

On Action Reversibility in Lifted Planning

Michael Morak¹, Lukáš Chrpá², Jakub Med² and Wolfgang Faber¹

¹University of Klagenfurt, Austria

²Czech Technical University in Prague

Abstract

Action reversibility deals with the problem of whether and under what circumstances we can undo effects of a given action by applying a sequence of other actions for each state of concern. Recent works established the notions of action reversibility in the propositional settings.

The present paper studies the concept of action reversibility in lifted STRIPS settings that represent the environment by first-order predicates and define action schemas defined over those predicates. In particular, we define the notions of lifted reversibility and lifted uniform reversibility. Although undecidable in its most general form, due to the undecidability of first-order logic, we show that several results from the propositional settings can be generalized to the lifted settings.

Keywords

Computational Planning, Reversibility, Reasoning about Actions and Change

1. Introduction

In the field of Automated Planning [1, 2], the aim traditionally is to find sequences of actions, also referred to as plans, that transform an initial state into some goal state. The states are usually represented by means of a logical formalism, while the actions change states by means of effects, and they may also require preconditions in order to be taken. The underlying assumption is usually that everything unaffected by action effects in a state remains as it is, yielding an inherently nonmonotonic setting. A simple, traditional formalism is propositional STRIPS [3], in which states are represented using Boolean variables. In practical settings, a more object-centric formalism is often needed, as actions and properties are often parameterised. For example, named items might be in a location or may be moved from one location to another. In that scenario, rather than having many propositional variables and actions, it is preferable to have one parameterised property and one parameterised action. In planning, this is called a *lifted setting*. The STRIPS formalism in its lifted form [3, 4] represents the environment by first-order logic predicates and action schemas in parameterised form that are instantiated with specific objects.

In the propositional (i.e., non-lifted) case, the question of action reversibility recently received attention [5, 6, 7]. It asks whether the effects of an action can later be undone by a sequence of other actions. However, for lifted planning, this question has not been investigated yet. It is our aim in this paper to set out core definitions in order to capture the two dominant versions of reversibility in use in propositional settings and generalize them to the lifted setting. Lifted action reversibility is domain-specific as it should apply to all sets of objects. For example, we can observe that we can undo the effects of a “load” action by an “unload” action regardless of a truck, package, or location (assuming that we unload the same package from the same truck at the same location). Hence, we can determine the (uniform) reversibility for a lifted action schema once rather than for each specific (grounded) action (the number of instances of a lifted action schema can be exponential with respect to the size of the domain description, including the available objects).

As it turns out, lifting the notions of reversibility from actions to action schemas is not as straightforward as it may seem at first glance. Especially for uniform reversibility, one needs to clarify what uniformity means precisely for sequences of action schemas. It also turns out that some computational tasks are undecidable because the underlying tasks in first-order logic are undecidable. However, we

24th International Workshop on Nonmonotonic Reasoning, July 17–19, 2026, Lisbon, Portugal

✉ michael.morak@aau.at (M. Morak); chrpaluk@cvut.cz (L. Chrpá); jakub.med@cvut.cz (J. Med); wolfgang.faber@aau.at (W. Faber)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

can show that under suitable restrictions, we indeed get decidability and study the computational complexity in these settings. In one setting, we get the same complexity as in the propositional case, which might be unexpected when not looking into the details.

Intuitively, the notion of action reversibility implies a certain level of safety as it is possible to undo the effects of an (undesirable) action or (exogenous) event. In cybersecurity, action reversibility can be useful as a remedy for attacks (by extending the work of Boddy et al. [8]) as well as for assessing the risk of actions. The usefulness of action reversibility for verifying security protocols has been demonstrated in cloud management, confirming the undoability of system operations (or informing the user about permanent and irreversible changes in the system) [9, 10]. The concept of action reversibility is also important in *planning against nature* [11, 12, 13] that deals with generating *safe* plans that are resilient against “acts of nature” (exogenous events that might occur regardless of actions of the agent).

Contributions. The contributions of this paper are as follows:

- We set out the core definitions of lifted reversibility and its uniform version.
- We list several known results from propositional reversibility and show how they can be lifted to our setting.
- We show that lifted reversibility is undecidable in its most general version, but obtain decidability results for the case where the domain of objects is fixed. We show that, in this case, determining reversibility is EXPSPACE-complete.
- We investigate the restrictive case of universal uniform reversibility, that is, where there is a guarantee that an action is always reversible via the same sequence of actions, independent of how and where the original action was applied. We show that this case, surprisingly, is equivalent to the propositional case and hence we inherit the PSPACE-completeness of this task (and NP-completeness in the case of action sequences of polynomially bounded length).

We start by recalling a lifted STRIPS setting in Section 2. The main definitions of reversibility for setting follow at the beginning of Section 3, lifting of some results from the propositional setting is discussed in Section 3.1, and complexity results for decidable problems are discussed in Section 3.2 before concluding in Section 4.

2. Preliminaries

We consider classical planning in a first-order (lifted) setting, as commonly formalized in STRIPS-style planning [3, 1, 14]. We assume the reader is familiar with first-order logic (FO) and only give some core definitions.

Objects, Variables, and Predicates. Let $\mathcal{U}$ be a countably infinite set of objects (the *object universe*) and $\mathcal{V}$ a countably infinite set of variables. A *substitution* σ is a mapping from variables in $\mathcal{V}$ to objects in $\mathcal{U}$. A *schema* $\mathcal{P}$ is a finite set of predicate symbols, where each predicate $p \in \mathcal{P}$ has a fixed arity $\text{arity}(p)$.

Atoms and Literals. An *atom* is an expression of the form $p(t_1, \dots, t_k)$, where $p \in \mathcal{P}$ is a k -ary predicate and each *term* t_i is either an object from $\mathcal{U}$ or a variable from $\mathcal{V}$. An atom is *ground* if it contains no variables.

Structures and First-Order Satisfaction. For a schema $\mathcal{P}$, a finite $\mathcal{P}$ -structure $\mathcal{I}$ is a pair $\mathcal{I} = \langle \mathcal{O}, I \rangle$, where $\mathcal{O} \subset \mathcal{U}$ is a finite set of objects and I is a set of ground atoms over $\mathcal{P}$ with terms from $\mathcal{O}$. We assume FO formulas ϕ are built only from atoms over variables and standard logical connectives and quantifiers. A structure $\mathcal{I} = \langle \mathcal{O}, I \rangle$ is a *model* of (or *satisfies*, denoted $\mathcal{I} \models \phi$) an FO logic sentence ϕ (i.e. a formula without free variables) if the following holds, inductively: $\mathcal{I} \models \exists X. \phi(X)$ holds iff $\phi(o)$ holds for some object $o \in \mathcal{O}$; $\mathcal{I} \models \forall X. \phi(X)$ holds iff $\phi(o)$ holds for all objects $o \in \mathcal{O}$; logical connectives $\vee, \wedge, \neg, \rightarrow$, etc., are interpreted in the standard way. Finally, for ground atoms, $\mathcal{I} \models p(o_1, \dots, o_k)$ iff $p(o_1, \dots, o_k) \in I$, where all o_i are objects from $\mathcal{O}$.

Action Schemas and Ground Actions. For a schema $\mathcal{P}$, an *action schema* a consists of three sets of atoms over $\mathcal{P}$ and terms $\mathcal{V}$: the *preconditions* $\text{pre}(a)$, the *add effects* $\text{add}(a)$, and the *delete effects* $\text{del}(a)$.

The variables appearing in any of these sets are denoted $\text{var}(a)$. A *ground action* a is an action schema where $\text{var}(a) = \emptyset$. Let $\mathcal{O} \subset \mathcal{U}$ be finite and let $\sigma : \text{var}(a) \rightarrow \mathcal{O}$ be a bijective substitution (i.e. mapping different variables to different objects). Applying σ to some action schema a yields a ground action $\sigma(a)$, with σ applied to the variables in each component of the action schema. The set of all ground actions that can be obtained from action schema a w.r.t. a set of objects $\mathcal{O}$ is denoted $\text{ground}(a, \mathcal{O})$. This notation naturally extends to sets of action schemas.

States and Semantics. Let $\mathcal{H}_{\mathcal{O}}$ (the *universe of facts*) be the finite set of ground atoms over schema $\mathcal{P}$ and the finite set of objects $\mathcal{O} \subset \mathcal{U}$. A *state* is a set $s \subseteq \mathcal{H}_{\mathcal{O}}$ of ground atoms. We assume a ground atom in $\mathcal{H}_{\mathcal{O}}$ is false in s iff it is not contained in s . A ground action a is *applicable* in a state s iff $\text{pre}(a) \subseteq s$. Applying an applicable action a in s , denoted $a[s]$, yields a successor state $s' = (s \setminus \text{del}(a)) \cup \text{add}(a)$. Applying a finite sequence of ground actions $\pi = (a_1, \dots, a_n)$ to state s , denoted $\pi[s]$, applies actions iteratively: $s_{i+1} = a_{i+1}[s_i]$, where $s_0 = s$ and $\pi[s] = s_n$.

Planning Domains, Tasks, and Plans. A *planning domain* for schema $\mathcal{P}$ is a tuple $\mathcal{D} = \langle \mathcal{P}, \mathcal{A} \rangle$, where $\mathcal{A}$ is a finite set of action schemas over $\mathcal{P}$. A *planning task* is a tuple $\Pi = \langle \mathcal{D}, \mathcal{O}, s_0, G \rangle$, where $\mathcal{O} \subset \mathcal{U}$ is a finite set of objects, s_0 is the initial state, and G is a set of ground atoms representing the goal condition. A *plan* for a planning task Π is a finite sequence of ground actions π such that $G \subseteq \pi[s_0]$, that is, the state resulting from the application of π satisfies goal condition G .

Propositional Reversibility [5]. Let S be a set of states from some universe of facts $\mathcal{H}_{\mathcal{O}}$. A ground action a from a set $\mathcal{A}$ of ground actions is called *S-reversible* w.r.t. $\mathcal{A}$ iff for every state $s \in S$ where a is applicable it holds that there exists a sequence of actions π such that $\pi[a[s]] = s$. If the sequence of actions π is the same for all states in S , a is called *uniformly S-reversible* w.r.t. $\mathcal{A}$ and π is called the *reverse plan*. The set S above can be specified by a propositional formula ϕ over $\mathcal{H}_{\mathcal{O}}$, yielding the notions of ϕ -reversibility and *uniform ϕ -reversibility*. When $S = 2^{\mathcal{H}_{\mathcal{O}}}$ (or, equivalently, $\phi \equiv \top$), a is called *universally* (uniformly) reversible.

3. Lifted Reversibility

In this section, we extend well-known notions of reversibility for ground planning [15, 16, 5] to the lifted setting. Let ϕ be a FO sentence over the predicates $\mathcal{P}$ from some planning domain $\mathcal{D} = \langle \mathcal{P}, \mathcal{A} \rangle$, using variables from $\mathcal{V}$.

Definition 1. An action schema a is ϕ -reversible w.r.t. planning domain $\mathcal{D}$ if and only if, for every finite $\mathcal{P}$ -structure $\mathcal{I} = \langle \mathcal{O}, I \rangle$ for $\mathcal{D}$ that satisfies ϕ , every ground action $a' \in \text{ground}(a, \mathcal{O})$ is $\{I\}$ -reversible w.r.t. the set $\text{ground}(\mathcal{A}, \mathcal{O})$.

The definition above captures our intuition: any model of the formula ϕ represents a state I and if some ground action a' obtained from action schema a is applicable in I there needs to be a sequence of ground actions obtainable from the action schemas in $\mathcal{A}$ that reverses a' for state I . Hence, lifted reversibility for action schemas is equivalent to (classical, ground) reversibility of all associated ground actions.

However, finding a notion of lifted uniform reversibility is somewhat more involved. In the ground case, uniform reversibility requires that the same sequence of ground actions always revert the initial ground action, independent of the state. In order to capture this notion faithfully in the lifted setting, we must not only specify a sequence of action schemas that can always be applied to revert the originally applied action schema, but we must also specify how those action schemas are grounded during the reversal; otherwise the same ground action could be reverted by two very different sequences of ground actions for two different states, which runs counter to the intuition of uniformity. The following definition lays out this idea. Let ϕ , $\mathcal{V}$, and $\mathcal{D} = \langle \mathcal{P}, \mathcal{A} \rangle$ be as for Definition 1.

Definition 2. An action schema a_0 is uniformly ϕ -reversible w.r.t. planning domain $\mathcal{D}$ iff there exists a set of variables $V \subseteq \mathcal{V}$, a sequence of action schemas $a_1, \dots, a_n$, and substitutions $\sigma_1, \dots, \sigma_n$, with $\sigma_i : \text{vars}(a_i) \rightarrow V$ mapping all the action schema's variables to variables from V , for $0 \leq i \leq n$, such that the following holds:

For each finite set of objects $\mathcal{O}$ and every ground action a obtained from a_0 w.r.t $\mathcal{O}$ there must exist a substitution $\rho : V \rightarrow \mathcal{O}$ with $\rho(\sigma_0(a_0)) = a$ such that for each $\mathcal{P}$ -structure $\langle \mathcal{O}, I \rangle$ that satisfies ϕ we have that if a is $\{I\}$ -reversible w.r.t. ground actions $\text{ground}(\mathcal{A}, \mathcal{O})$, then a must be $\{I\}$ -reversible via the sequence of ground actions $\rho(\sigma_1(a_1)), \dots, \rho(\sigma_n(a_n))$.

We first note that the above notions actually generalize the propositional notions of reversibility, as introduced by Morak et al. [5], which itself generalizes notions of reversibility previously introduced [15, 16]. To see this, it is sufficient to recall that propositional formulas can be viewed as FO sentences without variables and ground actions as action schemas without variables.

The main decision problem we would like to analyze in this paper is the question of *lifted (uniform) ϕ -reversibility*: given a planning domain $\mathcal{D} = \langle \mathcal{P}, \mathcal{A} \rangle$, a FO sentence ϕ over $\mathcal{P}$, and an action schema $a(X_1, \dots, X_n) \in \mathcal{A}$, decide whether $a(X_1, \dots, X_n)$ is ϕ -reversible w.r.t. $\mathcal{D}$. In general, this problem is undecidable by the undecidability of finite satisfiability of FO formulas, per Trakhtenbrot's Theorem:

Proposition 1. *Lifted ϕ -reversibility and lifted uniform ϕ -reversibility are undecidable.*

However, this undecidability stems solely from the fact that ϕ can be an arbitrary FO sentence. For several useful forms of ϕ , we obtain decidability and complexity results. In particular, we will be interested in the case where $\phi = \top$. This special case is called *universal reversibility*, that is, when an action is (uniformly) $\top$ -reversible, we call it *universally (uniformly) reversible*.

Before looking into these cases, however, we will first investigate how several known results for propositional reversibility can be transferred to the lifted setting. These will then aid us in our decidability and complexity investigation.

3.1. Lifting Known Reversibility Results

As already briefly mentioned, lifted reversibility is a strict generalization of propositional reversibility, since ground atoms can be treated in the same way as propositions. We can thus directly apply the results for propositional reversibility to determine whether a ground action is (uniformly) reversible with respect to a set of ground atoms, representing the universe of facts, and a set of ground actions.

Chrupa et al. [6] proved (for ground planning tasks in FDR representation) that a ground action a is universally uniformly reversible if a as well as actions in the reverse plan “touch” only facts that are present in the precondition of a . We restate this result here for the propositional (or ground) STRIPS setting.

Theorem 1 (Chrupa et al. [6]). *Let $\mathcal{D} = \langle \mathcal{F}, \mathcal{A} \rangle$ be a ground planning domain, that is, where $\mathcal{F}$ is a set of ground atoms, and $\mathcal{A}$ is a set of ground actions over $\mathcal{F}$. Then, an action $a_0 \in \mathcal{A}$ is universally uniformly reversible if and only if there exists a reverse plan $\pi = (a_1, \dots, a_n)$ such that for each i with $0 \leq i \leq n$ it holds that $\text{pre}(a_i) \cup \text{add}(a_i) \cup \text{del}(a_i) \subseteq \text{pre}(a_0)$.*

A similar observation can be derived for the lifted setting, since in the reverse action schema sequence, we cannot “touch” predicates that are not part of the precondition of the “to be reversed” action schema.

Theorem 2. *Let $\mathcal{D} = \langle \mathcal{P}, \mathcal{A} \rangle$ be a planning domain. An action schema $a_0 \in \mathcal{A}$ is universally uniformly reversible if and only if there exists a sequence of action schemas $a_1, \dots, a_n$ and substitutions $\sigma_1, \dots, \sigma_n$ with $\sigma_i : \text{var}(a_i) \rightarrow \text{var}(a)$ as in Definition 2 and for all i with $0 \leq i \leq n$ it holds that $\text{pre}(\sigma_i(a_i)) \cup \text{add}(\sigma_i(a_i)) \cup \text{del}(\sigma_i(a_i)) \subseteq \text{pre}(a_0)$, where σ_0 is the identity function.*

Proof. The “if” direction follows immediately from Definition 2. The “only if” direction can be seen as follows: Assume an action schema $a_0 \in \mathcal{A}$ is universally uniformly reversible but the condition in the theorem above does not hold. This means that there is an action schema a_i in the sequence of action schemas that contains some atom that cannot be mapped to the atoms in $\text{pre}(a_0)$. W.l.o.g., let a_i be the first action schema in the sequence where this atom $p(X_1, \dots, X_m)$ occurs. Now, let $\mathcal{O}$ be a finite set of objects and let ρ be a substitution from $X_1, \dots, X_m$ to $\mathcal{O}$. If the atom occurs in $\text{pre}(a_i)$, it is easy to create a state over $\mathcal{P}$ and $\mathcal{O}$ where $\rho(p(X_1, \dots, X_m))$ does not appear, but where a_0 is applicable. However, in

this state, a_0 cannot be reversible, since a_i is not applicable. Now, assume that $p(X_1, \dots, X_m)$ appears in the add or delete effects of a_i . It is again easy to create two states s_1 and s_2 over $\mathcal{P}$ and $\mathcal{O}$, that are the same except that in one $p(X_1, \dots, X_m)$ appears and in the other it does not, with a_0 applicable in both. Again, a_0 cannot be reversible, since applying a_1 through a_i then would lead to a single state s . But this implies that a_0 can only be reversible in at most one of s_1 and s_2 , contradicting our assumption that it is *universally* uniformly reversible. $\square$

In analogy to works concerning, for example, macro-action learning (see e.g. [17]) or action schema acquisition (see e.g. [18]), we can generalize the results of grounded (universal) uniform reversibility into the lifted settings. In particular, the idea stems from tracking the objects in all the involved actions and replacing those objects with variables, such that different objects are replaced with different variables. Such an approach is sound since the action schema and predicate grounding consider all possible substitutions for objects, and ground actions maintain the same structure as their original action schemas. Note that instantiating different variables to the same objects might, in some cases, compromise the integrity of reverse plans—an analogous observation was made for lifted macro-actions [17]. Therefore, we require bijective mappings from variables to objects when grounding (see $\text{ground}(\cdot, \cdot)$ in Section 2). This reflects standard practice, where distinct objects for distinct variables are often enforced either by types (in the typed STRIPS representation), or by construction (e.g. when a *move* action actually needs to change the location). Also, this is not a restriction in general. If different variables are allowed to unify (i.e., map to the same object) in an action schema, an equivalent set of action schemas can be constructed where a copy of the action schema is introduced with a single, fresh variable replacing the variables allowed to unify. In our setting, we get:

Theorem 3. *Let $\mathcal{D} = (\mathcal{P}, \mathcal{A})$ be a planning domain. Let $a \in \mathcal{A}$ be an action schema and $\nu : \text{var}(a) \rightarrow \mathcal{O}$ be a bijective substitution mapping the variables of a into a finite set of objects $\mathcal{O} \subset \mathcal{U}$. It holds that a is universally uniformly reversible if and only if ground action $\nu(a)$ is universally uniformly reversible w.r.t. $\text{ground}(\mathcal{A}, \mathcal{O})$.*

Proof. $\Rightarrow$: It immediately follows from Definition 2.

$\Leftarrow$: Let us assume that $\pi^g = \langle a_1^g, \dots, a_n^g \rangle$ is a reverse plan for $a^g = \nu(a)$. According to Theorem 1, we can observe that variables of all actions in π^g (and a^g) are mapped to $\mathcal{O}$, and that all ground atoms that all actions a_i^g contain are from $\text{pre}(a^g)$. From π^g , we can obtain a sequence of action schemas, i.e., $\pi = \langle a_1, \dots, a_n \rangle$ such that $a_i = \nu^{-1}(a_i^g)$. Note that $a_1, \dots, a_n$ are variants (i.e., the same action up to variable renaming) of action schemas in $\mathcal{A}$. Also, since ν^{-1} is inverse of ν , we have $\text{var}(a_i) \subseteq \text{var}(a)$.

Then we have to show that for every bijective substitution $\xi : \text{var}(a) \rightarrow \mathcal{O}'$, where $\mathcal{O}'$ is a set of objects, $\langle \xi(a_1), \dots, \xi(a_n) \rangle$ is a reverse plan for $\xi(a)$. The bijectivity of ξ is enforced by $\text{ground}(\mathcal{A}, \mathcal{O})$, since only substitutions that map distinct variables to distinct objects are used for grounding (recall Section 2). It can be observed that for a ground atom $f \in \text{pre}(a_i^g)$, there is a corresponding ground atom $f' \in \text{pre}(\xi(a_i))$, i.e., $f' = \xi(\nu^{-1}(f))$. A similar observation can be made for add and delete effects. Hence, we conclude that if $\langle a_1^g, \dots, a_n^g \rangle$ is a reverse plan for $a^g = \nu(a)$, then $\langle \xi(a_1), \dots, \xi(a_n) \rangle$ is a reverse plan for $\xi(a)$. $\square$

3.2. Decidable Lifted Reversibility

We now turn our attention to notions of lifted ϕ -reversibility that avoid the problem of undecidability caused by allowing arbitrary FO sentences. One straightforward scenario where lifted reversibility is decidable is when the set of objects is known, that is, there is a fixed finite set $\mathcal{O} \subset \mathcal{U}$ containing the objects under consideration for a particular lifted planning domain $\mathcal{D} = \langle \mathcal{P}, \mathcal{A} \rangle$. In this case, the problem of deciding reversibility clearly becomes decidable: both the formula ϕ and the action schemas can be grounded w.r.t. $\mathcal{O}$ and then reversibility can be tested for each ground action $a' \in \text{ground}(a, \mathcal{O})$ obtained from some action schema $a \in \mathcal{A}$. In fact, since for a fixed set of objects $\mathcal{O}$, both FO sentence ϕ and the planning domain can be grounded at the cost of an exponential blowup in size, we can obtain the following result by observing that lifted planning can be reduced to lifted reversibility.

Theorem 4. Let $\mathcal{D} = \langle \mathcal{P}, \mathcal{A} \rangle$ be a planning domain, ϕ a FO sentence, and $\mathcal{O} \subset \mathcal{U}$ a finite set of objects. Deciding whether an action schema $a \in \mathcal{A}$ is reversible in all models $\langle \mathcal{O}, I \rangle$ of ϕ is decidable and EXPSPACE-complete.

Proof. To show *membership* in EXPSPACE, first, the grounding ϕ' of ϕ w.r.t. $\mathcal{O}$ is obtained (replacing existential and universal quantification with disjunctions and conjunctions over all objects from $\mathcal{O}$ respectively, whereby ϕ' becomes a propositional formula such that $I \models \phi'$ iff $\langle \mathcal{O}, I \rangle \models \phi$), which is of exponential size. The following procedure runs in EXPSPACE: (1) select an action $a' \in \text{ground}(a, \mathcal{O})$; (2) check that a' is ϕ' -reversible w.r.t. $\text{ground}(\mathcal{A}, \mathcal{O})$, which can be done in PSPACE w.r.t. the size of $\text{ground}(\mathcal{A}, \mathcal{O})$ [5], which is exponential, yielding EXPSPACE overall; (3) repeat until all actions a' have been checked and answer yes if all of them are reversible, otherwise no. To show *hardness*, we can re-use the construction in the PSPACE-hardness proof for propositional reversibility [5] and generalize it to the lifted setting. In particular, take a (lifted) planning task $\Pi = \langle \mathcal{D}, \mathcal{O}, s_0, G \rangle$. Re-using the idea from Morak et al., we can create a planning domain $\mathcal{D}'$ that includes an additional starting action schema a_s and ending action schema a_e , such that a_s transforms an artificial starting state s_{init} into s_0 and a_e transforms any goal state into s_{init} , that is, a_e has G as its precondition. Note that it suffices that a_s and a_e be ground actions. With this, we have that a_s is reversible iff there exists a plan for Π , i.e. goal G can be reached from s_0 , an EXPSPACE-complete problem [4]. $\square$

Surprisingly, the computationally complex task of deciding lifted uniform ϕ -reversibility simplifies dramatically when $\phi = \top$, that is, for universal uniform reversibility, as the following result shows. This is due to the fact that, if any action schema is to be uniformly reversible in all possible finite domains of objects $\mathcal{O} \subset \mathcal{U}$, then this can only be the case when the action schema to be reversed fully specifies the starting state in its precondition.

Theorem 5. Universal uniform reversibility is decidable in PSPACE in general and in NP in case the length of the action schema sequence used to reverse the action in question is limited to a polynomial (in the size of the planning domain and $\mathcal{O}$).

Proof. Let action schema a_0 , sequence $a_1, \dots, a_n$ and substitutions $\sigma_1, \dots, \sigma_n$ as in Theorem 2. By this theorem, all action schemas a_i in this sequence need to be subsets of $\text{pre}(a)$ after variable renaming via σ_i . Hence, let $\mathcal{O} = \{c_X \mid X \in \text{var}(a)\}$ be a set of objects, representing Skolem terms that “freeze” the variables in a , and let ν be a substitution mapping variable X to object c_X . By Theorem 3, the substitution ν can be used to obtain a ground action $\nu(a_0)$ from a_0 and a reverse plan for $\nu(a_0)$ w.r.t. $\text{ground}(\mathcal{A}, \mathcal{O})$, namely, the sequence of ground actions $\nu(\sigma_1(a_1)), \dots, \nu(\sigma_n(a_n))$. The proof of Theorem 1 then states that $a_1, \dots, a_n$ are indeed a sequence of action schemas that reverse a_0 in all possible ways it can ever be applied in a planning task. But this then establishes equivalence between lifted universal uniform reversibility and propositional universal uniform reversibility, the latter a problem known to be PSPACE-complete in general (as we need to solve an underlying planning task [6]), and NP-complete in case where n is polynomial in the size of the input [5]. $\square$

4. Conclusions

In this paper, we laid out basic definitions of lifted reversibility, explored the connections to the propositional case, showed that the problem is undecidable in general, due to the undecidability of first-order logic, and offered several decidability and complexity results for the cases of reversibility for a fixed set of objects and the case of universal uniform reversibility.

In the future, we would like to explore in-depth which interesting decidable cases of lifted ϕ -reversibility exist (i.e. what form the formula ϕ takes in useful settings), explore how this problem can be solved in practice and implement solving systems to decide lifted reversibility.

Acknowledgements. This research was partially supported by the Czech Science Foundation (project no. 24-13337L) and Austrian Science Fund (FWF) under grant numbers 10.55776/PIN8782623 and 10.55776/COE12, under the joint bilateral project, and by the Grant Agency of Czech Technical University (project no. SGS24/140/OHK3/3T/13).

Declaration on Generative AI. The authors did not employ any generative AI tools or services.

References

- [1] M. Ghallab, D. S. Nau, P. Traverso, *Automated Planning—Theory and Practice*, Elsevier, 2004.
- [2] M. Ghallab, D. S. Nau, P. Traverso, *Automated Planning and Acting*, Cambridge University Press, 2016. URL: <http://www.cambridge.org/de/academic/subjects/computer-science/artificial-intelligence-and-natural-language-processing/automated-planning-and-acting?format=HB>.
- [3] R. Fikes, N. J. Nilsson, STRIPS: A new approach to the application of theorem proving to problem solving, *Artif. Intell.* 2 (1971) 189–208. doi:10.1016/0004-3702(71)90010-5.
- [4] K. Erol, D. S. Nau, V. S. Subrahmanian, Complexity, decidability and undecidability results for domain-independent planning, *Artif. Intell.* 76 (1995) 75–88. URL: [https://doi.org/10.1016/0004-3702\(94\)00080-K](https://doi.org/10.1016/0004-3702(94)00080-K). doi:10.1016/0004-3702(94)00080-K.
- [5] M. Morak, L. Chrpá, W. Faber, D. Fiser, On the reversibility of actions in planning, in: *Proc. KR*, 2020, pp. 652–661.
- [6] L. Chrpá, W. Faber, M. Morak, Universal and uniform action reversibility, in: *Proc. KR*, 2021, pp. 651–654.
- [7] J. Med, L. Chrpá, M. Morak, W. Faber, Weak and strong reversibility of non-deterministic actions: Universality and uniformity, in: S. Bernardini, C. Muise (Eds.), *Proceedings of the Thirty-Fourth International Conference on Automated Planning and Scheduling, ICAPS 2024, Banff, Alberta, Canada, June 1-6, 2024*, AAAI Press, 2024, pp. 369–377. URL: <https://doi.org/10.1609/icaps.v34i1.31496>. doi:10.1609/ICAPS.V34I1.31496.
- [8] M. S. Boddy, J. Gohde, T. Haigh, S. A. Harp, Course of Action Generation for Cyber Security Using Classical Planning, in: *Proceedings of the Fifteenth International Conference on Automated Planning and Scheduling*, June 5-10 2005, Monterey, California, USA, AAAI, 2005, pp. 12–21. URL: <http://www.aaai.org/Library/ICAPS/2005/icaps05-002.php>.
- [9] I. Weber, H. Wada, A. D. Fekete, A. Liu, L. Bass, Automatic Undo for Cloud Management via AI Planning, in: *Proceedings of the Eighth USENIX Conference on Hot Topics in System Dependability, HotDep’12*, USENIX Association, USA, 2012, p. 10. URL: <https://www.usenix.org/conference/hotdep12/workshop-program/presentation/weber>, hollywood, CA.
- [10] I. Weber, H. Wada, A. D. Fekete, A. Liu, L. Bass, Supporting Undoability in Systems Operations, in: *27th Large Installation System Administration Conference (LISA 13)*, USENIX Association, Washington, D.C., 2013, pp. 75–88. URL: <https://www.usenix.org/conference/lisa13/technical-sessions/presentation/weber>.
- [11] L. Chrpá, E. Karpas, On verifying linear execution strategies in planning against nature, in: *ICAPS*, AAAI Press, 2024, pp. 86–94.
- [12] L. Chrpá, E. Karpas, On generating robust plans and linear execution strategies in planning against nature, in: *ICAPS*, AAAI Press, 2025, pp. 169–177.
- [13] L. Chrpá, M. Pilát, J. Med, On eventual applicability of plans in dynamic environments with cyclic phenomena, in: *KR*, 2021, pp. 184–193.
- [14] M. Helmert, Concise finite-domain representations for PDDL planning tasks, *Artif. Intell.* 173 (2009) 503–535. URL: <https://doi.org/10.1016/j.artint.2008.10.013>. doi:10.1016/j.artint.2008.10.013.
- [15] T. Eiter, E. Erdem, W. Faber, Undoing the effects of action sequences, *J. Applied Logic* 6 (2008) 380–415. doi:10.1016/j.jal.2007.05.002.
- [16] J. Daum, Á. Torralba, J. Hoffmann, P. Haslum, I. Weber, Practical undoability checking via

contingent planning, in: Proc. ICAPS, 2016, pp. 106–114. URL: <http://www.aaai.org/ocs/index.php/ICAPS/ICAPS16/paper/view/13091>.

- [17] L. Chrupa, Generation of macro-operators via investigation of action dependencies in plans, *Knowledge Engineering Review* 25 (2010) 281–297.
- [18] B. Juba, H. S. Le, R. Stern, Safe learning of lifted action models, in: *Proceedings of the 18th International Conference on Principles of Knowledge Representation and Reasoning, KR 2021*, Online event, November 3-12, 2021, 2021, pp. 379–389.