



Universiteit
Leiden

The Netherlands

Formal models of software-defined networks

Feng, H.

Citation

Feng, H. (2024, December 3). *Formal models of software-defined networks*.
Retrieved from <https://hdl.handle.net/1887/4170508>

Version: Publisher's Version

License: [Licence agreement concerning inclusion of doctoral thesis in the Institutional Repository of the University of Leiden](#)

Downloaded from: <https://hdl.handle.net/1887/4170508>

Note: To cite this publication please use the final published version (if applicable).

Chapter 7

Towards causality reasoning for SCA

In this chapter, we introduce an NFA causal model based on counterfactuals, inspired by the seminal works on causal analysis by Halpern and Pearl, adapted to finite automata models and with safety properties defined by regular expressions [26]. The latter encodes undesired execution traces. We devise a framework that computes actual causes, or minimal traces that lead to states enabling hazardous behaviors. Furthermore, our framework exploits counterfactual information and identifies modalities to steer causal executions toward alternative safe ones. This can provide systems engineers with valuable data for actual debugging and fixing erroneous behaviors. Our framework employs standard algorithms from automata theory, thus paving the way to further generalizations from finite automata to richer structures like probabilistic, KAT, and NetKAT automata. The ultimate goal is to extend the framework to symbolic constraint automata [34], so as to be applied for causal reasoning on SDNs, for example by using our Reo model presented in Chapter 4.

7.1 Introduction

Causal models and associated causal inference machinery are precious tools for the interpretation and explanation of systems failures. Current testing and verification frameworks such as equivalence checking, for instance, assess whether or not systems comply with their specifications, and at most will produce a counterexample in case the system fails. Causal analysis, instead, plays an important role in explaining complex phenomena that are actual sources of hazards by adding, for example, additional information to

counterexamples on how to avoid the hazard.

A notion of causality often embraced and adopted by computer scientists was introduced by Halpern and Pearl in their seminal works [47, 46]. Their causal model encodes complex logical structures of multiple events that contribute to undesired effects, or hazards. In essence, the model is based on the so-called alternative worlds, originally proposed by Lewis [76]. In short, Lewis assumes the existence of worlds satisfying a sufficiency condition, where both the cause and the effect occur, and other worlds satisfying a necessity condition, in which neither the cause nor the effect occurs. This enables formulating the counterfactual argument, which defines a first condition to be satisfied by a cause, namely: when the presumed cause does not occur, the effect will not occur either. More complex aspects such as redundancy and preemption are also captured by the causal model in [47, 46]. For intuition, redundancy refers to simultaneous events that play the same role in enabling an undesired effect. Orthogonally, preemption refers to subsequent events that have the same power to enable the effect. In both cases, the counterfactual test alone cannot determine the actual cause. Last, but not least, causes in the spirit of [46] comply with a minimality requirement which guarantees that only the relevant set of causal events is identified.

Related work. Over time, several notions of causality have been proposed, each of which is tailored to the type of the system under analysis, and associated correctness specifications. Of particular interest for this chapter are the works in [74, 23, 25]. The aforementioned results propose trace-based adoptions of causality á la Halpern and Pearl, applicable to automata models. These, in combination with model checking-based methodologies, enabled computing causes for the violation of safety and liveness properties in Kripke structures and labeled transition systems, for instance.

Our work is closely related to the contribution in [25]. Given an automaton model, the naive goal is to identify the shortest sequence of actions that enable the effect, i.e., that can bring the system into a hazardous state. These are called “causal traces”. Note that, in contrast with the often tedious counterexamples identified by model-checkers, the minimality of causal traces implies concise descriptions of systems faults. Thus, causal traces encode essential information for systems engineers, for instance, and they can serve as a debugging aid. As previously stated, in the spirit of Halpern and Pearl, our definition of causality imposes a sufficiency condition: namely, whenever a causal trace is executed, the effect is reached as well. However, important information on how to avoid/fix hazardous behaviors can be extracted based on the aforementioned set of alternative worlds (or traces in our model), that do not lead to an undesired effect. Hence, we designed our causal model in the spirit of the counterfactual criterion of Lewis and identified modalities to avoid hazardous scenarios. Similarly to [74, 23, 25], we call these escape options – “events causal by their non-occurrence”. This information

can be exploited to steer an execution towards an alternative safe one, with immediate applicability in synthesizing schedulers, for instance.

A rich body of work successfully exploited the counterfactual argument for fault analysis and debugging techniques. Examples related to counterexample explanation in model checking are the works in [44, 43, 92], for instance. In [43] the authors propose a framework for understanding errors in ANSI C programs, based on distance metrics for program executions. In [44] the cause describing the error includes the identification of source code fragments crucial to distinguishing success from failure and differences in invariants between failing and non-failing runs. Distance criteria have also been exploited in [92], in combination with the so-called nearest neighbor queries to perform fault localization. The why-because-analysis in [71] was used to reason about aviation accidents, in a framework where Lamport’s Temporal Logic of Actions (TLA) described both the behavior of a system, the (history of) hazards and the sequence of the states leading to an accident. The work in [114] provides a comprehensive approach to systematic debugging including, among others, delta debugging – a technique for isolating minimal input to reproduce an error.

For finer notions of causal dependencies that distinguish between interleaving and true concurrency, for instance, we refer to event structures [8, 88]. Nevertheless, in our work, we adhere to the approaches in [74, 23, 25], and do not take into consideration the order of events along execution traces.

Our contributions. We propose a shifting from the bisimulation setting presented in [25] to a trace-based setting in the context of regular languages and automata theory. The benefits are multifold. For instance, the paradigm change facilitates the application of more standard algorithms from automata theory, in contrast with the rather ad-hoc procedures in [74, 23, 25]. Furthermore, the current framework enables the use of an expressive logic for defining safety properties in terms of regular expressions (or automata), instead of the ordinary Hennessy-Milner logic. The language-based approach to causality enables representing both hazards and causal explanations in terms of automata – a format better accepted by engineers. In addition, in this chapter, we use regular languages (or full regular expressions including Kleene-star) to encode the non-occurrence of events. Previous related works such as [74, 23, 25] can only provide finite sets of runs steering an execution towards an alternative safe one. Orthogonal to the aforementioned results, the current approach entails a “may” semantics of causality, instead of “must”; nevertheless, we believe that the approach can be easily modified to cater to the “must” version. Besides, in contrast with the results in [25], steering executions are guaranteed not to jump over hazardous states by simply concatenating sequences causal by their non-occurrence and the causal trace. The ultimate goal of the current work is to generalize from finite automata to richer structures like probabilistic automata and NetKAT

automata [4, 37].

Structure of the chapter. In Section 7.2 we provide an overview of regular languages and associated automata theory aspects. A running example is introduced in Section 7.3. Section 7.4 defines the language-based model of causality, whereas in Section 7.5 we show how to compute actual causes and safe computations. In Section 7.6 we provide an experimental evaluation of our method and in Section 7.7 we discuss how our model can be extended with tests and assignments. Section 7.8 concludes our work.

7.2 Preliminaries

In this section, we recall a few basic facts about regular languages, finite automata, and regular expressions [80].

Let A be a finite set of actions that we refer to as an *alphabet*. A *word* or *string* over A is a finite sequence $a_1 \dots a_n$ of elements from A . We denote by ε the empty word, i.e. the sequence of length 0, and write A^* to denote the set of (possibly empty) words over A . A *language* L is just a subset of words, that is $L \subseteq A^*$. We call a word w' to be a *prefix* of a word w whenever $w = w'w''$. A word w' is said to be a *sub-word* of a word w , if w' is obtained by deleting one or more elements of A at some not necessarily adjacent positions in w . We denote by $\text{sub}(w)$ the set of all sub-words of w . Note that $\text{sub}(\varepsilon) = \emptyset$. Also, $\varepsilon \in \text{sub}(w)$ but $w \notin \text{sub}(w)$ for every non empty word w .

A *finite automaton* (FA) is a 5-tuple $M = (S, A, i, \longrightarrow, F)$, where S is a finite set of states, $i \in S$ is the initial state, $F \subseteq S$ is the set of accepting states and $\longrightarrow \subseteq S \times A \times S$ is the transition relation. For simplicity, we write $s \xrightarrow{a} t$ whenever $(s, a, t) \in \longrightarrow$. A transition relation is called deterministic if for all $s \in S$ and $a \in A$ if $s \xrightarrow{a} t_1$ and $s \xrightarrow{a} t_2$ then $t_1 = t_2$.

A string $w \in A^*$ is accepted by an automaton M from a state s if either (1) $w = \varepsilon$ and $s \in F$, or (2) $w = aw'$ and there exist $s \xrightarrow{a} t$ such that w' is accepted by M from the state t . The *language accepted* by a FA M is the set $L(M) = \{w \in A^* \mid M \text{ accepts } w \text{ from } i\}$. Since for every FA M , we can build an FA N with a deterministic transition relation such that $L(M) = L(N)$, without loss of generality we will consider only finite automata with a deterministic transition relation.

A language L over the alphabet A is said to be *regular* if there exists a finite automaton M accepting it, that is $L(M) = L$. The class of all regular languages is closed under union, intersection, concatenation, complement, and Kleene star. Here language union and intersections are the usual set-theoretic operations, whereas concatenation of two languages L_1 and L_2 is given by the set $L_1 \cdot L_2 = \{w_1w_2 \mid w_1 \in L_1 \wedge w_2 \in L_2\}$. Finally, for a language L , its Kleene star closure is defined by $L^* = \bigcup_{n \in \mathbb{N}} L^n$ where $L^0 = \{\varepsilon\}$ and $L^{n+1} = L \cdot L^n$ for all $n \in \mathbb{N}$, thus denoting the concatenation of a language with

itself a finite number of time.

In this chapter, we are interested in system communicating by message passing, and thus we will always assume that the alphabet A is partitioned in three disjoint subsets A_I , A_O , and A_P of input, output, and private actions, respectively. Notationally, for $a \in A$, we write $a?$ if a is an input action in A_I and $a!$ if a is an output action in A_O , and use no markings for private actions in A_P . We use σ to denote an action that can be either input, output, or private.

Let A and B be two alphabets with disjoint private actions, and assume the set P is disjoint from Q . Given two finite automata $M = (P, A, i, \rightarrow_M, E)$ and $N = (Q, B, j, \rightarrow_N, F)$ their *parallel composition* is defined by the finite automaton $M \parallel N = (P \times Q, \Sigma, \langle i, j \rangle, \rightarrow, E \times F)$ where $\Sigma_I = (A_I \setminus B_O) \cup (B_I \setminus A_O)$, $\Sigma_O = (A_O \setminus B_I) \cup (B_O \setminus A_I)$, $\Sigma_P = (A_P \cup B_P) \cup (A_I \cap B_O) \cup (A_O \cap B_I)$, and $\rightarrow$ is the least transition relation such that

$$\begin{array}{c} \frac{p \xrightarrow{\sigma}_M p' \quad \sigma \notin B}{\langle p, q \rangle \xrightarrow{\sigma} \langle p', q \rangle} \qquad \frac{q \xrightarrow{\sigma}_N q' \quad \sigma \notin A}{\langle p, q \rangle \xrightarrow{\sigma} \langle p, q' \rangle} \\ \frac{p \xrightarrow{a?}_M p' \quad q \xrightarrow{a!}_N q'}{\langle p, q \rangle \xrightarrow{a} \langle p', q' \rangle} \qquad \frac{p \xrightarrow{a!}_M p' \quad q \xrightarrow{a?}_N q'}{\langle p, q \rangle \xrightarrow{a} \langle p', q' \rangle} \end{array}$$

The topmost rules are about either private actions that are not affected by the other automaton or communication actions that do not involve the other automaton. The two rules at the bottom are about complementary communication actions $a!$ and $a?$ that are synchronized resulting in the private action a . Note that when $A = B$ with $A_I = B_O$, $A_O = B_I$, and $A_P = B_P = \emptyset$ then parallel composition reduces to the product automata where all actions synchronize. In the case A is completely disjoint from B then parallel composition results in the so-called shuffle product. Other variations of synchronization products could be defined similarly, including multi-process synchronization, hiding of successful communication, value passing synchronization (for a finite value domain), and synchronization parameterized by a finite subset of actions.

For the characterization of the parallel composition of two languages, we need first to introduce the *projection* function. Given two alphabets A_1 and A_2 we define the projection $\pi_i: (A_1 \cup A_2)^* \rightarrow A_i^*$ by $\pi_i(\varepsilon) = \varepsilon$, and $\pi_i(\sigma \cdot w) = \sigma \cdot \pi_i(w)$ if $\sigma \in A_i$, and $\pi_i(w)$ otherwise. Because projections are surjective functions they have inverse π_i^{-1} returning the set of strings that are projected into a given one. More precisely, we define the *inverse projection* by $\pi_i^{-1}(w) = \{x \in (A_1 \cup A_2)^* \mid \pi_i(x) = w\}$ for every $w \in A_i^*$. Projections and their inverses can be extended to languages by applying them to all the strings in the language. In general we have that $\pi_i(\pi_i^{-1}(L)) = L$ but for the converse it only holds that $L \subseteq \pi_i^{-1}(\pi_i(L))$. Note that if two alphabets A_1 and A_2 have disjoint private actions and we partition $A_1 \cup A_2$ as in the alphabet of the parallel composition of two automata, then projections will assign private actions of $A_1 \cup A_2$ to either private, input or output actions in A_i unambiguously. Similarly, inverse projections assign private

actions to private actions but may assign input and output actions to private ones.

The *parallel composition* of two languages $L_1 \subseteq A_1^*$ and $L_2 \subseteq A_2^*$ is the language $L_1 \parallel L_2$ on the alphabet $A_1 \cup A_2$ defined as $\pi_1^{-1}(L_1) \cap \pi_2^{-1}(L_2)$. The intersection takes care that dual communication actions will be synchronized, and that disjoint private events will be shuffled with the others. As expected, we have that $L(M_1 \parallel M_2) = L(M_1) \parallel L(M_2)$, implying that regular languages are closed under parallel composition [96].

We conclude this section by introducing extended regular expressions, that we may use as alternative syntax to FAs to reason about causality in complex systems composed of several components potentially communicating with each other.

Given an alphabet A including communication actions, *extended regular expressions* are given by the following grammar:

$$e ::= 0 \mid 1 \mid a \mid a? \mid a! \mid e; e \mid e + e \mid e \parallel e \mid e^*, \quad (7.1)$$

where $a \in A_P$, $a?$ implies $a \in A_I$, and $a!$ implies $a \in A_O$. In process theoretic terms, 0 denotes no behavior, and 1 denotes a terminating process. The further building blocks of processes are (communication) actions. Processes can be composed sequentially, non-deterministically, in parallel, or can loop a finite number of times. Communication between process terms is performed based on synchronizations between opposite communication actions, that play thus a sender, respectively, receiver role. In the sequel, we often use A as shorthand for the regular expression obtained by the finite set of every action in A , and $\neg a$ as a shorthand for the set of every action in A except a . Note that in general, we could extend negation to all regular expressions, as regular languages are closed under complement.

Ordinary regular expressions are expressions without any parallel composition. Except for the parallel composition, we assume that an action cannot be used as input and output in the same ‘sequential’ expression, i.e., regular expression with no occurrence of the $\parallel$ operator. With this mild restriction, we can associate each regular expression e a language $L(e)$ inductively as follows:

$$\begin{array}{lll} L(0) = \emptyset & L(e_1; e_2) = L(e_1) \cdot L(e_2) & \\ L(1) = \{\varepsilon\} & L(e_1 + e_2) = L(e_1) \cup L(e_2) & L(e^*) = L(e)^* \\ L(a) = \{a\} & L(e_1 \parallel e_2) = L(e_1) \parallel L(e_2) & \end{array}$$

It is well known [63] that the language of an ordinary regular expression is regular. The same holds for our extended regular expressions, as we have seen that regular languages are closed under parallel composition. This implies that for every (extended) regular expression e there exists an automaton M such that $L(e) = L(M)$. We will not describe the construction here as it is outside the scope of this chapter.

and the only accepting state is ⑧ corresponding to $\langle 5, 4, 1 \rangle$. The red states ③ and ④ will be used in the next section as examples of states leading to a hazard situation: a car entering the crossing and not leaving it before the train enters the crossing too.

7.4 A Language-based causal model

In this section, we introduce a notion of causality with respect to a so-called *hazard*, or *effect* expressed in terms of regular expressions. The current causal framework is inspired by the model introduced in [25] and massaged into the setting of FAs to use trace semantics instead of bisimulation, and define different system properties in terms of regular expressions (such as reachability) instead of the ordinary Hennessy-Milner logic.

In short, a hazard is a regular language specified by a regular expression e (or the corresponding automaton). It is said to occur in a FA M representing our model whenever there is a finite (and possibly empty) string $c = a_0 \dots a_n$ in M such that after c we may observe the hazard, that is, $L(c; e) \cap L(M) \neq \emptyset$. In this case, we say that c may enable the hazard e in M . Additional conditions that have to be satisfied by c , such as minimality and non-occurrence of events, are formalized in Definition 7.1.

For an intuition, consider the railway crossing example of the previous section. A hazardous situation can happen whenever both the train and the car enter the crossing, and none of them leaves the crossing before the other one enters it. The regular expression encoding this hazard is:

$$e = (Cc; (-Cl)^*; Tc + Tc; (-Tl)^*; Cc); A^* \quad (7.2)$$

Note that the hazard situation can terminate with any string in A^* . This is to guarantee that after a trace c enables e , their concatenation will contain behaviors accepted by the automaton, and thus the hazard is observed. It is straightforward to see that in the FA in Fig. 7.1 it is possible to reach the above hazard with the string $c_1 = CaGo$ leading to the state ③, but also with the string $c_2 = CaGoTa$ leading to the state ④. The intersection of the language of the hazard e with that of the automaton M starting from either state ③ or ④ instead of ① is non-empty. Furthermore, state ③ and ④ are both reachable from the initial state ①.

We may say that c_1 does a better job at describing the relevant sequence of actions that, if triggered, lead to a hazard because it is a *minimal* sequence enabling it. Moreover, we see that it is possible to avoid the hazard by “decorating” the string c_1 with the strings Ta , $TcTl$ and, respectively, $CcCl$. This can result, for instance, in the string $w = TaCaTcTlGoCcCl$ which does not lead to a hazard. Sequences such as Ta , $TcTl$ and $CcCl$ are called *causal by non-occurrence* in works such as [23, 25]. Non-

occurrence is essential for describing how certain dangerous situations, if controllable, can be avoided within a system. This concept plays an important role in our definition of causality.

As formalized in Definition 7.1, the non-occurrence of events is captured in terms of the so-called *computations* [25]. The latter are strings in a regular language, typically denoted by π , built on top of a string $c = a_0 \dots a_n$, and “decorated” with strings $d_0^i, \dots, d_{n+1}^i$, with $i \in I$, where I is a finite set of integers, such that:

$$w \in \pi \Rightarrow w = d_0^i a_0 d_1^i \dots a_n d_{n+1}^i.$$

Intuitively, given a trace c that enables a hazard, strings in π describe all the alternative runs (such as w above) that execute all actions in c and avoid the hazard. The only requirement is that all strings specified by π are observable executions of M ; i.e., for a given FA M , $\pi \subseteq L(M)$. Notice that π being a regular language means that it can be expressed as a regular expression r , and because all strings in π contain c as subword, we have $r = \Sigma_{j,k} r_k^j$ with $r_k^j = r_0^j; a_0; r_1^j; \dots a_k; r_{k+1}^j$ for some finite indexes j and k and regular expressions r_{k+1}^j . For simplicity, we sometimes write r instead of π .

The next definition formally introduces *decorated causes* for an FA M with respect to a hazard e .

Definition 7.1 (Causality for FAs). *Let $M = (S, A, s_0, \rightarrow, F)$ be a FA, e be a regular expression over A , denoting a hazard, and $c \in A^*$. We say that the computation π built on top of c , with $\pi \subseteq L(M)$, is a decorated cause of the hazard e if*

AAC1: *The string c may enable e – $L(c; e) \cap L(M) \neq \emptyset$*

AAC2.1: *If the effect e is not observed then it has not been caused by c –*

$$\forall w \in L(M) \setminus L(A^*; e): (L(w; e) \cap L(M) = \emptyset) \Rightarrow (c \notin \text{sub}(w) \vee w \in \pi).$$

AAC2.2: *Strings of π are safe, i.e., they do not cause the effect e –*

$$\forall w \in \pi: w \notin L(A^*; e) \wedge (L(w; e) \cap L(M) = \emptyset)$$

AAC3: *Minimality –*

*for all $c' \in \text{sub}(c)$ there is no computation π' built on top of c' with $\pi' \subseteq L(M)$, that satisfies **AAC1**–**AAC2.2** with respect to the string c' and the hazard e .*

We call c as above a causal trace and sometimes write $\text{Cause}_c(e, M)$ to denote the corresponding decorated cause π . We let $\text{Causes}(e, M)$ be the union of all $\text{Cause}_c(e, M)$.

Intuitively, **AAC1** identifies a scenario where the string c enables the hazard e in M . Note that **AAC1** entails a “may” semantics of causality, instead of “must”, as c does not always have to lead to e . Catering for the “must” version requires modifying **AAC1** to $L(c; e) \subseteq L(M)$. **AAC2.1** is a necessity condition according to which, if a word w cannot enable e , then either w does not contain the causal trace c (meaning it

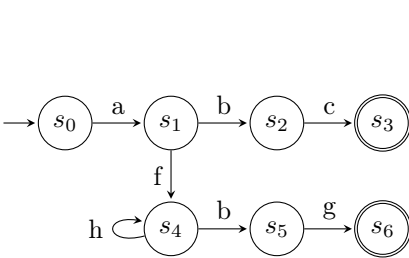
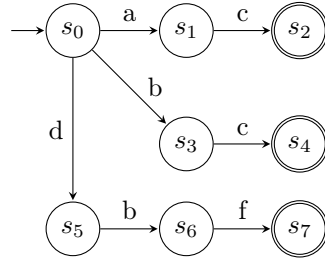
is an execution bringing not to the hazard), or it has been decorated with events that eliminate the possibility of executing the hazard. Note that **AAC2.1** can be equivalently expressed (by modus tollens) as a sufficiency condition stating that a string w enables the hazard e whenever the causal trace is contained in w but it is not decorated with elements causal by their non-occurrence that would avoid the execution of the hazard:

$$\forall w \in L(M) \setminus L(A^*; e) : (c \in \text{sub}(w) \wedge w \not\subseteq \pi) \Rightarrow (L(w; e) \cap L(M) \neq \emptyset)$$

AAC2.2 requires causal traces decorated with events causal by their non-occurrence to avoid the hazard. Furthermore, note that c itself cannot be a safe computation in π , because otherwise **AAC2.2** would contradict **AAC1**. Observe that **AAC2.2** is reminiscent of the traditional counterfactual criterion of Lewis, as it allows us to test the dependence of e on c under certain contingencies encoded, in our case, in terms of non-occurrence of events. We refer to [47] for more insight on the so-called *structural contingencies*. **AAC3** is the minimality condition that requires considering decorated causes entailed by the shortest causal traces c satisfying **AAC1**–**AAC2.2**.

We conclude the section with a few examples intended to clarify certain aspects of the above definition and the differences with the work [25]. To begin with, we illustrate the role played by loops in the decorations of computations.

Example 1. Consider the automaton M_1 in Figure 7.3 and let the hazard be expressed by the regular expression $e = c; A^*$, meaning that we have to avoid executing action c .


 Figure 7.3: Automaton M_1

 Figure 7.4: Automaton M_2

Clearly, the string ab is a possible cause for the hazard. Hence, $Cause_{ab}(e, M_1)$ for this example can be encoded via the regular expression: $a; f; h^*; b; g$. Note that as a result of considering the decorations as regular expressions, all finite repetitions of the loop are conveniently represented with the Kleene star operator. The work in [25] handles loops in the decorations by unfolding the loop only a finite number of times specified a priori, hence, only the string $afh^nb g$ would be describing hazard avoidance, for all $n \leq k$ and some fixed k .

In the second example, we consider the case when there are no possible decorations

to steer a causal trace away from its hazard.

Example 2. Consider the automaton M_2 in Figure 7.4 and let the hazard be as before expressed by the regular expression $e = c; A^*$.

In this example, there are two possible causal traces, namely, a and b . There are no possible decorations for the causal trace a to make it avoid the hazard, whereas, there exists a decoration for the causal trace b with $Cause_b(e, M_2) = d; b; f$. Whenever there are no computations π satisfying Definition 7.1 for e in M w.r.t. a trace c , we say that the hazard e , if enabled by c , is *unavoidable* in M .

In the above two examples, there was no actual difference if we had used c as a hazard instead of the regular expression $c; A^*$. In the next example, we show an FA where the two expressions entail different decorated causes.

Example 3. Consider the automaton M_3 in Figure 7.5 and the hazards $e = c; A^*$ and $e' = c$. For both hazards, ab is the causal trace, but

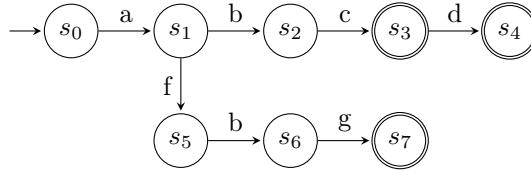


Figure 7.5: Example 3

$$\begin{aligned}
 Cause_{ab}(e, M_3) &= a; f; b; g \\
 Cause_{ab}(e', M_3) &= a; f; b; g + a; b; c; d
 \end{aligned}$$

Observe that the string $abcd$ is considered safe (i.e., avoids the hazard) according to $Cause_{ab}(e', M_3)$ but is not considered safe in $Cause_{ab}(e, M_3)$, where the string $afbg$ is considered safe in both cases. This is different than the usual notion of safety (modeled as in e and thus forbidding any possible continuation after the hazard) as e' allows to overpass the hazard if the system does not stop there. The expression e' asserts that the trace cannot halt with the action c . Accordingly, both $abcd$ and $afbg$ are valid strings that satisfy this condition and thus avoid the hazard e' . On the other hand, the expression e asserts that the action c followed by any possible sequence of actions (i.e., in A^*) constitutes a violation, hence, the action c cannot be observed at any point in execution. Therefore, only $afbg$ is a valid execution that will avoid the hazard e . It is essentially not possible to define properties similar to e with the approach in [25], as they allow jumping over a hazardous state while executing strings in π .

7.5 Computing causes

Given a FA $M = (S, A, i, \longrightarrow, F)$ and an effect specified by a regular expression e on A , we show an algorithm for computing the set $\text{Causes}(e, M)$ using standard operations on automata and graphs. The algorithm first computes the set of loop-free traces that lead to the hazard e . Then, for each one of them, it determines the associated computation satisfying conditions **AAC2.1** – **AAC2.2** in Definition 7.1. The union of all such computations will give a first approximation of the set $\text{Causes}(e, M)$. We will then show below how to obtain precisely the set $\text{Causes}(e, M)$ by requiring the minimality condition **AAC3** in Definition 7.1.

Algorithm 1: Computing Causes

Input: A FA $M = (S, A, i, \longrightarrow, F)$, an effect e .

Output: The set of decorated causes $\text{Causes}(e, M)$.

- (1) Compute the set of traces that lead to e by following the steps:
 - (1.1) For all $s \in S$, construct the FA $P_s = (S, A, s, \longrightarrow, F)$ and compute the following intersection:

$$L(P'_s) = L(P_s) \cap L(e).$$
 - (1.2) Construct the automaton $P = (S, A, i, \longrightarrow, F')$ where

$$F' = \{s \mid L(P'_s) \neq \emptyset\}.$$
 - (1.3) Compute all simple paths from the initial states i and a final state $f \in F$ in P .
 - (1.4) Let CausalTraces be the set of all strings in $L(P)$ labeling the paths computed in (1.3).
 - (2) For all $c = a_0 \dots a_n \in \text{CausalTraces}$, compute $\text{Cause}_c(e, M)$ by :

$$(L(A^*; a_0; A^*; \dots; A^*; a_n; A^*) \setminus \{c\}) \cap (L(M) \setminus (L(A^*; e) \cup L(P)))$$
 - (3) Return the union of all the languages computed in step (2) as $\text{Causes}(e, M)$.
-

Next, we discuss the underlying ideas behind the certain steps of Algorithm 1 and then provide a proof of correctness for the algorithm. We first compute all traces that enable e by constructing in steps (1.1) and (1.2) the automaton P that accepts exactly all traces in M possibly causing the effect e . The only difference between the automata P and M is their set of final states. The procedure for constructing P first involves constructing a set of automata P_s , for all the states s of the automaton M , such that

s is the initial state in P_s and accepts strings of the language of the hazard e . If the intersection of $L(P_s)$ with $L(e)$ is non-empty, then the corresponding state is considered as a final state in the automaton P (step (1.2)). As a result, the strings in $L(P)$ are exactly those strings bringing M to a state where the hazard is activated. For our railway crossing example in Section 7.3 with the hazard given by the regular expression in (7.2), the automaton P would be the one in Figure 7.2 with states ③ and ④ as the only final states.

In step (1.3) we compute *CausalTraces* as the subset of strings accepted by P via a simple path starting from the initial state and ending in a final state. These paths correspond to the set of loop-free traces that lead to the hazard e . While this condition does not guarantee minimality (see discussion below) it already reduces the set of possibly causal traces to a finite set. In general, $L(P)$ will be infinite, if it involves a loop in the automaton.

For each of the above finitely many causal traces, in step (2), we compute the set of associated computations. For a given possibly causal trace c , this is done by subtracting all the traces that enable the effect (i.e., $L(P)$) and all the traces that observe the effect (i.e., $L(A^*;e)$) from $L(M)$ and then take the intersection of the resulting language with the language resulted from c decorated with non-occurrence in all possible ways. Note that the intersection computed in step (2) may be empty, meaning that the hazard e is unavoidable when executing the actions of c . For our running example in Section 7.3, the possible causal traces computed by the algorithm are *CaGo* and *CaGoTa*. Examples of strings in the associated computations are *CaGoTaCcClTcTl* and *CaGoCcClTaTcTl*. Note that the first string avoids the hazard for both possibly causal traces, while the latter is a string that avoids the hazard for *CaGo*.

Finally, the union of the resulting languages in the step (2) of Algorithm 1 is returned as a first approximation of the set of all decorated causes of M for the hazard e . For this set, the following theorem guarantees that conditions **AAC1** – **AAC2.2** hold. However, condition **AAC3** may fail to hold.

Theorem 5. *The computations in $\text{Causes}(e, M)$ returned by Algorithm 1 satisfy conditions **AAC1** – **AAC2.2** by construction.*

Proof. The set $\text{Causes}(e, M)$ returned by Algorithm 1 is obtained as the union of all $\text{Causes}_c(e, M)$ for all $c \in \text{CausalTraces}$. Elements in this set are obtained in step (1.4). These strings are computed based on the language that the automaton P (constructed in step (1.2)) recognizes. By construction, $x \in L(P)$ implies there is $y \in L(e)$ such that $xy \in L(M)$. Hence $L(x;e) \cap L(M) \neq \emptyset$. Since $\text{CausalTraces} \subseteq L(P)$, condition **AAC1** holds.

In order to show that **AAC2.1** holds for some $c \in \text{CausalTraces}$, take a string x accepted by M that is not in $L(A^*;e)$. Assume that $L(x;e) \cap L(M) = \emptyset$. Then $x \notin L(P)$

because otherwise, as we have just seen above, there would exist $y \in L(e)$ such that $xy \in L(M)$. Therefore, $x \in L(M) \setminus (L(A^*;e) \cup L(P))$. Because $CausalTraces \subseteq L(P)$, it follows that $x \neq c$ for any possibly causal trace c . We have now two cases: for every $c \in CausalTraces$ either $c \in sub(x)$ or not. In the latter case **AAC2.1** holds. In the other case $c \in sub(x)$ and thus $x \in L(A^*;a_0;A^*; \dots; A^*;a_n;A^*)$, from which it follows based on step (2) that $x \in Causes_c(e, M)$, and thus **AAC2.1** holds.

It remains to show that **AAC2.2**. For some possible causal trace $c \in CausalTrace$ let $x \in Cause_c(e, M)$. We must show that $x \notin A^*e$ and that $L(x;e) \cap L(M) = \emptyset$. The first part of the conjunction in **AAC2.2** holds because the construction in step (2) $Cause_c(e, M)$ cannot contain strings from $L(A^*;e)$. Similarly, the second part of the conjunction holds because $L(P)$ is subtracted from $L(M)$ in the same step. $\square$

Condition **AAC3** does not necessarily hold for $Cause_c(e, M)$ used by the Algorithm 1. In fact, for possibly causal traces $x, y \in CausalTraces$, if $x \in sub(y)$ then any sub-string of x is also a sub-string of y . In other words, for $a_0 \dots a_n = x \neq y = b_0 \dots b_m$ we have

$$L(A^*;a_0;A^*; \dots; A^*;a_n;A^*) \subseteq L(A^*;b_0;A^*; \dots; A^*;b_m;A^*) \quad (7.3)$$

By step (2) of Algorithm 1 we thus have that $Causes_x(e, M) \subseteq Causes_y(e, M)$. Note that it must be the case that $m > n$ for $x \in sub(y)$. We can therefore easily compute the smallest sets of safe computations by removing from the set $CausalTraces$ all strings y that have another possible causal trace $x \in CausalTraces$ of smaller length as sub-word. In our running example, the trace *CaGo* is a sub-word of the other one *CaGoTa*, and indeed, the computation for *CaGoTa* is included in the computation for *CaGo* as well. Hence, only the causal trace *CaGo* satisfies the minimality condition **AAC3**.

7.6 Experimental evaluation

In this section, we provide an experimental evaluation and assess the applicability of our method. We developed a tool prototype implementing our approach and evaluated the time performance by computing the decorated causes on randomly generated FAs with growing size. The implementation is based on Python and closely follows Algorithm 1. The inputs to our tool are an FA and a regular expression which describes the effect on the given FA. The output of our tool is an automaton that characterizes the set of all decorated causes with respect to the given inputs. In our implementation, we utilized the BRICS automaton library [84] for performing standard automaton operations.

We evaluated our tool in the following experimental setting: we generated random FAs by using the libalf [18] framework. In the process of generating FAs, we fixed the size of the alphabet to 5. We then generated over 1000 FAs with an increasing number

of states and achieved a maximum of 300 states. Figure 7.6 shows an example of an FA with 5 states that was generated randomly by libalf. For each generated FA we also randomly computed an effect for which the decorated causes are determined. We fixed the size of the effect length to 3. All the experiments were conducted on a computer running Ubuntu 20.04.3 with an 8-core 1.8GHz Intel i7-10510U processor and 16 GB RAM.

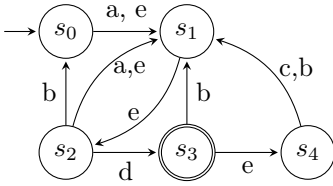


Figure 7.6: Randomly generated FA with 5 states.

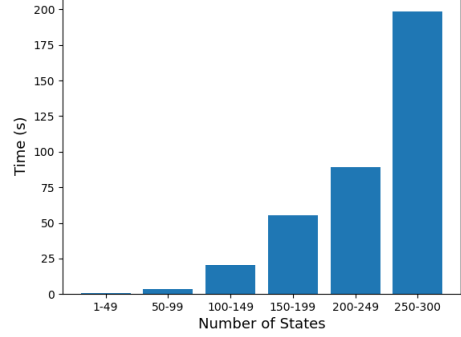


Figure 7.7: Experimental Results

The results of our experiments are displayed in Figure 7.7. We group the randomly generated FAs by their number of states and report the average running times in each group. We only report the times of the experiments in which the decorated causes were not empty. The results indicate that for relatively small FAs with less than 100 states, a result is obtained within 10 seconds. For larger FAs with 250 to 300 states, a result is obtained in 3 minutes on average and within 15 minutes at maximum. We remark that these results are obtained without any attempts to tailor the standard automaton operations to our setting.

	Number of States in the Input FA					
	1-49	50-99	100-149	150-199	200-249	250-300
# States	71	185	266	422	484	560
# Transitions	236	654	997	1565	1862	2177
# Potential Causes	81	328	10476	21932	44750	73318
# (Minimal) Causes	3	8	10	18	10	22

Table 7.1: Average size of obtained decorated causes.

In Table 7.1 we summarize some information on the automata that recognize the decorated causes returned by the algorithm. Depending on the number of states of the automata given as input, we report the average number of states and transitions of the returned automata, the average number of causes, and the average number of minimal causes obtained. As expected, the size of the automata of the output increased

linearly with that of the input. However, the number of potential causal traces computed increases exponentially. That is not the case for the number of minimal causal traces, as it increases only marginally when the size of the input increases. In fact, in the majority of the cases, the number of minimal causes is less than 5, regardless of the size of the given input automaton.

7.7 Extensions

To illustrate the generality of our causal model we briefly discuss possible extensions to consider the addition of tests and assignments.

Adding tests: KAT The set of regular expressions we considered in (7.1) can be extended with a set B of Boolean tests that we assume generated from a finite set At of atoms, meaning that every $b \in B$ is equivalent modulo the equations of the Boolean algebra to a finite disjunction of atoms in At . This way one can model basic programming constructs, like conditionals, loops, guarded actions, and assertions using tests in B and actions in A .

Kozen [67] showed that the above extensions of regular expressions, called KAT (Kleene algebra with tests) expressions, play the same role with regular sets of guarded strings as ordinary regular languages play for regular expressions. Here a *guarded string* is an ordinary string over the alphabet $A \cup At$, such that the symbols in A alternate with the atoms At . Formally, a guarded language is a subset of $(At \times A)^* \times At$.

A *deterministic KAT automaton* recognizing guarded strings [68] is just a deterministic finite automaton $(S, \Sigma, i, \longrightarrow, F)$ with $\Sigma = At \times A$ and $F \subseteq S \times At$. The only differences are thus the transitions that are now labeled by guarded actions (α, a) , and the accepting states, which are now labeled with atoms marking the end of an accepted string. The idea is that an action a is executed only when its guard α (pre-condition) is true, and a string is accepted only in states where the post-condition holds. We say that a guarded string $w \in (At \times A)^* \times At$ is accepted by a KAT automaton M from a state s if either (1) $w = \alpha$ and $(s, \alpha) \in F$, or (2) $w = (\alpha, a)w'$ and there exists $s \xrightarrow{\alpha, a} t$ such that w' is accepted by M from the state t . The *language accepted* by a KAT automaton M is the set $L(M) = \{w \in (At \times A)^* \times At \mid M \text{ accepts } w \text{ from } i\}$.

Our causal model for automata extends naturally to KAT automata by considering hazards e as KAT expressions and causes c as strings in $(At \times A)^*$. Safe computations in M for the hazard e with respect to c are non-empty strings of $L(M)$ satisfying **AAC1** as in Definition 7.1 but with respect to the alphabet $(At \times A)$ instead of A only. Also, the algorithm for computing causes needs no adjustment, but for the way how operations on automata are computed.

Adding assignments: NetKAT NetKAT[4] is a network programming model, which is used for specifying and verifying the packet-processing behavior of software-defined networks. In a nutshell, it is a variation on KAT that considers actions not as abstract elements of an alphabet A but rather as state transformers, like assignments, that are executed when a precondition α is satisfied and modified into a post-condition β .

For a given set of atoms At of a Boolean algebra B , a *deterministic NetKAT automaton* [37] is a deterministic FA $M = (S, \Sigma, i, \longrightarrow, F)$ such that $\sigma = At \times At$ and $F \subseteq S \times (At \times At)$. The transition relation $\longrightarrow$ is thus labeled by pairs of atoms (α, β) and so are the accepting states. The interpretation of these pair of atoms is that they represent pre-conditions and post-conditions of one-step executions.

A string $w \in At \times At \times At^*$ is accepted by M from a state s only when post-conditions match the subsequent pre-condition, meaning that either (1) $w = \alpha\beta$ and $(s, \alpha, \beta) \in F$, or (2) $w = (\alpha\beta)w'$ and there exists $s \xrightarrow{\alpha, \beta} t$ such that $\beta w'$ is accepted by M from the state t . Note that in the last condition is crucial that w' is not the empty string. The *language accepted* by a NetKAT automaton M is the set $L(M) = \{w \in At \times A \times At^* \mid M \text{ accepts } w \text{ from } i\}$.

As for KAT automata, our causal model for automata extends naturally to NetKAT automata too, with hazard represented by NetKAT expressions [37], causes as strings in At^* , and safe computations as strings in $L(M)$ that can be projected into a cause by deleting some atoms and satisfying the rest of the conditions of Definition 7.1.

7.8 Conclusions

In this chapter, we moved the causal model proposed in [25] from labeled transition systems to finite automata to obtain a language-based causal model for safety. The model is in line with the notion of causality described in a logical context in [46] in the sense that a hazard may be observed if and only if it has been caused. Analogously to the alternative worlds of Lewis [75], we also considered decorated causes as alternatives to causes in the sense that they allow executing all actions of a cause interleaved with other actions that guarantee hazard avoidance.

We treated only the case when causes *may* enable a hazard while strings of the decorated causes must avoid it. While it can be interesting to consider a stronger notion of causes as strings c that bring the automaton M to states where the hazard e is inevitable for any of its possible extensions (i.e., by changing **AAC1** to $L(c; e) \subseteq L(M)$), such a change would imply that there would be no causes in our railway system example.

We have also presented an algorithm to compute decorated causes, relying only on basic automata-theoretic operations. The algorithms could be improved, using model checking techniques for marking those states in which a hazard is enabled, and search

techniques to find the decorated causes avoiding marked states. Also, it would be interesting to move from automata back to labeled transition systems but remain in a trace setting, with hazards specified as LTL properties.

Finally, we briefly discussed extensions of our work to KAT and NetKAT automata. More work needs to be done here, both to precisely set the definitions and to show the applicability of the method to, for example, find causes of a hazard in a software-defined network.