TEST		GLUCOSE-D1	
	$\hat{\Delta}_1 \uparrow$	$\hat{\Delta}_2 \uparrow$	$\hat{\Delta}_1 \uparrow$	$\hat{\Delta}_2 \uparrow$	$\hat{\Delta}_1 \uparrow$	$\hat{\Delta}_2 \uparrow$
Best	0.6900	0.7000	0.5640	0.5640	0.6645	0.6968
-S	0.01	0.06	-	-	0.08	0.11
-Q	0.01	-	-	-	0.03	-
-C	-	-	0.01	0.01	0.09	0.13
-E	0.01	0.01	-	-	0.03	-

Table 2: **Single-component ablations on normalizations.** Marked in red are percentage decreases compared with the best result (i.e., computed as $(a - b)/a$).

formance when each normalization is removed in Table 2, where red marks the percentage decrease compared with the best result (**D** and **F** not shown as there is no change). Full ablations of all combinations of normalizations and more discussions are given in the Appendix. We observe that (i) certain normalizations benefit certain datasets; (ii) in general, improvements due to normalizations are only *marginal*, so long as the framework is appropriate.

6 Discussions and Open Problems

We articulate the central question of CCR and introduce ROCK, a novel framework for zero-shot CCR based on classical causal inference principles. ROCK sheds light on the CCR problem from new perspectives that are arguably more well-founded and demonstrates great potential for zero-shot CCR as shown by empirical studies of various datasets. There are several possible avenues for future works. (i) **Prompt engineering** for better temporal predictors and event sampler will likely benefit ROCK; (ii) **Implicit events and reporting biases** in training data are likely to bias the LMs. How to account for implicit events? (iii) **Contextual CCR** takes the *contexts* into consideration. How to incorporate contextual information for ROCK?

References

- Debarun Bhattacharjya, Tian Gao, and Dharmashankar Subramanian. 2020. [Order-dependent event models for agent interactions](#). In *Proceedings of the Twenty-Ninth International Joint Conference on Artificial Intelligence, IJCAI-20*, pages 1977–1983. International Joint Conferences on Artificial Intelligence Organization.
- Antoine Bosselut, Hannah Rashkin, Maarten Sap, Chaitanya Malaviya, Asli Celikyilmaz, and Yejin Choi. 2019. COMET: Commonsense Transformers for Automatic Knowledge Graph Construction. In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics (ACL)*.
- Mario Bunge. 1979. *Causality and modern science*, 4 edition. Routledge.
- Du-Seong Chang and Key-Sun Choi. 2004. Causal relation extraction using cue phrase and lexical pair probabilities. In *IJCNLP*.
- William G Cochran and S Paul Chambers. 1965. The planning of observational studies of human populations. *Journal of the Royal Statistical Society. Series A (General)*, 128(2):234–266.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. 2019. Bert: Pre-training of deep bidirectional transformers for language understanding. In *Proc. of the Annual Conference of the North American Chapter of the Association for Computational Linguistics*.
- Quang Xuan Do, Yee Seng Chan, and Dan Roth. 2011. Minimally supervised event causality identification. In *Proceedings of the Conference on EMNLP*. Association for Computational Linguistics.
- Amir Feder, Katherine A. Keith, Emaad Manzoor, Reid Pryzant, Dhanya Sridhar, Zach Wood-Doughty, Jacob Eisenstein, Justin Grimmer, Roi Reichart, Margaret E. Roberts, Brandon M. Stewart, Victor Veitch, and Diyi Yang. 2021. [Causal inference in natural language processing: Estimation, prediction, interpretation and beyond](#).
- Ronald A Fisher. 1958. Cancer and smoking. *Nature*, 182(4635):596–596.
- Leo Gao, Stella Biderman, Sid Black, Laurence Golding, Travis Hoppe, Charles Foster, Jason Phang, Horace He, Anish Thite, Noa Nabeshima, Shawn Presser, and Connor Leahy. 2020. The Pile: An 800gb dataset of diverse text for language modeling. *arXiv preprint arXiv:2101.00027*.
- Matt Gardner, Joel Grus, Mark Neumann, Oyvind Tafjord, Pradeep Dasigi, Nelson F. Liu, Matthew Peters, Michael Schmitz, and Luke S. Zettlemoyer. 2017. [Allennlp: A deep semantic natural language processing platform](#).
- Tobias Gerstenberg, Noah D. Goodman, David A. Lagnado, and Joshua B. Tenenbaum. 2021. A counterfactual simulation model of causal judgments for physical events. *Psychological review*.
- Andrew Gordon, Zornitsa Kozareva, and Melissa Roemmele. 2012. [SemEval-2012 task 7: Choice of plausible alternatives: An evaluation of commonsense causal reasoning](#). In **SEM 2012: The First Joint Conference on Lexical and Computational Semantics – Volume 1: Proceedings of the main conference and the shared task, and Volume 2: Proceedings of the Sixth International Workshop on Semantic Evaluation (SemEval 2012)*, pages 394–398, Montréal, Canada. Association for Computational Linguistics.
- C. W. J. Granger. 1969. [Investigating causal relations by econometric models and cross-spectral methods](#). *Econometrica*, 37(3):424–438.
- Mingyue Han and Yinglin Wang. 2021. [Doing good or doing right? exploring the weakness of commonsense causal reasoning models](#). In *Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing (Volume 2: Short Papers)*, pages 151–157, Online. Association for Computational Linguistics.
- Trevor Hastie, Robert Tibshirani, and Martin Wainwright. 2015. *Statistical Learning with Sparsity: The Lasso and Generalizations*. Chapman Hall.
- James J Heckman. 2005. Rejoinder: response to sobel. *Sociological Methodology*, 35(1):135–150.
- Austin Bradford Sir Hill. 1965. The environment and disease: Association or causation? *Journal of the Royal Society of Medicine*, 58:295 – 300.
- Paul W Holland. 1986. Statistics and causal inference. *Journal of the American statistical Association*, 81(396):945–960.
- Guido W Imbens and Donald B Rubin. 2015. *Causal inference in statistics, social, and biomedical sciences*. Cambridge University Press.
- Dongyeop Kang, Varun Gangal, Ang Lu, Zheng Chen, and Eduard Hovy. 2017. Detecting and explaining causes from text for a time series event. In *Conference on Empirical Methods on Natural Language Processing*.
- Pride Kavumba, Naoya Inoue, Benjamin Heinzerling, Keshav Singh, Paul Reisert, and Kentaro Inui. 2019. [When choosing plausible alternatives, clever hans can be clever](#). In *Proceedings of the First Workshop on Commonsense Inference in Natural Language Processing*, pages 33–42, Hong Kong, China. Association for Computational Linguistics.
- Katherine A Keith, David Jensen, and Brendan O’Connor. 2020. Text and causal inference: A review of using text to remove confounding from causal estimates. *arXiv preprint arXiv:2005.00649*.
- Benjamin Kuipers. 1984. Commonsense reasoning about causality: deriving behavior from structure. *Artificial intelligence*, 24(1-3):169–203.
- Zhiyi Luo, Yuchen Sha, Kenny Q Zhu, Seung-won Hwang, and Zhongyuan Wang. 2016. Commonsense causal reasoning between short texts. In *KR*, pages 421–431.
- John Stuart Mill. 1851. *A System of Logic, Ratiocinative and Inductive: Being a Connected View of the Principles of Evidence, and the Methods of Scientific Investigation*, volume 1 of *Cambridge Library Collection - Philosophy*. Cambridge University Press.
- Nasrin Mostafazadeh, Aditya Kalyanpur, Lori Moon, David Buchanan, Lauren Berkowitz, Or Biran, and Jennifer Chu-Carroll. 2020. [Glucose: Generalized and contextualized story explanations](#).

743	Jerzy S Neyman. 1923. On the application of probability	798
744	theory to agricultural experiments. Essay on principles.	799
745	Section 9. <i>Annals of Agricultural Sciences</i> , 10:1–51.	800
746	Qiang Ning, Zhili Feng, and Dan Roth. 2017. A Structured	801
747	Learning Approach to Temporal Relation Extraction . In	802
748	<i>Proc. of the Conference on Empirical Methods in Natural</i>	
749	<i>Language Processing (EMNLP)</i> , pages 1038–1048, Copen-	803
750	hagen, Denmark. Association for Computational Linguis-	804
751	tics.	805
752	Qiang Ning, Zhili Feng, Hao Wu, and Dan Roth. 2019a. Joint	806
753	reasoning for temporal and causal relations. <i>arXiv preprint</i>	807
754	<i>arXiv:1906.04941</i> .	
755	Qiang Ning, Sanjay Subramanian, and Dan Roth. 2019b. An	
756	Improved Neural Baseline for Temporal Relation Extrac-	
757	tion . In <i>Proc. of the Conference on Empirical Methods in</i>	810
758	<i>Natural Language Processing (EMNLP)</i> .	811
759	Qiang Ning, Hao Wu, Haoruo Peng, and Dan Roth. 2018.	812
760	Improving temporal relation extraction with a globally ac-	
761	quired statistical resource . In <i>Proceedings of the 2018</i>	
762	<i>Conference of the North American Chapter of the Asso-</i>	813
763	<i>ciation for Computational Linguistics: Human Language</i>	814
764	<i>Technologies, Volume 1 (Long Papers)</i> , pages 841–851,	815
765	New Orleans, Louisiana. Association for Computational	816
766	Linguistics.	817
767	Timothy J. O’Gorman, Kristin Wright-Bettner, and Martha	
768	Palmer. 2016. Richer event description: Integrating event	
769	coreference with temporal, causal and bridging annotation.	818
770	Judea Pearl. 1995. Causal diagrams for empirical research .	819
771	<i>Biometrika</i> , 82(4):669–688.	820
772	Judea Pearl and Dana Mackenzie. 2018. <i>The book of why: the</i>	821
773	<i>new science of cause and effect</i> . Basic Books.	822
774	Jonas Peters, Dominik Janzing, and Bernhard Schölkopf. 2017.	823
775	<i>Elements of Causal Inference: Foundations and Learning</i>	
776	<i>Algorithms</i> . The MIT Press.	824
777	Alec Radford, Jeff Wu, Rewon Child, David Luan, Dario	825
778	Amodei, and Ilya Sutskever. 2019. Language models are	826
779	unsupervised multitask learners.	
780	Hannah Rashkin, Maarten Sap, Emily Allaway, Noah A Smith,	827
781	and Yejin Choi. 2018. Event2mind: Commonsense in-	828
782	ference on events, intents, and reactions. <i>arXiv preprint</i>	829
783	<i>arXiv:1805.06939</i> .	830
784	Paul R Rosenbaum. 1984. The consequences of adjustment	831
785	for a concomitant variable that has been affected by the	832
786	treatment. <i>Journal of the Royal Statistical Society: Series</i>	833
787	<i>A (General)</i> , 147(5):656–666.	834
788	Paul R Rosenbaum. 1989. Optimal matching for observational	835
789	studies. <i>Journal of the American Statistical Association</i> ,	
790	84(408):1024–1032.	836
791	Paul R Rosenbaum. 2002. <i>Observational Studies</i> . Springer.	837
792	Paul R Rosenbaum and Donald B Rubin. 1983. The central	838
793	role of the propensity score in observational studies for	
794	causal effects. <i>Biometrika</i> , 70(1):41–55.	839
795	Dan Roth. 2017. Incidental Supervision: Moving beyond Su-	840
796	pervised Learning . In <i>Proc. of the Conference on Artificial</i>	841
797	<i>Intelligence (AAAI)</i> .	842
	Donald B Rubin. 1974. Estimating causal effects of treatments	843
	in randomized and nonrandomized studies. <i>Journal of</i>	844
	<i>Educational Psychology</i> , 66(5):688.	845
	Donald B Rubin. 1980. Bias reduction using mahalanobis-	846
	metric matching. <i>Biometrics</i> , pages 293–298.	847
	Donald B Rubin. 2005. Causal inference using potential out-	848
	comes: Design, modeling, decisions. <i>Journal of the Ameri-</i>	849
	<i>can Statistical Association</i> , 100(469):322–331.	850
	Bertrand Russell. 1912. On the notion of cause . <i>Proceedings</i>	851
	<i>of the Aristotelian Society</i> , 13:1–26.	852
	E. Sandhaus. 2008. The New York Times Annotated Corpus.	853
	<i>Linguistic Data Consortium, Philadelphia</i> .	854
	Maarten Sap, Hannah Rashkin, Derek Chen, Ronan Le Bras,	855
	and Yejin Choi. 2019. Social iqa: Commonsense reasoning	856
	about social interactions. In <i>EMNLP 2019</i> .	857
	Maarten Sap, Vered Shwartz, Antoine Bosselut, Yejin Choi,	
	and Dan Roth. 2020. Commonsense reasoning for natural	
	language processing. In <i>Proceedings of the 58th Annual</i>	
	<i>Meeting of the Association for Computational Linguistics:</i>	
	<i>Tutorial Abstracts</i> , pages 27–33.	
	Vered Shwartz, Peter West, Ronan Le Bras, Chandra Bhaga-	
	vatula, and Yejin Choi. 2020. Unsupervised commonsense	
	question answering with self-talk . In <i>Proceedings of the</i>	
	<i>2020 Conference on Empirical Methods in Natural Lan-</i>	
	<i>guage Processing (EMNLP)</i> , pages 4615–4629. Associa-	
	tion for Computational Linguistics.	
	Ieva Staliunaite, Philip John Gorinski, and Ignacio Iacobacci.	
	2021. Improving commonsense causal reasoning by adver-	
	sarial training and data augmentation. In <i>AAAI</i> .	
	Alexandre Tamborrino, Nicola Pellicanò, Baptiste Pannier,	
	Pascal Voitot, and Louise Naudin. 2020. Pre-training is	
	(almost) all you need: An application to commonsense	
	reasoning. <i>ArXiv</i> , abs/2004.14074.	
	Siddharth Vashishtha, Adam Poliak, Yash Kumar Lal, Ben-	
	jamin Van Durme, and Aaron Steven White. 2020. Tempo-	
	ral reasoning in natural language inference. In <i>Proceedings</i>	
	<i>of the 2020 Conference on Empirical Methods in Natural</i>	
	<i>Language Processing: Findings</i> , pages 4070–4078.	
	Ben Wang and Aran Komatsuzaki. 2021. GPT-J-6B: A 6	
	Billion Parameter Autoregressive Language Model. https://github.com/kingoflolz/mesh-transformer-jax .	
	Thomas Wolf, Lysandre Debut, Victor Sanh, Julien Chau-	
	mond, Clement Delangue, Anthony Moi, Pierric Cistac,	
	Tim Rault, Rémi Louf, Morgan Funtowicz, Joe Davison,	
	Sam Shleifer, Patrick von Platen, Clara Ma, Yacine Jer-	
	nite, Julien Plu, Canwen Xu, Teven Le Scao, Sylvain Gug-	
	ger, Mariama Drame, Quentin Lhoest, and Alexander M.	
	Rush. 2020. Transformers: State-of-the-art natural lan-	
	guage processing . In <i>Proceedings of the 2020 Conference</i>	
	<i>on Empirical Methods in Natural Language Processing:</i>	
	<i>System Demonstrations</i> , pages 38–45, Online. Associa-	
	tion for Computational Linguistics.	
	Tongshuang Wu, Marco Tulio Ribeiro, Jeffrey Heer, and	
	Daniel Weld. 2021. Polyjuice: Generating counterfac-	
	tuals for explaining, evaluating, and improving models . In	
	<i>Proceedings of the 59th Annual Meeting of the Association</i>	
	<i>for Computational Linguistics and the 11th International</i>	
	<i>Joint Conference on Natural Language Processing (Volume</i>	
	<i>1: Long Papers)</i> , pages 6707–6723, Online. Association	
	for Computational Linguistics.	

- 858 Hongming Zhang, Yintong Huo, Xinran Zhao, Yangqiu Song,
859 and Dan Roth. 2021. Learning contextual causality be-
860 tween daily events from time-consecutive images. In *Pro-
861 ceedings of the IEEE/CVF Conference on Computer Vision
862 and Pattern Recognition (CVPR) Workshops*, pages 1752–
863 1755.
- 864 Ben Zhou, Qiang Ning, Daniel Khashabi, and Dan Roth. 2020.
865 Temporal common sense acquisition with minimal supervi-
866 sion. *arXiv preprint arXiv:2005.04304*.

A Additional Experiment Details

A.1 Rule-of-Thumb for Choosing ϵ

In Table A.1 we show the best ϵ values when constrained in $\epsilon \in [0, 0.1]$. Hence we recommend setting ϵ to be reasonably small ϵ such as within (0.01, 0.1) when $p = 1$ and relatively smaller such as (0.005, 0.05) when $p = 2$. The optimal value depends on the implementation details of ROCK components and domains of CCR to be performed, yet these choices should result in a good start.

A.2 Further Discussions on Temporality Fine-Tuning

In Figure 3, we observe that, counterintuitively, without temporality fine-tuning, the best performances of balanced estimators (0.58) are higher than those with temporality fine-tuning (0.564). Although this gap is within one standard deviation of the random baseline (0.022) thus no statistically significant conclusions can be drawn, but it might hint that pretrained LMs may have already been very aware of temporality. Is this really the case? A closer look at the full ablation table to be introduced shortly in Table A.5 reveals that the stellar performance is attributed to one particular normalization, estimand normalization (E), which was actually detrimental to another dataset (GLUCOSE-D1). Hence we think this normalization may favor certain dataset over others, thus we think it is not recommendable to include this normalization when dealing with a new dataset.

A.3 Full Ablation on Normalizations

Recall in Section 5.4 we discussed six possible normalizations that may stabilize the estimation procedure:

- (D) **Direct Matching:** in (6), instead of forming the temporal propensity vectors $\mathbf{q}$ using conditional probabilities, we may directly match the vectors of probabilities $(f(A, X))_{X \in \mathcal{X}}$. This normalization is not well motivated but might be easier to compute under certain circumstances, hence we include it as a comparison.
- (F) **Temporality Pre-Filtering:** as the covariate sampler and temporal predictor are two different LMs, a sampled covariate might not be a preceding event judged by the temporal predictor. Thus, we can filter the covariates $\mathcal{X}$ before matching temporal propensities such

that we only keep covariates $X \in \mathcal{X}$ satisfying $f(X, E_1) > f(S, E_1)$.

- (S) **Score Normalization:** in Section 4 we use $s(E_1, E_2)$ for $f(E_1, E_2)$. We can also normalize it and form $f(E_1, E_2)$ through

$$f(E_1, E_2) = \frac{s(E_1, E_2)}{s(E_1, E_2) + s(E_2, E_1) + s(E_1, N) + s(N, E_1)} \quad (\text{A.1})$$

where N represents the null event when no additional information is given, set as an empty string. In practice, this normalization does not differ much from the normalization

$$f(E_1, E_2) = \frac{s(E_1, E_2)}{s(E_1, E_2) + s(E_2, E_1)}, \quad (\text{A.2})$$

which does not involve N . However, using N has the benefit of stabilizing the estimate $f(\cdot, \cdot)$ as in rare scenarios $s(E_1, E_2)$ and $s(E_2, E_1)$ may both close to zero.

- (Q) **Propensity Normalization:** in Equation (6), we can also normalize the estimates first before forming the q vectors via

$$\begin{aligned} P(X(0)) &= \frac{f(X, E_1)}{\sum_{X' \in \mathcal{X}} f(X', E_1)}, \\ P(X(0), A(1)) &= \frac{f(X, A)}{\sum_{X' \in \mathcal{X}} f(X', A)}, \end{aligned} \quad (\text{A.3})$$

where we estimate $P(X(0))$ as the relative frequency of $X(0)$ among all possible events in $\mathcal{X}$; and $P(X(0), A(1))$ among all possible (X, A) pairs.

- (C) **Co-Occurrence Stabilization:** on rare occasions, the fine-tuned temporal predictor may sometimes still fail to cover the connectives. We can stabilize $\mathbb{P}(X \prec A)$ by setting it to $(P(A(0), X(1)) + P(X(0), A(1)))/2$. This in effect results in an alternative estimand based on co-occurrences of events (instead of precedence) and can be viewed as a weaker causation in CCR.
- (E) **Estimand Normalization:** the score normalization (N) takes place at temporal propensity matching. We can normalize the temporal probability $\mathbb{P}(A \prec B)$ in the estimand Δ by dividing $(P(A(0), B(1)) + P(B(0), A(1)))$, thus setting

$$\mathbb{P}(A \prec B) = \frac{P(A(0), B(1))}{P(A(0), B(1)) + P(B(0), A(1))}. \quad (\text{A.4})$$

	COPA-DEV		COPA-TEST		GLUCOSE-D1	
	$\hat{\Delta}_1$	$\hat{\Delta}_2$	$\hat{\Delta}_1$	$\hat{\Delta}_2$	$\hat{\Delta}_1$	$\hat{\Delta}_2$
ϵ^*	0.043067	0.006029	0.059232	0.048837	0.046643	0.009374

Table A.1: Best choices of ϵ when $\epsilon < 0.1$.

	$\hat{\Delta}_1$	$\hat{\Delta}_2$	$\hat{\Delta}_{E_1}$	$\hat{\Delta}_A$	$\hat{\Delta}_X$
$(E_1, E_2^{(1)})$	-0.002	-0.002	0.106	0.002	0.106
$(E_1, E_2^{(2)})$	-0.001	-0.001	0.086	-0.012	0.086

Table A.2: Scores for Example A.1.

A.3.1 Ablation Results

We report ablations on all possible subset of normalizations together with temporality fine-tuning (-T, see Section 4 in Table A.5. Note that when **D** is enabled, **S** and **Q** are not active and when **C** is enabled, **E** is not active, thus resulting in a total of $2^2(2^2 + 1)(2^1 + 1) = 30$ combinations

Ablations resulting in the best performances are highlighted in blue and those resulting in the worst the performances are highlighted in red. Shaded rows are results without temporal fine-tuning (using top $k = 30$ tokens in mask language modeling). We summarize our observations as follows.

Improvements due to normalizations are marginal. The gap between best and worst performance are marginal, except for the GLUCOSE-D1 dataset, which is mainly caused by enabling estimand normalization (**E**). Without considering **E**, the worst result is 0.594 (**+Q** or **+FQ**). Furthermore, we note the gap between the best results and the results under no normalizations ($\emptyset$) is also marginal, indicating that for CCR it is more important to have a well-established baseline and temporal signal extractors than exploring different normalizations.

Furthermore, the outliers are interesting: enabling estimand normalization (**E**) has little or no effects on most datasets but can boost the performance on COPA-TEST under non fine-tuned temporal predictors (-T) while is detrimental to GLUCOSE-D1 under fine-tuned temporal predictors.

Rules-of-thumb for choosing normalizations. As a general rule-of-thumb, temporal score nor-

	$\hat{\Delta}_1$	$\hat{\Delta}_2$	$\hat{\Delta}_{E_1}$	$\hat{\Delta}_A$	$\hat{\Delta}_X$
$(E_1^{(1)}, E_2)$	-0.010	-0.010	0.068	0.036	0.068
$(E_1^{(2)}, E_2)$	0.002	0.001	0.098	0.035	0.098

Table A.3: Scores for Example A.2.

	$\hat{\Delta}_1$	$\hat{\Delta}_2$	$\hat{\Delta}_{E_1}$	$\hat{\Delta}_A$	$\hat{\Delta}_X$
$(E_1^{(1)}, E_2)$	0.056	-0.001	0.109	0.096	0.109
$(E_1^{(2)}, E_2)$	0.005	-0.010	0.279	0.118	0.279

Table A.4: Scores for Example A.3.

malization (**S**) should be enabled and the q vectors should be properly formed (without direct matching **D**); temporal pre-filtering (**F**) and propensity normalization (**Q**) in general do not affect the results significantly; co-occurrence stabilization (**C**) has greater positive effect on datasets when a weaker notion of causation are desirable (e.g., GLUCOSE-D1 we constructed); while estimand normalization (**E**) improves certain datasets (e.g., COPA-TEST without temporal fine-tuning), it has detrimental effects on some others (e.g., GLUCOSE-D1 with temporal fine-tuning), hence we recommend disabling it by default.

A.4 Full Examples

We also attach three full examples from our implementation of the ROCK. The problem instances are given below. For each instance, we tabulate 50 covariates sampled, all interventions generated, the corresponding $\|q(x; A) - q(x; E_1)\|_p$, and the temporal probabilities $\mathbb{P}(\cdot \prec E_2)$.

Example A.1 (Did E_1 cause $E_2^{(1)}$ or $E_2^{(2)}$?).

E_1 : The teacher assigned homework to the students.
 $E_2^{(1)}$: The students passed notes.
 $E_2^{(2)}$: The students groaned.

This is the 72-nd instance of COPA-DEV, the full tables for inferring the causation from E_1 to $E_2^{(1)}$ and E_1 to $E_2^{(2)}$ are given in Table A.6 and Table A.7 respectively. Different scores are shown in Table A.2. Note that this example is not easy:

Example A.2 (Did $E_1^{(1)}$ or $E_1^{(2)}$ cause E_2 ?).

$E_1^{(1)}$: I was preparing to wash my hands.
 $E_1^{(2)}$: I was preparing to clean the bathroom.
 E_2 : I put rubber gloves on.

This is the 63-rd instance of COPA-DEV, the full tables for inferring the causation from $E_1^{(1)}$ to E_2 and $E_1^{(2)}$ to E_2 are given in Table A.8 and

Table A.9 respectively. Different scores are shown in Table A.3.

Example A.3 (Did $E_1^{(1)}$ or $E_1^{(2)}$ cause E_2 ?).

$E_1^{(1)}$: His pocket was filled with coins.

$E_1^{(2)}$: He sewed the hole in his pocket.

E_2 : The man's pocket jingled as he walked.

This is the 79-th instance of COPA-DEV, the full tables for inferring the causation from $E_1^{(1)}$ to E_2 and $E_1^{(1)}$ to $E_2^{(2)}$ are given in Table A.10 and Table A.11 respectively. Different scores are shown in Table A.2.

Dataset	Score	Best	Worst	θ	+D	+F	+S	+Q	+C	+E	+DF	+DC	+DE	+FS	+FQ	+FC	+FE	+SQ	+SC	+SE	+QC	+QE	+DFC	+DFE	+FSQ	+FSQC	+FSQE	+FQC	+FQE	+SQC	+SQE	+FSQC	+FSQE
COPA-Dev	$\Delta_1 \uparrow$	0.690	0.620	0.670	0.660	0.670	0.680	0.650	0.650	0.680	0.660	0.650	0.680	0.680	0.650	0.650	0.680	0.670	0.620	0.680	0.660	0.660	0.650	0.680	0.670	0.660	0.690	0.660	0.660	0.660	0.690	0.660	0.690
	$\Delta_2 \uparrow$	0.700	0.630	0.630	0.650	0.630	0.690	0.630	0.660	0.640	0.650	0.650	0.680	0.690	0.630	0.660	0.640	0.670	0.630	0.700	0.660	0.660	0.650	0.680	0.670	0.640	0.700	0.660	0.640	0.700	0.640	0.700	
	$\Delta_3 \uparrow$	0.564	0.528	0.542	0.548	0.542	0.540	0.548	0.544	0.554	0.548	0.564	0.554	0.540	0.548	0.564	0.554	0.532	0.564	0.532	0.558	0.560	0.564	0.554	0.532	0.560	0.538	0.560	0.560	0.538	0.560	0.528	
COPA-Test	$\Delta_1 \uparrow$	0.564	0.526	0.554	0.542	0.554	0.540	0.544	0.544	0.548	0.542	0.564	0.548	0.540	0.544	0.564	0.548	0.538	0.564	0.534	0.562	0.556	0.564	0.546	0.538	0.562	0.536	0.562	0.556	0.562	0.536	0.526	
	$\Delta_2 \uparrow$	0.665	0.503	0.600	0.606	0.600	0.594	0.594	0.613	0.503	0.606	0.639	0.503	0.594	0.594	0.613	0.503	0.594	0.639	0.510	0.613	0.510	0.639	0.503	0.594	0.665	0.516	0.613	0.510	0.665	0.516	0.665	
	$\Delta_3 \uparrow$	0.697	0.503	0.594	0.606	0.600	0.594	0.594	0.619	0.510	0.600	0.639	0.503	0.606	0.594	0.619	0.510	0.600	0.697	0.510	0.619	0.516	0.639	0.503	0.600	0.690	0.516	0.619	0.516	0.690	0.516	0.690	
GLUCOSE-D1	$\Delta_1 \uparrow$	0.620	0.550	0.590	0.580	0.580	0.580	0.580	0.570	0.620	0.550	0.560	0.610	0.580	0.580	0.570	0.620	0.580	0.560	0.620	0.560	0.620	0.560	0.610	0.580	0.560	0.610	0.560	0.560	0.610	0.560	0.610	
	$\Delta_2 \uparrow$	0.630	0.530	0.610	0.600	0.610	0.600	0.580	0.600	0.630	0.530	0.550	0.600	0.600	0.580	0.600	0.630	0.580	0.580	0.610	0.580	0.620	0.550	0.600	0.580	0.560	0.610	0.580	0.620	0.560	0.610	0.610	
	$\Delta_3 \uparrow$	0.580	0.484	0.494	0.486	0.494	0.522	0.498	0.484	0.574	0.486	0.496	0.574	0.522	0.498	0.484	0.574	0.514	0.512	0.570	0.506	0.580	0.496	0.574	0.514	0.512	0.570	0.506	0.580	0.512	0.570	0.570	
COPA-Test (-T)	$\Delta_1 \uparrow$	0.574	0.484	0.494	0.502	0.494	0.530	0.508	0.484	0.574	0.502	0.492	0.570	0.530	0.508	0.484	0.574	0.528	0.524	0.570	0.494	0.574	0.492	0.570	0.528	0.522	0.570	0.494	0.574	0.522	0.570	0.522	
	$\Delta_2 \uparrow$	0.606	0.510	0.568	0.555	0.568	0.574	0.568	0.535	0.606	0.555	0.510	0.587	0.574	0.568	0.535	0.606	0.568	0.529	0.606	0.542	0.594	0.510	0.587	0.568	0.535	0.600	0.542	0.594	0.535	0.600	0.535	
	$\Delta_3 \uparrow$	0.619	0.503	0.568	0.555	0.568	0.587	0.561	0.503	0.619	0.555	0.516	0.587	0.581	0.561	0.503	0.619	0.581	0.529	0.613	0.548	0.594	0.516	0.587	0.581	0.535	0.613	0.548	0.594	0.535	0.613	0.535	
GLUCOSE-D1 (-T)	$\Delta_1 \uparrow$	0.619	0.503	0.568	0.555	0.568	0.587	0.561	0.503	0.619	0.555	0.516	0.587	0.581	0.561	0.503	0.619	0.581	0.529	0.613	0.548	0.594	0.516	0.587	0.581	0.535	0.613	0.548	0.594	0.535	0.613	0.535	
	$\Delta_2 \uparrow$	0.619	0.503	0.568	0.555	0.568	0.587	0.561	0.503	0.619	0.555	0.516	0.587	0.581	0.561	0.503	0.619	0.581	0.529	0.613	0.548	0.594	0.516	0.587	0.581	0.535	0.613	0.548	0.594	0.535	0.613	0.535	
	$\Delta_3 \uparrow$	0.619	0.503	0.568	0.555	0.568	0.587	0.561	0.503	0.619	0.555	0.516	0.587	0.581	0.561	0.503	0.619	0.581	0.529	0.613	0.548	0.594	0.516	0.587	0.581	0.535	0.613	0.548	0.594	0.535	0.613	0.535	

Table A.5: **Full ablation studies on normalizations.** Ablations resulting in the best performances are highlighted in blue and those resulting in the worst the performances are highlighted in red. Shaded rows are results without temporal fine-tuning (using top $k = 30$ tokens in mask language modeling). (i) The gap between best and worst performance are marginal, except for the GLUCOSE-D1 dataset, which is mainly caused by enabling estimand normalization (E). Without considering E, the worst result is 0.594 (+Q or +FQ). (ii) In general, temporal fine-tuning helps. The only exception on COPA-Test is caused by the estimand normalization (E). (iii) As a general rule-of-thumb, it does not hurt to start with no normalizations enabled.

Sampled Covariates X	$ q(x;A) - q(x;E_1) _p$	E_1 and Interventions A	$\mathbb{P}(\cdot \rightarrow E_2)$
X₁: He had written a brief book summary of the book and, using a set of questions. X₂: There was homework help, help desk, and online support. X₃: No one did the work on time, and no one received good grades for it. X₄: The kids had to do their school homework online. X₅: This was the norm. X₆: The class would sit quietly and listen to their teacher talk. X₇: It was Free Time. X₈: Homework was only assigned when the teacher had a class with a lot of work for the students to. X₉: He did not give homework to his students. X₁₀: There was a long period of time when nobody ever did any homework. X₁₁: She had been teaching them during class for weeks. X₁₂: It was just a fun afternoon with the kids, and then it turned into a time of dr. X₁₃: Every night, a student would get to work on their homework. X₁₄: The students had to listen to music and watch a video, respectively, before they could do their. X₁₅: The students had to do the homework themselves. X₁₆: The children used to go to school in the morning and study their books until the evening. X₁₇: The teacher assigned homework to the entire class. X₁₈: Beth student was given a piece of paper with some number on it. X₁₉: They could play without limits. X₂₀: I thought that the homework was just a part of my study in each class. X₂₁: Students who have not completed their homework will not be allowed to go to the next class. X₂₂: The assignment was simple, they were just to read the assigned reading. X₂₃: However, he handed out the following set of questions, which the teacher posed one by one to. X₂₄: Students were not given much homework. X₂₅: Students were only encouraged to work on assignments and were not explicitly told to do extra. X₂₆: Only the school's teacher did so. X₂₇: He would just talk to them or read articles or give his own opinion on the subject. X₂₈: There was no homework. X₂₉: The students had to read the textbook and test their knowledge of the material. X₃₀: They were assigned to do some homework. X₃₁: Teachers would typically assign the work to the students, but this teacher assigned it to the students and. X₃₂: There were no homework assignments at all. X₃₃: The students would go to the Internet and download games. X₃₄: The assignment had already been completed. X₃₅: He asked his students on the first day of class to write down on A4 paper any questions. X₃₆: No homework was assigned. X₃₇: The students were all in the classroom, sitting in rows like the soldiers in the First World War. X₃₈: I just gave them a paper with one page written on it. X₃₉: There was no homework. X₄₀: The students were told the homework, and the students were to do the homework on their own.	0 0.0185 0.0508 0.0894 0.0279 0.1053 0.1201 0.0591 0.0365 0.0201 0.0870 0.1321 0.0820 0.1524 0.1349 0.1399 0.0468 0.0485 0.0382 0.0135 0.0488 0.0512 0.0515 0.0617 0.0298	E₁:The teacher assigned homework to the students. A₁:The professor assigned homework to the students. A₂:The professor supported the tourists assigned homework to the students. A₃:The tourists ran, or the teacher assigned homework to the students. A₄:The teacher took homework to the students. A₅:The teacher was assigning Justin with the homework to the students. A₆:The teacher replaced the carpet for the library last night because the carpet was old homework to the students. A₇:The teacher assigned to read the children's book came to the students. A₈:The teacher assigned less homework to the students. A₉:The teacher assigned tests to the students. A₁₀:No one was assigned homework to the students. A₁₁:Unless the senator performed, the teacher assigned homework to the students. A₁₂:Noelle Leong on the other hand assigned homework to the students. A₁₃:The teacher didn't give homework to the students. A₁₄:The teacher didn't assign homework to the students. A₁₅:The teacher didn't tell anyone homework to the students. A₁₆:The teacher assigned nothing to the students. A₁₇:The teacher assigned no children to the students. A₁₈:The teacher assigned no class to the students. A₁₉:The student assigned homework to the students. A₂₀:The professor assigned homework to the students. A₂₁:The teacher worked on the algebraic homework to the students. A₂₂:The teacher wrote homework to the students. A₂₃:The teacher read homework to the students. A₂₄:The teacher assigned tests to the students. A₂₅:The teacher assigned to the classroom stopped to the students. A₂₆:The teacher assigned anger to the students.	0.5031 0.4993 0.5082 0.4987 0.5177 0.3935 0.5093 0.5120 0.5135 0.4999 0.4886 0.3867 0.4874 0.5524 0.5198 0.5140 0.5175 0.5230 0.5167 0.4958 0.4993 0.5334 0.5151 0.5231 0.4999 0.5244 0.5065

Table A.6: **Example 1a:** the first plausible pair of the 72-th instance in COPA-DEV, matched interventions are highlighted. Here E_1 : The teacher assigned homework to the students. and E_2 : The students passed notes.

Sampled Covariates $\mathcal{X}$	$\ q(x; A) - q(x; E_1)\ _p$	E_1 and Interventions $\mathcal{A}$	$\mathbb{P}(\cdot \rightarrow E_2)$
X₁: He had written a brief book summary of the book and, using a set of questions. X₂: There was homework help, help desk, and online support. X₃: No one did the work on time, and no one received good grades for it. X₄: The kids had to do their school homework online. X₅: This was the norm. X₆: The class would sit quietly and listen to their teacher talk. X₇: It was free time. X₈: Homework was only assigned when the teacher had a class with a lot of work for the students to. X₉: He did not give homework to his students. X₁₀: There was a long period of time when nobody ever did any homework. X₁₁: She had been teaching them during class for weeks. X₁₂: It was just a fun afternoon with the kids, and then it turned into a time of dr. X₁₃: Every night, a student would get to work on their homework. X₁₄: The students had to listen to music and watch a video, respectively, before they could do their. X₁₅: The students had to do the homework themselves. X₁₆: The children used to go to school in the morning and study their books until the evening. X₁₇: The teacher assigned homework to the entire class. X₁₈: Each student was given a piece of paper with some number on it. X₁₉: They could play without limits. X₂₀: I thought that the homework was just a part of my study in each class. X₂₁: Students who have not completed their homework will not be allowed to go to the next class. X₂₂: The assignment was simple, they were just to read the assigned reading. X₂₃: However, he handed out the following set of questions, which the teacher posed one by one to. X₂₄: Students were not given much homework. X₂₅: Students were only encouraged to work on assignments and were not explicitly told to do extra. X₂₆: Only the school's teacher did so. X₂₇: He would just talk to them or read articles or give his own opinion on the subject. X₂₈: There was no homework. X₂₉: The students had to read the textbook and test their knowledge of the material. X₃₀: Teachers would typically assign the work to the students, but this teacher assigned it to the students and. X₃₁: There were no homework assignments at all. X₃₂: The students would go to the Internet and download games. X₃₃: The assignment had already been completed. X₃₄: He asked his students on the first day of class to write down on A4 paper any questions. X₃₅: No homework was assigned. X₃₆: The students were all in the classroom, sitting in rows like the soldiers in the First World War. X₃₇: I just gave them a paper with one page written on it. X₃₈: There was no homework. X₃₉: The students were told the homework, and the students were to do the homework on their own.	0 0.0185 0.0508 0.0894 0.0279 0.1063 0.1201 0.0591 0.0365 0.0201 0.0870 0.1521 0.0820 0.1524 0.1349 0.1999 0.0468 0.0485 0.0362 0.0301 0.0135 0.0488 0.0512 0.0515 0.0201 0.0647 0.0208	E₁: The teacher assigned homework to the students. A₁: The professor assigned homework to the students. A₂: The professor supported the tourists assigned homework to the students. A₃: The tourists ran, or the teacher assigned homework to the students. A₄: The teacher took homework to the students. A₅: The teacher was assigning Justin with the homework to the students. A₆: The teacher replaced the carpet for the library last night because the carpet was old homework to the students. A₇: The teacher assigned to read the children's book came to the students. A₈: The teacher assigned less homework to the students. A₉: The teacher assigned tests to the students. A₁₀: No one was assigned homework to the students. A₁₁: Unless the senator performed, the teacher assigned homework to the students. A₁₂: Noelle Leong on the other hand assigned homework to the students. A₁₃: The teacher didn't give homework to the students. A₁₄: The teacher didn't assign homework to the students. A₁₅: The teacher didn't tell anyone homework to the students. A₁₆: The teacher assigned nothing to the students. A₁₇: The teacher assigned no children to the students. A₁₈: The teacher assigned no class to the students. A₁₉: The student assigned homework to the students. A₂₀: The professor assigned homework to the students. A₂₁: The teacher worked on the algebraic homework to the students. A₂₂: The teacher wrote homework to the students. A₂₃: The teacher read homework to the students. A₂₄: The teacher assigned tests to the students. A₂₅: The teacher assigned to the classroom stopped to the students. A₂₆: The teacher assigned anger to the students.	0.5308 0.5263 0.5217 0.5280 0.5340 0.5396 0.5015 0.5346 0.5515 0.5249 0.5832 0.5454 0.5288 0.5914 0.5866 0.6164 0.5487 0.5566 0.5477 0.5180 0.5263 0.5349 0.5318 0.5370 0.5249 0.5263 0.5291

Table A.7: **Example 1b:** the second plausible pair of the 72-th instance in COPA-DEV, matched interventions are highlighted. Here E_1 : The teacher assigned homework to the students . and E_2 : The students groaned.

Sampled Covariates, $\mathcal{X}$	$\ q(x; A) - q(x; E_1)\ _p$	E_1 and Interventions, A	$\mathbb{P}(\cdot \rightarrow E_2)$
X_1 : I had scrubbed my face, arms, and chest, using a baby shampoo called "Sun. X_2 : There was the bathroom, and that was a little bit trickier. X_3 : I had got dressed. X_4 : I had been standing up because my knees hurt, and they were stiff. X_5 : I had put on a new pair of latex gloves-I'm very careful about hand cleaning- X_6 : I wanted to take out my medicine and check all my symptoms. X_7 : I had put a couple of paper towels in the drawer by the sink. X_8 : I had been sitting in the armchair by the fire. X_9 : I had decided to make a cup of tea. X_{10} : I had been playing with my son, watching an old video on YouTube, and I. X_{11} : I had been brushing the sand from my clothes. X_{12} : I went through the washing ceremony to check the level of purity in my body. I washed my. X_{13} : I washed my face. X_{14} : I had just finished eating my breakfast. X_{15} : I prepared a simple salad and some rolls on the table. X_{16} : I scrubbed my hands with a little bit of soap. X_{17} : I turned to the side of the mirror, and I had a look. X_{18} : I always take my shoes off. X_{19} : However, I removed some leftover food from the table, where the two men had been eating. X_{20} : I took a few deep breaths and had a conversation with my heart. X_{21} : I had changed into the outfit I was wearing: a pretty, pale pink T-shirt and. X_{22} : I was clearing away breakfast things. X_{23} : I used to take a shower, and now it was time to do that again. X_{24} : I'd washed my face. X_{25} : I needed to check my phone. X_{26} : I'd taken the time to put on another pair of socks, and the socks for that matter. X_{27} : I washed my hands more than a thousand times. X_{28} : I needed to put on a gown and cap, and to check the medications I had received for. X_{29} : I'd been sitting at my desk, answering emails and making phone calls. X_{30} : I'd touched the wall for some reason. X_{31} : I used to dry then properly. X_{32} : I made sure my hands were Clean. X_{33} : As a last resort, I would always scrub the top of my hands with a nail brush to. X_{34} : I had looked into the bathroom mirror. X_{35} : I was talking to him. X_{36} : I'd decided to drink some water, which was, of course, a bad idea. X_{37} : I was standing in the room, with the window open. X_{38} : Of course, I had put my coat on. X_{39} : I used to wash my hands. X_{40} : Though, I'd pulled a towel off the rack and was drying my hair. X_{41} : I was taking a shower. X_{42} : I'd been wiping my palms on the sides of shorts and shirt, like a dirty secret. X_{43} : I had been holding a glass of orange juice, which I had drained, and a bowl of. X_{44} : I brushed my teeth and brushed my hair. I even dried my hair, and then I went. X_{45} : I would take a breath. X_{46} : I'd brushed my teeth, put on deodorant, shaved, dried my hair. X_{47} : I had brushed my teeth, applied makeup, and removed my contacts. X_{48} : I was making some tea. X_{49} : I'd removed all my jewelry. X_{50} : I had to take my shoes off.	0 0.2485 0.1702 0.2153 0.0752 0.1014 0.1210 0.1067 0.1424 0.3054 0.0700 0.0586 0.0912 0.1014 0.0157 0.0398 0.0324 0.1373 0.1263 0.1601 0.3740 0.3322 0.1291 0.1204 0.4810 0.0940 0.0700 0.2772 0.2068 0.0970 0.0000 0.1389 0.1404 0.0324 0.2359 0.0919 0.0966 0.0000 0.0682 0.0000 0.1424 0.0325 0.0000 0.0659 0.0754 0.0489 0.0783 0.0000	E_1 : I was preparing to wash my hands. A_1 : I was running low and got my hands wet but not the shoes because the hands were preparing to wash my hands. A_2 : I was standing close to the sink and was suddenly wet from the rain because the sink was well lit preparing to wash my hands. A_3 : I wanted to get rid of the smell of bleach and use water instead because the water was clean and preparing to wash my hands. A_4 : The person was preparing to wash my hands. A_5 : I've was preparing to wash my hands. A_6 : I guess was preparing to wash my hands. A_7 : I was about to start using dish soap to wash my hands. A_8 : I was going to wash my hands. A_9 : I was late for work so I was running and picking up the dishes. So I got to wash my hands. A_{10} : Berty was preparing to wash my hands. A_{11} : I was preparing to use lube to get out of my hands. A_{12} : I was preparing to take a vitamin c and a calcium supplement. I took my hands. A_{13} : I was preparing to cook dinner my hands. A_{14} : I was preparing to wash my feet. A_{15} : I was preparing to wash my face and hair. A_{16} : I didn't preparing to wash my hands. A_{17} : I don't know how to describe an incredible preparing to wash my hands. A_{18} : I didn't wash my hands with soap but instead used conditioner because the soap was preparing to wash my hands. A_{19} : Maybe it's because my dog died recently, or because my wife was sick that week. I don't know, but I was preparing to wash my hands. A_{20} : I didn't wash my hands was preparing to wash my hands. A_{21} : I can't wash my hands was preparing to wash my hands. A_{22} : I was not supposed to wash my hands. A_{23} : I was not supposed to be doing dishes after dinner, so I was going to wash my hands. A_{24} : I was not sure if I needed to use soap or Vingar to clean to wash my hands. A_{25} : Berty was preparing to wash my hands. A_{26} : No part of the preparation except was preparing to wash my hands. A_{27} : However, as I mentioned already, time to was preparing to wash my hands. A_{28} : I was preparing to wash my eyes but switched to water my hands. A_{29} : I was preparing to wash my hands. A_{30} : I was preparing to skip the soap and water. I couldn't because the soap wasn't needed for my hands. A_{31} : I was preparing to wash my hands and I was doing too much. A_{32} : I was preparing to wash the clothes. A_{33} : I was preparing to wash my hands but had to get water since the hands were not touching. A_{34} : I tried to wash preparing to wash my hands. A_{35} : I needed to get rid of my feet preparing to wash my hands. A_{36} : I was preparing to wash my hands. A_{37} : He was preparing to wash my hands. A_{38} : I was preparing to wash my hands. A_{39} : I was going to wash my hands. A_{40} : I was preparing to cook to wash my hands. A_{41} : It was preparing to wash my hands. A_{42} : The towel was preparing to wash my hands. A_{43} : I was preparing to put my hands. A_{44} : I was preparing to ditch my hands. A_{45} : I was preparing to wash my hands.	0.4817 0.4177 0.4072 0.3779 0.4754 0.4066 0.3801 0.4440 0.3987 0.3302 0.4601 0.4817 0.4764 0.4861 0.4917 0.4923 0.5002 0.4339 0.4227 0.4513 0.2685 0.2046 0.2412 0.4055 0.3228 0.4754 0.4001 0.4067 0.4536 0.4884 0.4817 0.4039 0.4780 0.5002 0.4759 0.5058 0.4861 0.4817 0.4971 0.4817 0.3987 0.4913 0.4817 0.4992 0.5118 0.4752 0.4679 0.4817

Table A.8: **Example 2a**: the first plausible pair of the 63-th instance in COPA-DEV, matched interventions are highlighted. Here E_1 : I was preparing to wash my hands. and E_2 : I put rubber gloves on.

Sampled Covariates $\mathcal{X}$	$\ g(x; A) - g(x; E_1)\ _p$	E_1 and Interventions A	$\mathbb{P}(c = E_2)$
X₁: I had scrubbed the kitchen floor and the sink and, uh, that kind of thing. X₂: There was the need to remove the rubbish from the front garden. X₃: I had been walking up and down the hall, running my hands over the wood-paneled. X₄: I had put a load of clothes, some books and some DVD's into the washing machine. X₅: I wanted to take out the trash and empty all the containers. X₆: I had put a couple of washcloths on the bathroom counter. X₇: I had had to deal with the dirty clothes hamper and the clothes on the floor. X₈: I had decided to clean the stove. X₉: I had decided to take a short break and get some ice-cream. X₁₀: I had vacuumed the room. X₁₁: I went down to the living room and turned on the TV. X₁₂: I needed to find a bottle opener. X₁₃: I prepared a simple salad and some crackers and cheeses in our very cute wooden bowl. X₁₄: I scrubbed down the sink, the counter, the tile and my hands. X₁₅: I turned on the TV to catch my favourite news programme, when the host came on and said. X₁₆: I always take my shower. X₁₇: I took a bath and washed my hair. X₁₈: I had to clean the house. X₁₉: I showered, put on a clean pair of pants and a shirt. X₂₀: I did my normal routine. X₂₁: I'd washed the coffee table, and before that, I'd vacuumed the floor. X₂₂: I needed to flush the toilet. X₂₃: I went to the kitchen to put away another load of dishes. X₂₄: I needed to put on a pair of rubber gloves. X₂₅: I'd already wiped the kitchen counter. X₂₆: I would take out the trash. X₂₇: I used to clean the living room and the kitchen, and even the bathroom sometimes. X₂₈: I made sure the refrigerator was stocked. X₂₉: As a last resort, I would always open the medicine cabinet and remove any expired birth control pills. X₃₀: I had to clear away the table, then rinse dishes and clean the table. X₃₁: I checked on the kids, who were doing what they normally did. X₃₂: The kitchen. X₃₃: I'd decided to organize the cabinets in the kitchen, because organizing might calm me down. X₃₄: I was doing laundry. X₃₅: Of course, I had put my makeup on. X₃₆: I turned on the radio. X₃₇: Though, I'd go to the kitchen and make a small snack for myself. X₃₈: I was taking a shower. X₃₉: I made myself a green smoothie. X₄₀: I would remove the dishes stored in the kitchen, and then, at last, I would clean. X₄₁: I brushed my teeth and took a shower, but I was not very interested in either task. X₄₂: I would take a good look at all of the objects in the room, cataloging and organizing. X₄₃: I'd brushed my teeth, put on deodorant and makeup, and made sure. X₄₄: I took a shower. X₄₅: I was making some soup. X₄₆: I had to take the covers off of my bed and the bedspread. X₄₇: I had to clear all the stuff out of my room. X₄₈: I took a shower and had a leisurely breakfast. X₄₉: I'd had a shower, washed my face, dried it with a towel then put. X₅₀: I was going to get us a drink of water and maybe some dinner.	0 0.2191 0.1228 0.0598 0.0703 0.0325 0.0651 0.1247 0.1402 0.0952 0.0000 0.0678 0.0513 0.0583 0.1601 0.0437 0.0549 0.0798 0.1496 0.0645 0.2527 0.0717 0.0505 0.1418 0.1548 0.1166 0.0745 0.0505 0.0593 0.1011 0.0773 0.1399 0.0549 0.1130 0.0430 0.0000 0.0814 0.1086 0.0651 0.0890 0.0454 0.3539 0.0952 0.1586 0.0000 0.0293 0.0931 0.1011	E₁: I was preparing to clean the bathroom. A₁: I was building the car instead of the house since the house was incompatible with cleanliness. preparing to clean the bathroom. A₂: I was so bad at the job that I was preparing to clean the bathroom. A₃: I was doing a job preparing to clean the bathroom. A₄: A woman was preparing to clean the bathroom. A₅: Kevin was preparing to clean the bathroom. A₆: Emily was preparing to clean the bathroom. A₇: I was going to take a bath instead of to clean the bathroom. A₈: I was able to do the cleaning in time, and was pretty good at it, although the to clean the bathroom. A₉: I was going to clean the bathroom. A₁₀: I was preparing to clean the bathroom. A₁₁: Bill was preparing to clean the bathroom. A₁₂: Empty was preparing to clean the bathroom. A₁₃: I was preparing to sleep the bathroom. A₁₄: I was preparing to cook dinner the bathroom. A₁₅: I was preparing to cook dinner for my family. the bathroom. A₁₆: I was preparing to clean the kitchen floor. A₁₇: I was preparing to clean the kitchen table. A₁₈: I wasn't preparing to clean the bathroom. A₁₉: I didn't want to wash either the towels or the sponge. I was preparing to clean the bathroom. A₂₀: I was not preparing to clean the bathroom. A₂₁: When I was done, I was preparing to clean the bathroom. A₂₂: I wasn't was preparing to clean the bathroom. A₂₃: no one was preparing to clean the bathroom. A₂₄: I was not rushing too much and didn't get to clean the bathroom. A₂₅: I was not able to do the dishes, so I had to do the dishes. I had no problem washing to clean the bathroom. A₂₆: I was not going to to clean the bathroom. A₂₇: I was not supposed to was preparing to clean the bathroom. A₂₈: no one was preparing to clean the bathroom. A₂₉: No one was preparing to clean the bathroom. A₃₀: I was preparing to cook dinner the bathroom. A₃₁: I was preparing to skip the whole the bathroom. A₃₂: I was preparing to wash my hands, but I forgot to take the bathroom. A₃₃: I was preparing to clean the kitchen table. A₃₄: I was preparing to clean the bathroom all the time. I was not able to clean the bathroom all of the time and it would take forever. A₃₅: I was preparing to clean the kitchen counter. A₃₆: I was preparing to clean the bathroom. A₃₇: I smelled the cookies and wondered if mom was preparing to clean the bathroom. A₃₈: I knew it could be used by anyone, just a stranger, preparing to clean the bathroom. A₃₉: Emily was preparing to clean the bathroom. A₄₀: I was showering was preparing to clean the bathroom. A₄₁: she was preparing to clean the bathroom. A₄₂: I was hoping someone would fill me in on the details, so I tried to write to clean the bathroom. A₄₃: I was going to clean the bathroom. A₄₄: I was going through the old photos and couldn't decide which mirror to use for my mirror. The Mirror was too old to clean the bathroom. A₄₅: I was preparing to clean the bathroom. A₄₆: I was preparing to put the clothes on the bathroom. A₄₇: I was preparing to put in the sewing machine the bathroom. A₄₈: I was preparing to cook dinner the bathroom.	0.5023 0.3765 0.4935 0.5171 0.5083 0.4880 0.4522 0.3744 0.3202 0.4437 0.5023 0.4967 0.4727 0.5102 0.4308 0.4288 0.5098 0.5073 0.4573 0.3829 0.4836 0.2748 0.4897 0.4948 0.4809 0.4008 0.4419 0.5107 0.4948 0.4911 0.5102 0.4662 0.4426 0.5073 0.4897 0.5059 0.5023 0.4830 0.4470 0.4522 0.5224 0.5137 0.3144 0.4437 0.4813 0.5023 0.5046 0.4958 0.5102

Table A.9: **Example 2b**: the second plausible pair of the 63-th instance in COPA-DEV, matched interventions are highlighted. Here E_1 : I was preparing to clean the bathroom. and E_2 : I put rubber gloves on .

Sampled Covariates X	$\ q(x;A) - q(x;E_1)\ _p$	E_1 and Interventions A	$\mathbb{P}(x \prec E_2)$
X_1: He'd had nothing in his pockets but his father's pocket watch, and some old coins he'd. X_2: There had been the time he'd been a little boy, about four years old, and. X_3: He had been a very small fish in a very small pond. X_4: It had contained a knife and a set of keys-but there was nothing in it now. X_5: It seemed only to contain his breath and blood. X_{10}: He was a slave, and his owner used him roughly when displeased. X_{11}: He had put a couple of coins in his pouch. X_{12}: His wallet had been empty. X_{13}: He had been a slave-hunter to the west, a murderer in the service of his city. X_{14}: It had been half-filled with tobacco and a dirty handkerchief. X_{15}: He had been working as a waiter at the Côte du Rhône at the Hotel Mont. X_{16}: He couldn't even have given the goldfish. X_{17}: It was empty, but he could think of no problem more urgent than collecting them. X_{18}: He'd been a beggar. X_{19}: He had been a farmer. X_{20}: The contents of his pockets would have been only a pair of pants, a shirt and, if it. X_{21}: He'd hidden them at the old folk's home, where he'd lived. X_{22}: He had only had the money to keep him going. X_{23}: He didn't feel too well. X_{24}: He'd been just like them. X_{25}: The police had come and taken the man's wallet. X_{26}: He had been wearing a thick bracelet with a chain and gold links, a gift from the wife of. X_{27}: He'd been wearing a plain suit and polished his shoes all the time, that would no longer do. X_{28}: He used to keep spare coins in his shirt pocket, and use those for the fare. X_{29}: He'd been a thief, a beggar and a soldier. X_{30}: No one had seen it. X_{31}: He had not worked for a long time. X_{32}: The two had been making their way through a series of shops. X_{33}: He'd been standing with his back to the wall, holding an umbrella over his head. X_{34}: He had been trying to buy himself with the money that came from the sale of his books. X_{35}: Of course, he had been a slave. X_{36}: He used to have some. X_{37}: He'd carried his entire life in his pockets. X_{38}: He was wearing a shapeless dark coat with a dark suit underneath, a black hat, and. X_{39}: He'd been a poor kid that the state had taken from his mother and put in one home after. X_{40}: A few years back, he had an old steel frame. X_{41}: He would have sold his coat, and his shirt, and then the little shoes upon his feet. X_{42}: He must have been a poor man, and probably very pious.	0 0.1069 0.0706 0.0620 0.1457 0.0626 0.0038 0.0747 0.0364 0.1108 0.2276 0.1331 0.1170 0.1539 0.1500 0.1126 0.3496 0.1112 0.0897 0.0000 0.0820 0.2621 0.0233 0.0481 0.0753 0.0370	E_1: His pocket was filled with coins. A_1: His pocket however had been filled with coins. A_2: His pocket contained nine corsage filled with coins. A_3: His pocket had a large amount of space filled with coins. A_4: His pocket was lost when he accidentally took some with coins. A_5: His pocket was empty with coins. A_6: His pocket was filled with sandals and at least one with coins. A_7: A cowboy on the back of a wagon was filled with coins. A_8: A pocket iron was filled with coins. A_9: Mark was filled with coins. A_{10}: His pocket did not work as well as the shirt, because the shirt was filled with coins. A_{11}: His pocket wasn't filled with coins. A_{12}: His pocket was not filled with coins. A_{13}: His pocket was empty but his pocket had nine with coins. A_{14}: His pocket was not holding the lotion with coins. A_{15}: His pocket was not touched with coins. A_{16}: No matter how you feel about country music [...] the fact that it featured the really catchy John Denver does not appeal to me was filled with coins. A_{17}: No pocket book was filled with coins. A_{18}: Nothing was filled with coins. A_{19}: His pocket was filled with coins. A_{20}: His pocket had been filled with coins. A_{21}: His pocket was ruined and he decided to find a wallet. Instead because the pocket might have coins with coins. A_{22}: His pocket was full with coins. A_{23}: Her bag was filled with coins. A_{24}: Someone was filled with coins. A_{25}: Her wallet was filled with coins.	0.2080 0.1494 0.3912 0.3036 0.0620 0.2749 0.2507 0.2331 0.1933 0.1925 0.0094 0.1870 0.2477 0.1345 0.0834 0.1852 0.0209 0.1071 0.2345 0.2080 0.3352 0.2535 0.0913 0.3068 0.1762 0.2306 0.0217

Table A.10: **Example 3a:** the first plausible pair of the 79-th instance in COPA-DEV, matched interventions are highlighted. Here E_1 : His pocket was filled with coins. and E_2 : The man's pocket jingled as he walked.

Sampled Covariates $\mathcal{X}$	$\ q(x;A) - q(x;E_1)\ _p$	E_1 and Interventions A_1	$\mathbb{P}(x \prec E_2)$
X₁: He'd had nothing in his pocket to worry about. X₂: There had been no hole, just a thin line of cloth. X₃: He cut off all his fingers, including those on his left hand. X₄: He had put a knife in the pocket of his coveralls. X₅: He was a young, handsome fellow with dark blue eyes and black curly hair. X₆: He had put a couple of sticks in his pouch. X₇: His wallet had been in his hand, and he had thrown the wallet on the ground as. X₈: He had been a stranger to himself. X₉: It had been in his mouth. X₁₀: He went down to the river to check the damage. X₁₁: He stuffed a paper bag in place of his wallet. X₁₂: It was a secret, what with the police and all. X₁₃: He had been afraid to kill anyone. X₁₄: He'd hidden his heart, but not in a safe. X₁₅: He'd thought she'd just pulled it out of thin air, but there it was. X₁₆: He didn't feel too bad about taking the wallet from your wallet, but after he se. X₁₇: He'd been just like any other. X₁₈: The hole had been on the inside of his coat. X₁₉: He had been wearing his uniform cap with the rank insignia, and he put it on. X₂₀: He hadn't even looked at it. X₂₁: He'd sewn up the hole in his leg. X₂₂: There was nothing in it, not even the thorns, and there was nothing there but. X₂₃: No one had seen it. X₂₄: He had not wanted to talk about it. X₂₅: The two men had argued. X₂₆: He'd been thinking of the boy's parents. X₂₇: He had been trying to reach the hospital. X₂₈: He had put a small packet of powder in her shoe, had used a hairpin to. X₂₉: He'd been a good kid that the others seemed to like. X₃₀: He had hidden the bullet in his leg. X₃₁: As a little girl, before she had even known what it meant to be. X₃₂: He'd hidden it in the bottom of a pot of ointment. X₃₃: He had done nothing, but he could do nothing now but lie and wait, and be. X₃₄: It was not easy to find a place to hide the gun and if he was asked about. X₃₅: He'd had a small knife to cut his clothes, but he didn't. X₃₆: He always bought new jeans every time he went to Kmart. X₃₇: He'd had nothing. X₃₈: He had sewn the other pocket. X₃₉: The pocket was for the book. X₄₀: He had to get the tire.	0 0.1075 0.2456 0.0682 0.0680 0.0875 0.0849 0.2149 0.0930 0.0707 0.1181 0.2376 0.1161 0.1128 0.1479 0.0869 0.1401 0.3286 0.0871 0.0764 0.0456 0.0715 0.1560 0.0614 0.0555 0.0588 0.0311 0.0409	E₁: He sewed the hole in his pocket. A₁: Then he sewed the hole in his pocket. A₂: The boy was grumpy in high school, but happy at school, so the teacher taught him sewed the hole in his pocket. A₃: He cut pieces from the plate sewed the hole in his pocket. A₄: He stuffed the toy gun with the hole in his pocket. A₅: He burnt the hole by pulling the hole in his pocket. A₆: He pulled off the blanket and got a the hole in his pocket. A₇: He sewed the quilt better than a teepee because the teepee was a sloppy job in his pocket. A₈: He sewed with a towel more in his pocket. A₉: He sewed the hole with a wire instead of a plier in his pocket. A₁₀: He couldn't sewed the hole in his pocket. A₁₁: No matter how you feel about country music(I for one can't stand it despite my Houston roots), this only instilled sewed the hole in his pocket. A₁₂: No one knew how to sewed the hole in his pocket. A₁₃: He never filled the hole in his pocket. A₁₄: He couldn't bend the iron rod and instead tied the hole in his pocket. A₁₅: He had the hole in his pocket. A₁₆: He sewed no better than the tieclaxu which cut off his eye in his pocket. A₁₇: He sewed no better with the machine than with the method, because the machine was not precise in his pocket. A₁₈: He sewed not only the hole but also the whole ball inside the hole in his pocket. A₁₉: Jack sewed the hole in his pocket. A₂₀: Someone sewed the hole in his pocket. A₂₁: He screwed up sewed the hole in his pocket. A₂₂: He flunked out of high school, ended up in a strange town, and started writing about the weird the hole in his pocket. A₂₃: He stabbed Wisbech with a mop the hole in his pocket. A₂₄: He pulled the hole in his pocket. A₂₅: He sewed the turkey with a T-shirt in his pocket. A₂₆: He sewed the ball necklace in his pocket. A₂₇: He sewed the rope with a chisel in his pocket.	0.4818 0.3724 0.1477 0.5023 0.5053 0.3826 0.4970 0.2377 0.3728 0.4299 0.3049 0.1855 0.2986 0.2366 0.2361 0.3282 0.1917 0.0641 0.4182 0.4594 0.4086 0.4874 0.3031 0.4916 0.4673 0.5179 0.4765 0.4956

Table A.11: **Example 3b**: the second plausible pair of the 79-th instance in COPA-DEV, matched interventions are highlighted. Here E_1 : He sewed the hole in his pocket. and E_2 : The man's pocket jingled as he walked.