

Complete Game Logic with Sabotage

Noah Abou El Wafa

noah.abouelwafa@kit.edu

Karlsruhe Institute of Technology

Karlsruhe, Germany

Carnegie Mellon University

Pittsburgh, USA

André Platzer

platzer@kit.edu

Karlsruhe Institute of Technology

Karlsruhe, Germany

Carnegie Mellon University

Pittsburgh, USA

ABSTRACT

We introduce *Sabotage Game Logic* (GL_S), a simple and natural extension of Parikh’s Game Logic with a single additional primitive, which allows players to lay traps for the opponent to avoid. GL_S can be used to model infinite sabotage games, in which players can change the rules during game play. In contrast to Game Logic, which is strictly less expressive, GL_S is exactly as expressive as the modal μ -calculus. This reveals a close connection between the entangled nested recursion inherent in modal fixpoint logics and adversarial dynamic rule changes characteristic for sabotage games.

Additionally we present a natural Hilbert-style proof calculus for GL_S and prove completeness. The completeness of an extension of Parikh’s calculus for Game Logic follows.

CCS CONCEPTS

• **Theory of computation** $\rightarrow$ **Logic; Proof theory; Modal and temporal logics.**

KEYWORDS

game logic, μ -calculus, proof theory, completeness, expressiveness, sabotage games

1 INTRODUCTION

Games such as the Ehrenfeucht-Fraïssé game are invaluable tools in the study of logic [19] and some deep results about logic can be proved with the help of games [1, 27]. Logic can even be given meaning via games in the form of game-theoretical semantics [28].

On the contrary, logical methods are frequently used to study games [15]. Logic and games meet most directly in logics specifically designed for the study of games, such as Game Logic (GL) due to Parikh [36], which allows reasoning about the existence of winning strategies in a game. This requires giving exact meaning to general games, a nontrivial task, except for games that are limited to a fixed number of rounds. Nested alternating least and greatest fixpoints can provide the correct denotational semantics for games, when they are used to reflect the alternating responsibilities of the respective players at their decision points in the dynamic games [36].

Fixpoints generally play an important role in logic, for example in modal fixpoint logics such as the modal μ -calculus (L_μ). L_μ is where logic, games and fixpoints begin to converge. In fact Game Logic can be expressed in the modal μ -calculus using alternating fixpoint formulas to directly capture the semantics of alternating game play. However this first encounter is imperfect. After 25 years it was shown that GL is in fact strictly less expressive than L_μ [8].

The purpose of this paper is to remedy this situation by unifying the three fundamental concepts of logic, games, and fixpoints in a small and natural extension of Game Logic, which is shown to be *equivalent* to the fixpoint logic L_μ and to have a complete proof calculus. This identification of fixpoints with games eliminates the difference between interactive game play and alternating fixpoints. The key insight behind this paper is that because Game Logic can already express sufficient adversarial dynamics to express the alternating fixpoints of L_μ and is merely lacking a suitable way of referring to fixpoints by their respective fixpoint variables, this can be alleviated in a parsimonious and purely game-theoretic way. This is done in *Sabotage Game Logic* (GL_S), a new extension of GL. In sabotage game logic reference can be expressed, not through the unstructured use of fixpoint variables as is done in the modal μ -calculus, but by using a simple game operator $\sim a$ that *changes* the rules of subsequent game play. Playing the game $\sim a$ has the effect that the game a is reserved exclusively for the Angel player in the future. This can be used to change the rules of a game dynamically according to rules that are explicit in the original game. This simple and natural mechanism of imperative game play is expressively equivalent to the functional mechanism of unstructured nested named recursion with the fixpoint variables in the alternating fixpoints of L_μ . The role the sabotage $\sim a$ plays in establishing the equiexpressiveness reveals an interesting connection between games with sabotage and the nesting of fixpoints in the modal μ -calculus which have previously been studied separately.

Formulas of the modal μ -calculus are frequently easiest to understand through their corresponding validity or model-checking parity games [35]. This is complicated by the unstructured goto-like action a fixpoint variable induces. Sabotage game logic avoids this problem, as GL_S formulas describe two-player games built up from simple connectives and, instead of fixpoint variables, players only need to consider the previously committed acts of sabotage, making sabotage game logic a very intuitive logic with very high expressive power. By the equivalence of L_μ and GL_S , many desirable properties of the modal μ -calculus, such as decidability and small model property, can be transferred to sabotage game logic for free. Moreover completeness of an axiomatization of GL_S can be obtained through the translation. This is in contrast to the original axiomatization for game logic, for which completeness is still open after four decades [30]. GL_S promises to be a useful tool for understanding GL. This is evidenced by the completeness of an extension of Parikh’s axiomatization for GL obtained from the complete proof calculus for GL_S . To the best of our knowledge this is the only complete proof calculus for game logic to date. The embedding from sabotage game logic to the modal μ -calculus also suggests the possibility that the same property can be expressed

significantly more concisely in sabotage game logic than in the modal μ -calculus.

In summary, the contributions of this paper are threefold. Firstly, GL_s , a new minimal, natural, concise and intuitive extension of game logic well-suited to logically studying sabotage games, is introduced. Secondly, it is shown that GL_s is expressively exactly as powerful as the modal μ -calculus and, consequently, many desirable logical properties of L_μ transfer to GL_s . Thirdly, a sound proof calculus for GL_s is presented, proved complete and completeness is transferred to obtain a complete extension of Parikh's GL calculus.

Outline. The required preliminaries are recalled in Section 2. Section 3 introduces sabotage game logic (GL_s) and another extension recursive game logic (RGL) of game logic that will play an intermediary role in translating between the modal μ -calculus and sabotage game logic. In essence RGL adds completely recursive games to game logic by naively importing the notion of fixpoint variable reference from modal μ -calculus. Section 4 briefly recalls the definitions of some modal fixpoint logics such as the modal μ -calculus formally. In Section 5 the expressive power of fragments of recursive game logic is compared to modal fixpoint logics. Subsequently, proof calculi for sabotage game logic and the fragment right-linear game logic of recursive game logic are introduced and completeness for right-linear game logic is proved in Section 6. Finally in Section 7 equiexpressiveness of sabotage game logic and the modal μ -calculus is established and completeness of the proof calculus for GL_s is proved. Completeness of an extension of Parikh's calculus for GL is obtained as a corollary. All proofs are in Appendix D.

Related Work. Sabotage games have been considered to model algorithms under adversarial conditions and in learning [26, 44]. Previous work on using modal logic for the sabotage game using Sabotage Modal Logic (SML) [6, 44] differs from sabotage game logic. Unlike in GL_s , where obstruction is modelled as changing the meaning of the game described syntactically, in SML sabotage is described as changing the structure of interpretation. The sabotage μ -calculus was investigated for modelling infinite sabotage games [41]. In contrast to GL_s of the model changing nature satisfiability problem for SML is undecidable and lacks the finite model property [32].

For examples of applications of Game Logic see [37, 38]. The relation of the games, game logic, fixpoints and modal μ -calculus has been considered in [8, 10, 17, 20, 25, 29, 35, 36, 42]. Equiexpressiveness and relative completeness of game logic and modal μ -calculus in the first-order case was shown in [45]. The modal μ -calculus and its relation to model checking is well-studied [11, 12, 21, 39]. Completeness for game logic was conjectured [36]. A completeness proof based on a cut-free complete calculus for L_μ [2] was suggested [23]. It was recently shown not to work [30].

2 PRELIMINARIES

Effectivity Function. An *effectivity function* [22] is a monotone function $w : \mathcal{P}(X) \rightarrow \mathcal{P}(X)$. It will be used for the denotational semantics of a game, where $w(Y)$ denotes the set of all states (winning region) from which a given player can win into the region $Y \subseteq X$. These are naturally monotone, as any point in the winning region for a set $A \subseteq B$ is also in the winning region for the larger

goal B . Let $\mathcal{W}(X)$ be the set of such effectivity functions ordered by point-wise inclusion, i.e. $w \subseteq u$ if $w(A) \subseteq u(A)$ for all $A \subseteq X$.

- Definition 2.1.* (1) Given a set $A \subseteq X$ the intersection effectivity function is $A?(B) = A \cap B$ and the constant effectivity function $\bar{A}(B) = A$.
- (2) For an effectivity function $w \in \mathcal{W}(X)$ its dual is $w^d(A) = X \setminus w(X \setminus A)$.
- (3) An effectivity function $w \in \mathcal{W}(X)$ is *relational* if there is a relation $R \subseteq X \times X$ such that

$$w(A) = R \circ A = \{r : \exists s \in A \ rRs\}$$

- (4) For $w \in \mathcal{W}(X)$ let

$$\mu A.w(A) = \bigcap \{A \in \mathcal{P}(X) : w(A) \subseteq A\}$$

- (5) For a function $\Delta : \mathcal{W}(X) \rightarrow \mathcal{W}(X)$ let

$$\mu w.\Delta(w) = \bigcap \{p \in \mathcal{W}(X) : \Delta(w) \subseteq p\}$$

be the least fixpoint as usual.

As usual $\mu A.w(A)$ and $\mu q.\Delta(q)$ are the least fixpoints of w and Δ respectively, provided Δ is monotone [4]. In the sequel it will be necessary to work with both fixpoints of monotone functions on a power set and fixpoints of monotone functions on the set of monotone functions on a power set. Under some conditions the latter can be viewed pointwise as the former.

LEMMA 2.2. *Suppose $\Delta : \mathcal{W}(X) \rightarrow \mathcal{W}(X)$ is monotone and $\Delta(w)(A) = \Delta(w(\bar{A}))(\bar{A})$ for all $w \in \mathcal{W}(X)$ and all $A \in \mathcal{P}(X)$.*

- (1) $(\mu u.\Delta(u))(A) = \mu B.(\Delta(\bar{B}))(A)$
- (2) $(\nu u.\Delta(u))(A) = \nu B.(\Delta(\bar{B}))(A)$

See proof on page 16.

Neighbourhood and Kripke Structures. Both game logic and the modal μ -calculus can be interpreted over arbitrary coalgebraic structures [16]. While the modal μ -calculus is commonly interpreted over Kripke structures, game logic was originally interpreted over the more general class of neighbourhood models [34]. The results in this paper hold for both classes of models equally.

Definition 2.3. A *monotone neighbourhood structure* is a triple $\mathcal{N} = (|N|, v_N, \rho_N)$ consisting of a set $|N|$ and functions

$$v_N : \mathbb{A} \rightarrow \mathcal{P}(|N|) \quad \rho_N : \mathbb{G} \rightarrow \mathcal{W}(|N|)$$

assigning a valuation to atomic propositions from $\mathbb{G}$ and an effectivity function to atomic games from $\mathbb{G}$. The structure $\mathcal{N}$ is a *Kripke structure* if each $\rho_N(a)$ is relational.

3 EXTENSIONS OF GAME LOGIC

Throughout the paper fix a countably infinite set $\mathbb{A}$ of propositional constants, a countably infinite set $\mathbb{V}$ of fixpoint variables and a countably infinite set $\mathbb{G}$ of atomic games.

3.1 Sabotage Game Logic

Sabotage game logic (GL_s) is an extension of game logic defined by adding the atomic games $\sim a$. Formulas and games of GL_s are given

by the following grammar:

$$\begin{aligned}\varphi &::= P \mid \neg\varphi \mid \varphi \vee \psi \mid \langle\alpha\rangle\varphi \\ \alpha &::= a \mid \sim a \mid ?\varphi \mid \alpha \cup \beta \mid \alpha; \beta \mid \alpha^* \mid \alpha^d\end{aligned}$$

The formula $\langle\alpha\rangle\varphi$ expresses that player Angel has a winning strategy in the game α to reach one of the states in which φ is true. The test game $?\varphi$ is lost prematurely by Angel unless the formula φ is true in the current state. The choice game $\alpha \cup \beta$ allows Angel to choose between playing α or β . To play the sequential game $\alpha; \beta$ is to play β after α unless one of the players has already lost while playing α . The repetition game α^* allows Angel to decide after each round of α whether she wants to stop playing or repeat α unless one of the two players has already lost. The dual operator d switches the roles of the players. Any decision or test within the scope of a dual operator is not taken by Angel, but her opponent Demon.

The additional primitive $\sim a$ is the trap-setting game. When $\sim a$ is played, Angel lays a trap at a . It has the effect that should Demon, anytime in the subsequent play, try to play a (by reaching a^d), he falls into the trap and loses the game prematurely. The trap stays in effect throughout the formula until it is turned into a trap for Demon by playing $\sim a^d$ according to the rules of the game. However once a trap has been set for an atomic game it can only change hands, but will not be played normally again.

Viewed differently, if a player plays $\sim a$ the player claims the atomic game a for herself. The opponent is not allowed to play it and forfeits the game by trying. Formally $\sim a$ means that the game a belongs to Angel until the next time the game $\sim a^d$ is played. Playing $\sim a^d$ dually means the game a belongs to Demon until it returns to Angel. The effect of the claim is that the rules for playing a and a^d in the future *change* as in Table 1.

Table 1: Effect of Rule Changes

Owner of a	Rules for a	Rules for a^d
Neither $\emptyset$	a played normally	a^d played normally
Angel $\diamond$	a skipped	Angel wins a^d
Demon $\square$	Demon wins a	a^d skipped

Abbreviations and Conventions. As in game logic define the dual game connectives for demons choice, test, or repetition. That is let $\alpha \cap \beta$ abbreviate $(\alpha^d \cup \beta^d)^d$. This leaves the choice of whether to play α or β to Angel's opponent. Analogously let $!\varphi$ stand for the tests Demon needs to pass $(?\varphi)^d$ and $\alpha^\times$ for demonic repetition $(\alpha^d)^d$. The propositional connectives $\wedge, \rightarrow, \leftrightarrow$ and $\top, \perp$ are defined as usual. Sequential composition $;$ binds stronger than choice $\cup$, $\cap$ and conjunction and disjunction bind stronger than implication.

Infinite Plays. A subtlety in Game Logic is that it is possible for the two players to play indefinitely. For example Angel could potentially choose to repeat a $(?\top)^*$; $?\perp$ game indefinitely, instead of choosing to stop the repetition and thereby losing prematurely due to the rules of the $?$ game. This is desirable as it allows modelling infinite games. Intuitively the semantics will be defined so that the

player who causes the game to be infinite (by repeating a subgame infinitely often that is not contained in another subgame that is repeated infinitely often) loses.

Game Logic. Syntactically Parikh's Game logic (GL) [36] is the fragment of GL_S without games of the form $\sim a$ and $\sim a^d$. The semantics of sabotage game logic will be defined so that they agree with the usual semantics. Hence sabotage game logic is a genuine extension of game logic.

Examples. To illustrate the role of the trap-setting game consider the following two games

$$(\sim a \cap \sim a^d); a \quad (\sim a \cup \sim a^d); a; !\perp$$

Demon has a winning strategy in the first game. He simply chooses to play the game $\sim a^d$ first, thereby laying a trap for Angel at a . Subsequently Angel will need to try to play a . However the trap Demon set causes her to lose at this point. In the second game in contrast, Angel has a winning strategy. If Angel chooses first to play $\sim a$, she sets a trap for Demon at a . This also means that Angel can skip the following game of a and go straight to playing $!\perp$. At this point Demon loses the game, since he can not pass the test $!\perp$.

Example: The Poison Game in GL_S . The poison game is a graph game introduced [18] to capture the notion of a perfect kernel. To begin the game Angel picks a vertex on the graph. The players alternate to choose adjacent vertices, tracing a path through the graph. However Demon poisons every vertex he chooses to move to. Angel loses if she chooses to go to a poisoned vertex, but Demon is immune to the poison he leaves behind. It is Angel's objective to survive and Demon's to poison Angel. The value of the game lies in the fact that Angel has a winning strategy in the poison game played on a progressively and outwardly finite directed graph iff the graph has a local kernel [18].

Various approaches to capture the poison game in modal logic have been suggested [6, 9]. However the previously considered logics have undecidable satisfiability problems. In GL_S , instead of viewing the poisoning as a model change, which is difficult to capture by a well-behaved logic [3], poisoning can be understood as a describable rule change.

Formally we list the vertices of the graph as $a_1, \dots, a_n$ and view them simultaneously as atomic games to be interpreted as moving to the vertex of the same name. Now consider the games $\alpha_\diamond \equiv a_1 \cup \dots \cup a_n$ and $\alpha_\square \equiv ((a_1; \sim a_1) \cup \dots \cup (a_n; \sim a_n))^d$. In $\alpha_\diamond$ Angel chooses which vertex to go to next. (The semantics enforce that she loses should she try to go to a non-adjacent vertex.) Analogously $\alpha_\square$ allows Demon to make a move, which is followed by him poisoning the vertex he goes to. The two players now play the game

$$\varphi \equiv \langle(\alpha_\square; \alpha_\diamond)^\times\rangle \top.$$

interpreted over an n -element graph (W, E) viewed as a Kripke frame $\mathcal{K}$ with states $|\mathcal{K}| = W$ and the reachability relations

$$xR_{w_i}y \quad \text{if} \quad (x, y) \in E \text{ and } y = w_i.$$

In words, the formula $\langle a_i \rangle \top$ holds in exactly those states from which vertex w_i is reachable along an E -edge. The formula φ interpreted in this way describes the poison game, as φ is satisfiable in this structure (so there is a state from which Angel can win the

game $(\alpha_\square; \alpha_\diamond)^\times$ iff Angel has a winning strategy in the original poison game.

3.2 Recursive Game Logic

We introduce a second extension, *recursive game logic* (RGL), of game logic that allows arbitrary recursive games. This large increase in expressive power allows RGL to serve as the essential technical intermediary connecting GL_s to the modal μ -calculus. Recursive Game Logic (RGL) is defined by the grammar

$$\begin{aligned}\varphi &::= P \mid \neg\varphi \mid \varphi \vee \psi \mid \langle\alpha\rangle\varphi \\ \alpha &::= a \mid x \mid \alpha^d \mid ?\varphi \mid \alpha \cup \beta \mid \alpha; \beta \mid rx.\alpha\end{aligned}$$

where $P \in \mathbb{A}$, $a \in \mathbb{G}$, $x \in \mathbb{V}$. The additional restriction is that variables x are not free in games $?\varphi$ and can only be bound by $rx.\alpha$ if x appears only in the scope of an even number of d in α . Syntactically the only difference between recursive game logic and game logic is that repetition games α^* have been replaced by named subgames of the form $rx.\alpha$ and games x to recursively call these named subgames have been added. The ordinary repetition game α^* can be viewed as an abbreviation for $rx.(\alpha; x \cup ?\top)$

Intuitively a recursive game $rx.\alpha$ is played just like α until the variable x is encountered. In this case the game is interrupted and the players will begin another subgame of $rx.\alpha$. At some stage the players may finish playing this subgame. They then continue in the state they reached to play the original game that was interrupted previously.

The abbreviations for the usual propositional symbols and the demonic connectives are defined just as in GL_s. Additionally the dual version of a named subgame is defined as $\lambda x.\alpha$ to mean $(\lambda x.\alpha^d)^d$. This game is played similarly to $rx.\varphi$. The only difference is which of the players is held responsible if the game is played infinitely long. If the largest subgame repeated infinitely often during a play in a game of the form $rx.\alpha$, then Angel loses the game. If the largest such game is of the form $\lambda x.\alpha$, Demon loses.

3.2.1 Examples. An example of a fully recursive game is $rx.(\top \cup a; x; b)$. Angel can win this game relative to the winning condition that some formula φ holds in the final state exactly if there is some n such that after n rounds of playing a she can win a game of n consecutive rounds of b into φ . This game can not be described in game logic, which lacks facilities to retain the number of games b that still have to be played after Angel chooses the left side in her choice once.

3.3 Semantics of Game Logics

A denotational semantics for recursive game logic and sabotage game logic can be defined in a simple and compositional way. Superficially both semantics are quite different from the usual semantics of game logic. However it will be shown that for GL formulas the semantics of RGL and GL_s agree with the original semantics of GL.

3.3.1 Semantics of Recursive Game Logic. Because recursive game logic contains games of the form $rx.a; (x \cup c); b$ unlike in game logic the semantics of such a game can no longer be defined as the fixpoint of a function between power sets, as the play of c that will

take place after Angel chooses to play b for the first time must be taken into account.

Formally the semantics of recursive game logic with respect to both a monotone neighbourhood structure and a valuation. A *valuation* is a function $I : \mathbb{V} \rightarrow \mathcal{W}(|\mathcal{N}|)$ assigning an interpretation to every variable x . Given a valuation I a variable $x \in \mathbb{V}$ and a $w \in \mathcal{W}(|\mathcal{N}|)$ let $I[x \mapsto w]$ denote the valuation that agrees with I , except that $I[x \mapsto w](x) = w$.

Definition 3.1. For any monotone neighbourhood structure $\mathcal{N}$ and any valuation I define the *semantics* $\mathcal{N}[\![\varphi]\!]^I \in \mathcal{P}(|\mathcal{N}|)$ and $\mathcal{N}[\![\alpha]\!]^I \in \mathcal{W}(|\mathcal{N}|)$ by mutual induction on RGL formulas φ and RGL games α :

$$\begin{aligned}\mathcal{N}[\![P]\!]^I &= v_{\mathcal{N}}(P) & \mathcal{N}[\![\neg\varphi]\!]^I &= |\mathcal{N}| \setminus \mathcal{N}[\![\varphi]\!]^I \\ \mathcal{N}[\![\varphi \vee \psi]\!]^I &= \mathcal{N}[\![\varphi]\!]^I \cup \mathcal{N}[\![\psi]\!]^I & \mathcal{N}[\![\langle\alpha\rangle\varphi]\!]^I &= \mathcal{N}[\![\alpha]\!]^I(\mathcal{N}[\![\varphi]\!]^I) \\ \mathcal{N}[\![a]\!]^I &= \rho_{\mathcal{N}}(a) & \mathcal{N}[\![x]\!]^I &= I(x) \\ \mathcal{N}[\![?\varphi]\!]^I(A) &= \mathcal{N}[\![\varphi]\!]^I \cap A & \mathcal{N}[\![\alpha; \beta]\!]^I &= \mathcal{N}[\![\alpha]\!]^I \circ \mathcal{N}[\![\beta]\!]^I \\ \mathcal{N}[\![rx.\alpha]\!]^I &= \mu u. \mathcal{N}[\![\alpha]\!]^{I[x \mapsto u]} & \mathcal{N}[\![\alpha \cup \beta]\!]^I &= \mathcal{N}[\![\alpha]\!]^I \cup \mathcal{N}[\![\beta]\!]^I \\ \mathcal{N}[\![\alpha^d]\!]^I &= (\mathcal{N}[\![\alpha]\!]^I)^d\end{aligned}$$

For closed formulas the superscript I is dropped. As usual the notation $\mathcal{N} \models \varphi$ means that $\mathcal{N}[\![\varphi]\!]^I = |\mathcal{N}|$ for all valuations I . Moreover write $\models \varphi$ if $\mathcal{N} \models \varphi$ for all *monotone neighbourhood structures* $\mathcal{N}$, and $\models_K \varphi$ if $\mathcal{N} \models \varphi$ for all *Kripke structures* $\mathcal{N}$.

The semantics of named subgames are well-defined and the meaning of games $rx.\alpha$ can be seen to be extremal fixpoints by monotonicity of the function $u \mapsto \mathcal{N}[\![\alpha]\!]^{I[x \mapsto u]}$. The proof is postponed to Lemma 3.6.

3.3.2 Semantics of Sabotage Game Logic. The semantics of a game of sabotage game logic depends on the traps that players have already laid in the run of the game so far. To keep track of these, games and formulas of sabotage game logic must be evaluated in a context. A *context* is a function $c : \mathbb{G} \rightarrow \{\emptyset, \diamond, \square\}$ indicating which player has previously laid a trap (see Table 1). All contexts are assumed to have finite support, that is $c(a) = \emptyset$ for all but finitely many a . Let C be the set of all contexts and let $c_0(a) = \emptyset$ for all a be the constant context without any traps laid. For any set $U \subseteq |\mathcal{N}| \times C$ and any context $c \in C$ let $U \upharpoonright c = \{\omega : (\omega, c) \in U\}$ be the projection on $|\mathcal{N}|$.

To interpret an atomic game a it is necessary to consider the context in which it is played. If one of the players has already laid a trap the normal rules no longer apply. Formally any effectivity function $w : \mathbb{G} \rightarrow \mathcal{W}(|\mathcal{N}|)$ is *lifted* to $\widehat{w}(a) \in \mathcal{W}(|\mathcal{N}| \times C)$ by defining $(\omega, c) \in \widehat{w}(a)(U)$ iff

- (1) $c(a) = \emptyset$ and $\omega \in w(a)(U \upharpoonright c)$ or
- (2) $c(a) = \diamond$ and $(\omega, c) \in U$

Hence Angel can win the game a from a position ω in context c , into the set U if a is not claimed (i.e. $c(a) = \emptyset$) and additionally she can win a game of a played according to the usual rules into $U \upharpoonright c$. If a belongs to Angel (i.e. $c(a) = \square$), she can also win if the current state ω and context c are already in U . However if a belongs to Demon (i.e. $c(a) = \diamond$), Angel has already lost. This formalizes the effects of the rule changes as described in Table 1. For any context c write $\bar{c}$ for the dual context where angelic traps become demonic

traps and vice versa:

$$\bar{c}(a) = \begin{cases} \emptyset & \text{if } c(a) = \emptyset \\ \sqcap & \text{if } c(a) = \diamond \\ \diamond & \text{if } c(a) = \sqcap \end{cases}$$

For a set $A \subseteq W \times C$ write $A^C = \{(\omega, c) : (\omega, \bar{c}) \notin A\}$. For a function $w \in \mathcal{W}(W \times C)$ define the *sabotage dual*

$$w^D(A) = w(A^C)^C.$$

The sabotage dual extends the notion of the ordinary dual to sabotage games. In particular $(\omega, c) \in \widehat{w}(a)^D(U)$ iff

- (1) $c(a) = \emptyset$ and $\omega \in w(a)^D(\{v : (v, c) \in U\})$ or
- (2) $c(a) = \diamond$ or
- (3) $c(a) = \sqcap$ and $(\omega, c) \in U$

The semantics of formulas and games of recursive game logic with respect to a monotone neighbourhood structure is defined by mutual induction.

Definition 3.2. For GL_s formulas and any monotone neighbourhood structure $\mathcal{N}$ the *semantics* is defined as a set $\mathcal{N}[\![\varphi]\!]_s \in \mathcal{P}(|\mathcal{N}| \times C)$ as follows

$$\begin{aligned} \mathcal{N}[\![P]\!]_s &= v_N(P) \times C & \mathcal{N}[\![\langle \alpha \rangle \varphi]\!]_s &= \mathcal{N}[\![\alpha]\!]_s(\mathcal{N}[\![\varphi]\!]_s) \\ \mathcal{N}[\![\neg \varphi]\!]_s &= \mathcal{N}[\![\varphi]\!]_s^C & \mathcal{N}[\![\varphi \vee \psi]\!]_s &= \mathcal{N}[\![\varphi]\!]_s \cup \mathcal{N}[\![\psi]\!]_s \end{aligned}$$

The *semantics* of games α is defined as an effectivity function $\mathcal{N}[\![\alpha]\!]_s \in \mathcal{W}(|\mathcal{N}| \times C)$ by

$$\begin{aligned} \mathcal{N}[\![a]\!]_s &= \widehat{\rho_N}(a) \\ \mathcal{N}[\![\sim a]\!]_s(A) &= \{(\omega, c) : (\omega, c_a^{\diamond}) \in A\} \\ \mathcal{N}[\![? \varphi]\!]_s(A) &= \mathcal{N}[\![\varphi]\!]_s \cap A \\ \mathcal{N}[\![\alpha \cup \beta]\!]_s &= \mathcal{N}[\![\alpha]\!]_s \cup \mathcal{N}[\![\beta]\!]_s \\ \mathcal{N}[\![\alpha^*]\!]_s(A) &= \mu B. A \cup \mathcal{N}[\![\alpha]\!]_s(B) \\ \mathcal{N}[\![\alpha; \beta]\!]_s &= \mathcal{N}[\![\alpha]\!]_s \circ \mathcal{N}[\![\beta]\!]_s \\ \mathcal{N}[\![\alpha^d]\!]_s &= \mathcal{N}[\![\alpha]\!]_s^D \end{aligned}$$

The semantics of $\sim a$ illustrates the role of the context. Playing the game $\sim a$ changes the context and assigns $\diamond$ the game a to keep track of the trap Angel has laid there.

Unlike for sabotage modal logic [5] the semantics is not defined in terms of a changing model. Instead the state space is enlarged to contain the states of the structure and independently keep track of the acts of model change or, here, traps laid. The definition is similar in spirit to the modified semantics for the sabotage μ -calculus [5]. Unlike in the definition of the modal μ -calculus augmented with sabotage [41] the traps set will persist throughout multiple repetitions of a game α^* instead of being reset without cause.

3.3.3 Dual Normal Form. For some proofs it is important that negation is only applied to propositional atoms and the duality operator is only applied to atomic games and free variables. Formulas and games that satisfy this condition are said to be in *normal form*.

Definition 3.3. By mutual recursion on GL_s formulas and games define the *syntactic complement* $\bar{\varphi}$ of a sabotage game logic formula and the *syntactic dual* α^d of a sabotage game logic game as follows:

$$\bar{P} = \neg P \quad \overline{\neg P} = P$$

$$\begin{aligned} \overline{\varphi \wedge \psi} &= \bar{\varphi} \vee \bar{\psi} & \overline{\varphi \vee \psi} &= \bar{\varphi} \wedge \bar{\psi} \\ \overline{\langle \alpha \rangle \varphi} &= \langle \alpha^d \rangle \bar{\varphi} & (\alpha; \beta)^d &= \alpha^d; \beta^d \\ (a)^d &= a^d & (a^d)^d &= a \\ (\sim a)^d &= \sim a^d & (\sim a^d)^d &= \sim a \\ (? \varphi)^d &= !\varphi & (!\varphi)^d &= ?\varphi \\ (\alpha \cup \beta)^d &= \alpha^d \cap \beta^d & (\alpha \cap \beta)^d &= \alpha^d \cup \beta^d \\ (\alpha^*)^d &= (\alpha^d)^\times & (\alpha^\times)^d &= (\alpha^d)^* \end{aligned}$$

The syntactic complement and dual semantically correspond to set complements and dual functions:

LEMMA 3.4 (DUALITY). $\mathcal{N}[\![\bar{\varphi}]\!]_s = \mathcal{N}[\![\varphi]\!]_s^C$ for any GL_s formula φ and $\mathcal{N}[\![\alpha^d]\!]_s = \mathcal{N}[\![\alpha]\!]_s^D$ for any game α .

See proof on page 16.

By inductively replacing negations in $\neg \varphi$ by $\bar{\varphi}$ and duals α^d by α^d any GL_s formula and any game can be transformed into an equivalent formula or game in normal form.

Definition 3.3 can be easily modified for formulas of recursive game logic by defining

$$(x)^d = x \quad (x^d)^d = x \quad (rx.\alpha)^d = !x.(\alpha)^d \quad (!x.\alpha)^d = rx.(\alpha)^d$$

Again the syntactic complement and dual semantically correspond to set complements and dual functions.

LEMMA 3.5. $\mathcal{N}[\![\bar{\varphi}]\!]^I = |\mathcal{N}| \setminus \mathcal{N}[\![\varphi]\!]^I$ for any RGL formula φ and $\mathcal{N}[\![\alpha^d]\!]^I = (\mathcal{N}[\![\alpha]\!]^I)^d$ for any RGL game α .

See proof on page 16.

As was the case for GL_s , through inductively replacing negation $\neg \varphi$ and duality α^d by their syntactic versions $\bar{\varphi}$ and $(\alpha)^d$ every RGL formula and every RGL game can be turned into an equivalent formula or game in normal form respectively, since bound variables appear only within the scope of an even number of d operators.

In the sequel we assume all formulas and games of GL_s and RGL are in normal form. Because every formula is equivalent to one in normal form this does not restrict the generality.

3.3.4 Semantic Compatibility. Using the normal form, observe that the semantics of named games are indeed extremal fixpoint. This relies on the monotonicity:

LEMMA 3.6. If $rx.\alpha$ is a game of recursive game logic, then $F : q \mapsto \mathcal{N}[\![\alpha]\!]^{I[x \mapsto q]}$ is monotone.

PROOF. The game α is equivalent to a formula in normal form, in which x^d does not appear, since it must be in the scope of an even number of d operators in α . Monotonicity of F follows by induction on such a α . $\square$

Unlike for game logic the semantics of repetition games above was not defined as the fixpoint of a set operator. However Definition 3.1 agrees with the definition of the standard semantics for game logic for all GL games.

LEMMA 3.7. If α is a GL game then

$$\mathcal{N}[\![\alpha^*]\!]^I(A) = \mu B. A \cup \mathcal{N}[\![\alpha]\!]^I(B)$$

For a proof see Lemma B.2 in Appendix B.

Although Definitions 3.1 and 3.2 are quite different the semantics of recursive game logic and of sabotage game logic agree on their common fragment, GL.

PROPOSITION 3.8. *If φ is a formula and α a game of GL then*

$$\mathcal{N}[\![\varphi]\!] = \mathcal{N}[\![\varphi]\!]_s \upharpoonright c_0 \quad \mathcal{N}[\![\alpha]\!](U \upharpoonright c_0) = \mathcal{N}[\![\alpha]\!]_s(U) \upharpoonright c_0.$$

PROOF. Use Lemma 3.7 for the case of fixpoints in a mutual induction on formulas and games of game logic. $\square$

Write $\mathcal{N} \models \varphi$ to abbreviate $\mathcal{N}[\![\varphi]\!]_s \supseteq |\mathcal{N}| \times \{c_0\}$. This captures the intended semantics of φ as being evaluated when no traps have been set initially, by requiring the formula to hold in every state in the special context c_0 in which no game has been claimed. Moreover write $\models \varphi$ if $\mathcal{N} \models \varphi$ for all *monotone neighbourhood structures* $\mathcal{N}$ and $\models_K \varphi$ if $\mathcal{N} \models \varphi$ for all *Kripke structures* $\mathcal{N}$. This overloading of notation for game logic formulas is justified by Proposition 3.8.

4 MODAL FIXPOINT LOGICS

The Modal μ -Calculus. This section recalls two modal fixpoint logics. Of particular interest is the *modal μ -calculus* (L_μ) [11], because of its desirable logical properties. It has decidable satisfiability and model checking problems, the finite model property and comes with a natural complete proof calculus. The syntax of L_μ is given by the following grammar:

$$\varphi ::= P \mid \neg P \mid x \mid \langle a \rangle \varphi \mid [a] \varphi \mid \varphi \vee \psi \mid \varphi \wedge \psi \mid \mu x. \varphi \mid \nu x. \varphi$$

for $P \in \mathbb{A}$, $a \in \mathbb{G}$, $x \in \mathbb{V}$. The modal μ -calculus extends basic (multi)-modal logic with fixpoint operators $\mu x. \varphi$ and $\nu x. \varphi$. These denote the least and greatest fixpoints of φ in the sense that $\mu x. \varphi(x)$ is equivalent to $\varphi(\mu x. \varphi)$. The syntax enforces that fixpoint variables x can appear only positively in order to ensure that the semantics of fixpoint operators $\mu x. \varphi$ denote the desired extremal fixpoints.

Fixpoint Logic with Chop. An interesting extension of the modal μ -calculus is fixpoint logic with chop [33]. Although it lacks some of the nice properties of modal μ -calculus, its high expressiveness is useful to establish a close correspondence with the game logics from the previous section via a natural translation. The syntax of fixpoint logic with chop (FLC) [33] is given by the following grammar:

$$\varphi ::= \text{id} \mid P \mid \neg P \mid x \mid \langle a \rangle \varphi \mid [a] \varphi \mid \varphi \vee \psi \mid \varphi \wedge \psi \mid \varphi \circ \psi \mid \mu x. \varphi \mid \nu x. \varphi$$

for $P \in \mathbb{A}$, $a \in \mathbb{G}$, $x \in \mathbb{V}$. Fixpoint logic with chop is conceptually close to the modal μ -calculus. However fixpoint variables do not range over predicates (elements of $\mathcal{P}(|\mathcal{N}|)$) anymore, but over predicate transformers (monotone functions in $\mathcal{W}(|\mathcal{N}|)$) instead. Consequently formulas denote predicate transformers which admit a natural notion of concatenation $\circ$ and identity transformation id . As in the modal μ -calculus the definition syntactically restricts to positive appearances of x , to ensure the well-definedness of the semantics of the fixpoint operator.

Semantics of Fixpoint Logic with Chop. The semantics of fixpoint logic with chop is defined with respect to monotone neighbourhood structures and a valuation $I : \mathbb{V} \rightarrow \mathcal{W}(|\mathcal{N}|)$. By structural induction on formulas φ define the set $\mathcal{N}[\![\varphi]\!]^I \in \mathcal{W}(|\mathcal{N}|)$

$$\begin{aligned} \mathcal{N}[\![\text{id}]\!]^I &= \text{id} & \mathcal{N}[\![x]\!]^I &= I(x) \\ \mathcal{N}[\![\neg P]\!]^I &= |\mathcal{N}| \setminus v_{\mathcal{N}}(P) & \mathcal{N}[\![P]\!]^I &= v_{\mathcal{N}}(P) \\ \mathcal{N}[\![\varphi \vee \psi]\!]^I &= \mathcal{N}[\![\varphi]\!]^I \cup \mathcal{N}[\![\psi]\!]^I & \mathcal{N}[\![\varphi \wedge \psi]\!]^I &= \mathcal{N}[\![\varphi]\!]^I \cap \mathcal{N}[\![\psi]\!]^I \\ \mathcal{N}[\![\langle a \rangle \varphi]\!]^I &= \rho_{\mathcal{N}}(a) \circ \mathcal{N}[\![\varphi]\!]^I & \mathcal{N}[\![[a] \varphi]\!]^I &= \rho_{\mathcal{N}}(a)^d \circ \mathcal{N}[\![\varphi]\!]^I \\ \mathcal{N}[\![\mu x. \varphi]\!]^I &= \mu q. \mathcal{N}[\![\varphi]\!]^{I[x \mapsto q]} & \mathcal{N}[\![\nu x. \varphi]\!]^I &= \nu q. \mathcal{N}[\![\varphi]\!]^{I[x \mapsto q]} \\ \mathcal{N}[\![\varphi \circ \psi]\!]^I &= \mathcal{N}[\![\varphi]\!]^I \circ \mathcal{N}[\![\psi]\!]^I \end{aligned}$$

The semantics of μ and ν formulas denotes extremal fixpoints, since the semantics are monotone:

LEMMA 4.1. *The function $F : q \mapsto \mathcal{N}[\![\varphi]\!]^{I[x \mapsto q]}$ is monotone.*

PROOF. Monotonicity holds, because by definition of the syntax $\neg x$ can not appear in a formula. $\square$

The semantics of a formula of fixpoint logic with chop is defined as a *monotone* function. To assign a truth value the function can be evaluated at $\emptyset$. Write $\mathcal{N} \models \varphi$ if $\mathcal{N}[\![\varphi]\!]^I(\emptyset) = |\mathcal{N}|$ for all I . By monotonicity of the semantics this ensures that $\mathcal{N} \models \varphi$ iff $\mathcal{N}[\![\varphi]\!]^I(U) = |\mathcal{N}|$ for all I and all $U \subseteq |\mathcal{N}|$. Moreover write $\models \varphi$ if $\mathcal{N} \models \varphi$ for all *monotone neighbourhood structures* $\mathcal{N}$ and $\models_K \varphi$ if $\mathcal{N} \models \varphi$ for all *Kripke structures* $\mathcal{N}$.

The semantics of L_μ formulas with respect to the FLC semantics coincide with the usual semantics of modal μ -calculus. (See Lemma B.1 in Appendix B.)

Negation in Fixpoint Logic with Chop. The negation of a formula of fixpoint logic with chop is defined syntactically as usual:

$$\begin{aligned} \overline{\neg P} &= P & \overline{\langle a \rangle \varphi} &= [a] \overline{\varphi} & \overline{\varphi \vee \psi} &= \overline{\varphi} \wedge \overline{\psi} & \overline{\mu x. \varphi} &= \nu x. \overline{\varphi} \\ \overline{[a] \varphi} &= \langle a \rangle \overline{\varphi} & \overline{\varphi \wedge \psi} &= \overline{\varphi} \vee \overline{\psi} & \overline{\nu x. \varphi} &= \mu x. \overline{\varphi} \\ \overline{x} &= x \end{aligned}$$

The syntactic definition of negation corresponds semantically to complementation:

LEMMA 4.2. $\mathcal{N}[\![\overline{\varphi}]\!]^I(\emptyset) = |\mathcal{N}| \setminus \mathcal{N}[\![\varphi]\!]^{I^c}(\emptyset)$ for all L_μ formulas φ , where I^c is the pointwise complement of I .

PROOF. By a straightforward induction on formulas. $\square$

The Modal $$ -Calculus.* Restricting the fixpoints in FLC to structured ones as they appear in game logic yields a logic we call the modal $*$ -calculus, which is the exact modal fixpoint logic equivalent of game logic. The syntax of *modal $*$ -calculus* (L_*) is defined as

$$\varphi ::= \text{id} \mid P \mid \neg \varphi \mid \varphi \vee \psi \mid \langle a \rangle \varphi \mid \varphi \circ \psi \mid \varphi^*$$

This can be viewed as a fragment of FLC by interpreting $\neg \varphi$ as $\overline{\varphi}$ and φ^* as an abbreviation for $\mu x. \text{id} \vee \varphi \circ x$ where x is some fresh variable. Disjunctions $\varphi \vee \psi$ do not strictly need to be added as primitives, since they are definable in L_* as $(\varphi \circ \perp)^* \circ \psi$.

5 EXPRESSIVENESS

The semantics of game logic and modal μ -calculus are in many ways similar and game logic can express large parts of the modal μ -calculus. In particular it spans the entire fixpoint alternation hierarchy of the modal μ -calculus [7]. Nevertheless, game logic is less expressive than modal μ -calculus [8]. This section introduces natural translations to show that at the level of fixpoint logic with chop and recursive game logic, modal fixpoint logics and game logics can be identified completely. As a consequence the exact modal fixpoint logic corresponding to game logic and the fragment of recursive game logic corresponding to modal μ -calculus can be determined.

A formula φ of RGL is *well-named* if it does not bind the same variable twice and no variable appears both free and bound. Every formula is equivalent to a well-named formula by bound renaming.

5.1 Equiexpressiveness of FLC and RGL

For formulas φ, ψ of fixpoint logic with chop denote by $\varphi \frac{\psi}{x}$ the formula obtained from φ by syntactically replacing all free occurrences of x in φ by ψ . The same notation is used for syntactic substitution in formulas and games of game logics.

5.1.1 Translation from fixpoint logic with chop to recursive game logic. To express any formula φ of FLC equivalently in recursive game logic a translation $\varphi^\#$ of any FLC formula φ in the form of a RGL game is defined inductively:

$$\begin{aligned} (\text{id})^\# &= ?\top & (x)^\# &= x \\ (P)^\# &= ?P; !\perp & (\neg P)^\# &= ?\neg P; !\perp \\ (\varphi \vee \psi)^\# &= \varphi^\# \cup \psi^\# & (\varphi \wedge \psi)^\# &= \varphi^\# \cap \psi^\# \\ (\langle a \rangle \varphi)^\# &= a; \varphi^\# & ([a] \varphi)^\# &= a^d; \varphi^\# \\ (\mu x. \varphi)^\# &= rx. \varphi^\# & (vx. \varphi)^\# &= \imath x. \varphi^\# \\ (\varphi \circ \psi)^\# &= \varphi^\#; \psi^\# \end{aligned}$$

The RGL formula corresponding to φ is $\varphi^\# \equiv \langle \varphi^\# \rangle \perp$.

PROPOSITION 5.1 (CORRECT $\#$). *For any FLC formula φ the translation satisfies $\mathcal{N}[\varphi]^\dagger = \mathcal{N}[\varphi^\#]^\dagger$. Hence $\mathcal{N}[\varphi]^\dagger(\emptyset) = \mathcal{N}[\varphi^\#]^\dagger$.*

PROOF. By structural induction on φ . $\square$

5.1.2 Translation from recursive game logic to fixpoint logic with chop. Conversely any formula of recursive game logic can be expressed equivalently in fixpoint logic with chop. To do this, fix two fresh variables u, v . Intuitively the purpose of these variables is to mark the end of the game, so that it can later be replaced by its continuation. The difference between the two variables is that v marks games that end in fixpoint variables, while u marks the end of all other games. This distinction will only be important later when considering a particular subclass of formulas.

For any rGL formula φ and any RGL game α define FLC formulas φ^b and α^b by structural induction:

$$\begin{aligned} (P)^b &= P & (\neg P)^b &= \neg P \\ (\varphi \vee \psi)^b &= \varphi^b \vee \psi^b & (\varphi \wedge \psi)^b &= \varphi^b \wedge \psi^b \\ (\langle \alpha \rangle \varphi)^b &= \alpha^b \frac{\varphi^b}{u, v} & (a)^b &= \langle a \rangle u & (a^d)^b &= [a] u \\ (? \psi)^b &= \psi^b \wedge u & (! \psi)^b &= \neg(\psi^b) \vee u \\ (\alpha \cup \beta)^b &= \alpha^b \vee \beta^b & (\alpha \cap \beta)^b &= \alpha^b \wedge \beta^b \\ (x)^b &= x \circ v & (\alpha; \beta)^b &= \alpha^b \frac{\beta^b}{u, v} \\ (rx. \alpha)^b &= (\mu x. \alpha^b \frac{\text{id}}{u, v}) \circ u & (\imath x. \alpha)^b &= (vx. \alpha^b \frac{\text{id}}{u, v}) \circ u \end{aligned}$$

Note that $\varphi \frac{\psi}{u, v}$ denotes the formula obtained by replacing all appearances of u and v in φ by ψ . This is different from successive substitution $\varphi \frac{\psi}{u} \frac{\psi}{v}$. The substitutions here are always admissible, that is no fixpoint construct captures a free variable. In fact none of the variables that are substituted (u, v) even appears in the context of a fixpoint in the translation. (By choice of u, v , the variables do not appear in the original formula.)

PROPOSITION 5.2 (CORRECT b). *For any well-named RGL formula φ and any well-named RGL game α*

- (1) $\mathcal{N}[\varphi]^\dagger = \mathcal{N}[\varphi^b]^\dagger(A)$ for any $A \subseteq |\mathcal{N}|$
- (2) $\mathcal{N}[\alpha]^\dagger \circ w = \mathcal{N}[\alpha^b]^\dagger[u, v \mapsto w]$

See proof on page 16.

COROLLARY 5.3 (ROUNDTRIP). $\models \varphi \leftrightarrow \varphi^b$ and $\models \psi \leftrightarrow \psi^b$ for all well-named FLC formulas φ and all well-named RGL formulas ψ

5.2 The Modal μ -Calculus as a Game Logic

In this section the precise extension of game logic that corresponds to modal μ -calculus is identified. Although the lack of fixpoint variables of the modal μ -calculus in game logic was remedied by introducing named subgames, the modal μ -calculus can only be understood as a game logic, in which games are played in a tail-recursive way. This is required to capture the regularity of the modal μ -calculus in the context of recursive game logic.

A game α of recursive game logic is *right-linear* if no subgame $\beta; \gamma$ of α has a free fixpoint variable in β . A formula φ of recursive game logic is *right-linear* if all its subgames are right-linear. The fragment of recursive game logic consisting only of right-linear formulas and games is called *right-linear game logic* (rGL).

The translation $\#$ transforms formulas of the modal μ -calculus to right-linear game logic, since the only sequential games $\alpha; \beta$ introduced in the translation of α are of the form $a, a^d, ?P$ or $?P$. We can modify b to ensure that it carries out a converse translation. For any rGL formula φ and any rGL game α define L_μ formulas φ^d and α^d by structural induction. The definition of α^d is identical to the definition of α^b , except for the following cases

$$(x)^d = x \quad (\alpha; \beta)^d = \alpha^d \frac{\beta^d}{u} \quad (rx. \alpha)^d = \mu x. \alpha^d$$

Note that $\varphi^{\downarrow}$ is a modal μ -calculus formula, as it does not mention $\circ$. This is a generalization of the translation employed in [23].

PROPOSITION 5.4 (CORRECT $\downarrow$). *For any well-named closed right-linear RGL formula φ the L_μ formula $\varphi^{\downarrow}$ satisfies $\mathcal{N}[\![\varphi^{\downarrow}]\!] = \mathcal{N}[\![\varphi^b]\!]$. Moreover $\mathcal{N}[\![\varphi]\!] = \mathcal{N}[\![\varphi^b]\!](\emptyset)$.*

See proof on page 17.

COROLLARY 5.5 (EQUIEXPRESSIVENESS FOR L_μ). *Right-linear game logic and modal μ -calculus are equiexpressive.*

PROOF. It is easy to see that $\varphi^\#$ is a formula of right-linear game logic if φ is a formula in modal μ -calculus. This shows that right-linear game logic is at least as expressive as modal μ -calculus. The converse follows with Proposition 5.4. $\square$

5.3 Game Logic as a Fixpoint Logic

Recall that the modal $*$ -calculus is the fragment of the fixpoint logic with chop, which contains no fixpoints except in the form φ^* . Because the fixpoint structure in the modal $*$ -calculus mirrors the structure in game logic, the translations between RGL and FLC also show the equiexpressiveness of the modal $*$ -calculus and GL. This identifies the exact modal fixpoint logic corresponding to Parikh's original game logic.

The technical notion of formula separability will be used for the proof. A formula φ of the modal μ -calculus is *separable* if it contains fixpoints only in the forms $\mu x.(\psi \vee \rho)$ and $\nu x.(\psi \wedge \rho)$ where ρ does not mention x and ψ has no variable other than x free. Let L_s denote the set of separable formulas of the modal μ -calculus.

- LEMMA 5.6.** (1) *If φ is a L_* formula, then $\varphi^\#$ is a GL formula.*
 (2) *If φ is a well-named GL formula, then $\varphi^{\downarrow}$ is a L_s formula.*
 (3) *Any L_s formula is equivalent to a L_* formula.*

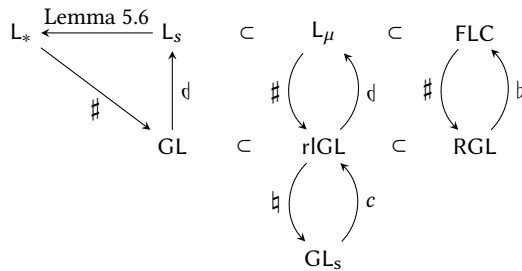
See proof on page 17.

THEOREM 5.7 (EQUIEXPRESSIVENESS FOR GL). *Game logic (GL), the modal $*$ -calculus (L_*) and the separable fragment of the modal μ -calculus (L_s) are equiexpressive.*

The equivalence between the separable fragment of modal μ -calculus and game logic has been shown [14, Theorem 3.3.10]. Theorem 5.7 adds to this equivalence the modal $*$ -calculus. It is still open whether game logic is equivalent to the two variable fragment. With the above this can be reduced to the question of whether every L_μ formula is expressible in L_* .

5.4 Summary of Expressiveness

The following illustrates the relations of the logics considered:



The equivalences for sabotage game logic will be proved in the sequel. All inclusions in the illustration are strict. Game Logic is strictly less expressive than the modal μ -calculus [8], hence it is also less expressive than right-linear game logic. Fixpoint logic with chop and thus recursive game logic are strictly more expressive than the modal μ -calculus and right-linear game logic [33].

In the definition of RGL and sabotage game logic the games could contain tests of arbitrary formulas. For example $\langle ?(\langle a \rangle P) \rangle P$ is a well-formed Game Logic formula. This *rich-test* version is in contrast to the *poor-test* version in which only tests of literals (i.e. formulas P and $\neg P$) are allowed.

COROLLARY 5.8 (TESTS). *The poor-test versions of game logic, right-linear game logic and recursive game logic are equiexpressive with their respective rich-test versions.*

PROOF. This can be seen by translating into the corresponding fragment of fixpoint logic with chop via $\flat$ or $\downarrow$, since the backward translation via $\#$ yields an equivalent (Corollary 5.3) poor-test formula, since the translation $\#$ only introduces tests of literals. $\square$

6 PROOF CALCULI

This section introduces natural proof calculi for right-linear game logic and sabotage game logic. We also recall Kozen's original calculus for the modal μ -calculus and its monotone variant, since completeness for the above game logics is obtained from completeness for the modal μ -calculus.

6.1 Proof Calculi for the Modal μ -Calculus

Because we consider the modal μ -calculus interpreted generally over neighbourhood structures, we recall the *monotone modal μ -calculus* mL_μ [24], the restriction of Kozen's calculus for the modal μ -calculus for neighbourhood structures. The monotone modal μ -calculus consists of all propositional tautologies together with all instances of the following axioms:

$$\begin{aligned} (\text{fp}) \quad & \varphi \frac{\mu x. \varphi}{x} \rightarrow \mu x. \varphi \\ (\alpha) \quad & \sigma x. \varphi \leftrightarrow \sigma y. \varphi \frac{y}{x} \quad (y \text{ fresh}, \sigma \in \{\mu, \nu\}) \end{aligned}$$

The rules of the proof calculus are:

$$\begin{aligned} (\text{MP}) \quad & \frac{\varphi \quad \varphi \rightarrow \psi}{\psi} & (\mu) \quad \frac{\varphi \frac{\psi}{x} \rightarrow \psi}{\mu x. \varphi \rightarrow \psi} & (\text{Ma}) \quad \frac{\varphi \rightarrow \psi}{\langle a \rangle \varphi \rightarrow \langle a \rangle \psi} \end{aligned}$$

We write $mL_\mu \vdash \varphi$ if there is a Hilbert style proof of φ in the monotone modal μ -calculus. Note that this is a subset of Kozen's proof calculus for the modal μ -calculus [31]. Adding the following axioms yields the full Kozen calculus.

$$\begin{aligned} ([\wedge]) \quad & [a] \varphi \wedge [a] \psi \leftrightarrow [a] (\varphi \wedge \psi) \\ (\text{K}) \quad & [a] (\varphi \vee \psi) \rightarrow \langle a \rangle \varphi \vee [a] \psi \\ ([\top]) \quad & [a] \top \end{aligned}$$

We write $L_\mu \vdash \varphi$ if there is a Hilbert-style proof in this calculus of the formula φ . Both proof calculi are complete:

PROPOSITION 6.1 (ENQVIST, SEIFAN, VENEMA [24]). *The monotone modal μ -calculus is sound and complete with respect to monotone neighbourhood structures. That is $mL_\mu \vdash \varphi$ iff $\models \varphi$ for L_μ formulas φ*

PROPOSITION 6.2 (WALUKIEWICZ [46]). *Kozen's calculus is sound and complete with respect to Kripke structures. That is $L_\mu \vdash \varphi$ iff $\models_K \varphi$ for L_μ formulas φ .*

6.2 Proof Calculi for Game Logics

Parikh proposed a similar Hilbert-style proof calculus for game logic [36], which can be extended to right-linear game logic. Consider the axioms

$$\begin{aligned} (\neg) \quad & \langle \neg \rangle \varphi \leftrightarrow \neg \langle \neg \rangle \neg \varphi \\ (?) \quad & \langle ? \rangle \varphi \leftrightarrow \varphi \wedge \psi \\ (\cup) \quad & \langle \alpha \cup \beta \rangle \varphi \leftrightarrow \langle \alpha \rangle \varphi \vee \langle \beta \rangle \varphi \\ (;) \quad & \langle \alpha ; \beta \rangle \varphi \leftrightarrow \langle \alpha \rangle \langle \beta \rangle \varphi \\ (fp_G) \quad & \langle \alpha \frac{rx.\alpha}{x} \rangle \varphi \rightarrow \langle rx.\alpha \rangle \varphi \\ (\alpha_G) \quad & \langle \sigma x.\alpha \rangle \varphi \leftrightarrow \langle \sigma y.\alpha \frac{y}{x} \rangle \varphi \quad (y \text{ fresh}, \sigma \in \{r, \iota\}) \end{aligned}$$

and the rules

$$\begin{aligned} (MP_G) \quad & \frac{\varphi \quad \varphi \rightarrow \psi}{\psi} & (M_G) \quad & \frac{\varphi \rightarrow \psi}{\langle \alpha \rangle \varphi \rightarrow \langle \alpha \rangle \psi} \\ (\mu_G) \quad & \frac{\langle \alpha \frac{\beta;?;\psi;\perp}{x} \rangle \varphi \rightarrow \langle \beta \rangle \psi}{\langle rx.\alpha \rangle \varphi \rightarrow \langle \beta \rangle \psi} \end{aligned}$$

We write $rlGL \vdash \varphi$ to mean that there is a Hilbert style proof of the formula φ consisting only of *right-linear game logic* formulas. A more general proof calculus for full recursive game logic is of interest as well. (Note however that μ_G is not sound for this language, since the soundness proof relies on right-linearity.) Because there can not be a recursive and complete such calculus, we restrict attention to the calculus for right-linear game logic.

The new rule μ_G is an adaptation of the Park fixpoint induction rule to the setting of right-linear games. Completeness requires the rule M_G only for games of the form a , a^d and x , since the more general rule is derivable. With the more general definition however it is clear that if $GL \vdash \varphi$, then by substituting free occurrences of x across the entire proof also $GL \vdash \varphi \frac{\alpha}{x}$.

If there is a proof of $rlGL \vdash \varphi$ consisting only of game logic formulas we write $GL \vdash \varphi$. Observe that this is equivalent to Parikh's calculus for game logic [36]. The rule μ_G restricted to GL formulas is interderivable with the standard game logic induction rule

$$(\mu_G^*) \quad \frac{\rho \vee \langle \alpha \rangle \psi \rightarrow \psi}{\langle \alpha^* \rangle \rho \rightarrow \psi}$$

in the context of the axioms $\cup$, $;$, $?$, $!$ and M_G . Similarly the axiom fp_G restricted to GL formulas is interderivable with

$$(*_G) \quad \varphi \vee \langle \alpha \rangle \langle \alpha^* \rangle \varphi \rightarrow \langle \alpha^* \rangle \varphi$$

in the presence of axioms $\cup$, $\neg$, $;$ and $?$.

The following axioms and rule are derivable from $\neg$:

$$\begin{aligned} (!) \quad & \langle ! \rangle \varphi \leftrightarrow (\varphi \rightarrow \psi) \\ (\cap) \quad & \langle \alpha \cap \beta \rangle \varphi \leftrightarrow \langle \alpha \rangle \varphi \wedge \langle \beta \rangle \varphi \\ (v_G) \quad & \frac{\langle \beta \rangle \psi \rightarrow \langle \alpha \frac{\beta;?;\psi;\perp}{x} \rangle \rho}{\langle \beta \rangle \psi \rightarrow \langle vx.\alpha \rangle \rho} \end{aligned}$$

Rule v_G captures the coinductive nature of the greatest fixpoints that define Demons's games, similarly to how the induction rule μ_G formalizes the inductive nature of Angel's games.

Originally Parikh's calculus was designed for game logic interpreted over monotone neighbourhood structures. An extension of Parikh's calculus with the axioms $[\wedge]$, K and $[\top]$ from Kozen's calculus turns it into a suitable axiomatization for game logic over Kripke structures. The appearance of a in these axioms ranges only over *atomic games* and not over arbitrary games for which these axioms would be unsound. If there is a proof of φ in the $rlGL$ calculus from the axioms $[\wedge]$, K and $[\top]$ we write $rlGL + G \vdash \varphi$. The proof calculus is sound for monotone neighbourhood structures and Kripke structures:

LEMMA 6.3 (SOUNDNESS). *For any formula φ of $rlGL$*

- (1) *if $rlGL \vdash \varphi$ then $\models \varphi$*
- (2) *if $rlGL + G \vdash \varphi$ then $\models_K \varphi$*

PROOF. For most axioms and rules the proof of soundness is exactly the same as for game logic [36]. Soundness of μ_G follows as in game logic with Lemma 3.7. $\square$

LEMMA 6.4. *The right-linearity axiom and the reverse fixpoint axiom are derivable in the $rlGL$ calculus for $rlGL$*

$$\begin{aligned} (RL) \quad & \langle \alpha \frac{\beta}{x} \rangle \varphi \leftrightarrow \langle \alpha \frac{\beta;?;\varphi;\perp}{x} \rangle \varphi \quad (\alpha \text{ is right-linear}) \\ (rfp_G) \quad & \langle rx.\alpha \rangle \varphi \rightarrow \langle \alpha \frac{rx.\alpha}{x} \rangle \varphi \end{aligned}$$

PROOF. Axiom RL can be derived by induction on α . For axiom rfp_G by RL and μ_G , it suffices to prove

$$\langle \alpha \frac{rx.\alpha}{x} \rangle \varphi \rightarrow \langle \alpha \frac{rx.\alpha}{x} \rangle \varphi$$

in the calculus. By induction one proves more generally for any right-linear game β that

$$\langle \beta \frac{rx.\alpha}{x} \rangle \varphi \rightarrow \langle \beta \frac{rx.\alpha}{x} \rangle \varphi.$$

For the case that $\beta \equiv x$ this holds by axiom fp_G . $\square$

6.2.1 Proof Calculus for Sabotage Game Logic. In addition to the proof calculus for right-linear game logic, the proof calculus for game logic can be extended to a proof calculus for sabotage game logic. This involves adding axioms for the rule-change games:

$$\begin{aligned}
(\sim) \langle \sim a; \langle \alpha \frac{a^d \tilde{\gamma}}{x} \frac{\tilde{\delta}}{y} \rangle \rangle \varphi &\leftrightarrow \langle \alpha \frac{\sim a; \tilde{\gamma}}{x} \frac{\tilde{\delta}}{y} \rangle \varphi & (\dagger) \\
(\wr) \langle \sim a^d; \langle \alpha \frac{a^d \tilde{\gamma}}{x} \frac{\tilde{\delta}}{y} \rangle \rangle \varphi &\leftrightarrow \langle \alpha \frac{\wr a^d; \tilde{\gamma}}{x} \frac{\tilde{\delta}}{y} \rangle \varphi & (\dagger) \\
(\Downarrow) \langle \eta; \sim a_i; \beta_i \rangle (\langle \alpha \frac{\beta}{w} \rangle \varphi &\leftrightarrow \langle \alpha \frac{(\bigcup_{i=1}^n a_i \beta_i) \beta}{w} \rangle \varphi) \quad (\eta = \sim a_1^d; \dots; \sim a_n^d, \text{ and } \ddagger) \\
(\cong) \langle \sim a; \alpha \frac{a}{x} \rangle \varphi &\leftrightarrow \langle \sim a; \alpha \frac{a \sim a}{x} \rangle \varphi \\
(\simeq) \langle \alpha \frac{\tilde{\gamma} \tilde{\delta} \beta}{x} \rangle \varphi &\leftrightarrow \langle \alpha \frac{\tilde{\gamma} \tilde{\delta} \beta}{x} \rangle \varphi \quad (a, a^d \notin \alpha, \varphi \text{ and } \eta_i, \delta_i \in \{\sim a, \sim a^d\}) \\
(\sim_1) \langle \alpha \frac{a \beta}{x} \frac{\sim a^d}{y} \rangle \varphi &\leftrightarrow \langle \alpha \frac{a \beta}{x} \frac{\sim a^d \eta}{y} \rangle \varphi \quad (\eta \in \{\sim b, \sim b^d\} \text{ and } \sim a, b, b^d \notin \alpha, \varphi)
\end{aligned}$$

The notation in axioms should be understood schematically. In $\sim$ and $\wr$ the game α is a GL_s game with two lists of free variables $\tilde{x}, \tilde{y}, \tilde{z}$ occurring only right-linearly¹ and $\tilde{\beta}, \tilde{\gamma}, \tilde{\delta}$ are lists of sabotage game logic games of the same dimensions that are substituted into α . The axioms $\sim$ and $\wr$ require the side condition $\dagger$ that $a, a^d, \sim a, \sim a^d$ do not appear in α and that a, a^d do not appear in φ . The axioms $\sim$ and $\wr$ capture the effect of a rule change deep within a formula. The axiom $\Downarrow$ allows to reason about more complex branching behaviour. It requires as a side condition $\ddagger$ saying that each β_i is a formula of the form $\sim b_1; \dots; \sim b_j; \sim b_{j+1}^d; \dots; \sim b_m^d$, such that sabotages of the atomic games a_i, b_j appear in α only in the form $\eta; \sim a_i; \beta_i$. Axiom $\simeq$ allows the uniform removal of unneeded sabotages.

The proof calculus for sabotage game logic consists of the six axioms $\sim, \wr, \Downarrow, \simeq, \sim_1$ and $\cong$ together with $\neg, ?, \cup, :, *_G, \text{MP}_G, \mu_G^*$ and M_G . If there is a Hilbert-style proof of φ in this calculus consisting only of GL_s formulas, write $\text{GL}_s \vdash \varphi$. The proof calculus for GL_s can also be adjusted to work for Kripke structures. This extension adds the axioms $[\wedge], \text{K}$ and $[\top]$. Write $\text{GL}_s + G \vdash \varphi$ if there is a proof of φ in this extension.

LEMMA 6.5 (SOUNDNESS). *For any formula φ of GL_s*

- (1) *if $\text{GL}_s + G \vdash \varphi$ then $\models_K \varphi$*
- (2) *if $\text{GL}_s \vdash \varphi$ then $\models \varphi$*

PROOF. Soundness of the common part of the proof calculus goes through exactly as for game logic [36]. Soundness of $\sim$ and $\wr$ is proved by induction on α using the side condition $\dagger$. $\square$

See full proof on page 17.

6.3 Equivalence of Proof Calculi

This section shows that the translations between right-linear game logic and the modal μ -calculus shows not only equiexpressiveness, but also that the proof calculi are equivalent. This will enable the transfer of completeness from modal μ -calculus to right-linear game logic.

Rank. In order to flatten the mutual inductions that arise naturally, because of the mutually inductive definition of formulas and games, a well-order on all formulas is required. This order is given by the rank of a formula of game logic defined in Appendix A. Many of the proofs about the calculi for game logic are carried out by induction on the rank of a formula.

¹No variable appears on the left hand side of a composition, in the scope of a loop or in a test.

Correctness of Translation. The translations between right-linear game logic and the modal μ -calculus have been proved sound semantically. In order to use these to relate the proof calculi, the soundness of the translation needs to be proved in the calculus itself. Since each calculus can only talk about formulas in its respective language the correct notion of correctness here is that of Corollary 5.3.

LEMMA 6.6 (PROVABLE CORRECTNESS). (1) $\text{rGL} \vdash \varphi \leftrightarrow \varphi^{\sharp\sharp}$
for any well-named rGL formula φ

(2) and $\text{mL}_\mu \vdash \varphi \leftrightarrow \varphi^{\sharp\sharp\flat}$ for any well-named L_μ formula φ

See proof on page 18.

The key then is that proofs in the modal μ -calculus can be turned into right-linear game logic proofs. Since the modal μ -calculus is complete and has the same expressive power as right-linear game logic it follows that any formula φ is provable up to translation.

THEOREM 6.7 (EQUIVALENCE OF CALCULI). *Right-linear game logic and the modal μ -calculus prove the same formulas modulo translation.*

- (1) $\text{rGL} \vdash \varphi$ iff $\text{mL}_\mu \vdash \varphi^{\flat}$ for closed well-named rGL formulas φ
- (2) $\text{L}_\mu \vdash \varphi$ iff $\text{rGL} \vdash \varphi^{\sharp}$ for well-named L_μ formulas φ

PROOF. For the forward direction of (1) observe that Lemma 6.3 and Corollary 5.5 imply $\models \varphi^{\flat}$. Hence $\text{mL}_\mu \vdash \varphi^{\flat}$ follows from Proposition 6.1. The forward direction of (2) is proved in the appendix as Lemma D.2.

For the backward direction of (1) consider $\text{mL}_\mu \vdash \varphi^{\flat}$. By the forward direction of (2) also $\text{rGL} \vdash \varphi^{\sharp\sharp\flat}$. Then by Lemma 6.6 conclude $\text{rGL} \vdash \varphi$. The backward implication of (2) is similar. $\square$

6.4 Completeness for rGL

Because the translations are semantically correct and preserve provability by the results of Section 6.3, completeness of rGL follows.

THEOREM 6.8 (rGL COMPLETENESS). *For a closed rGL formula φ*

- (1) $\text{rGL} + G \vdash \varphi$ iff $\models_K \varphi$
- (2) $\text{rGL} \vdash \varphi$ iff $\models \varphi$

PROOF. (2) By α_G assume without loss of generality that φ is well-named. Consider the following chain of equivalences

$$\begin{aligned}
&\models \varphi \\
\text{iff } &\models \varphi^{\flat} && (\text{Corollary 5.5}) \\
\text{iff } &\text{mL}_\mu \vdash \varphi^{\flat} && (\text{Proposition 6.1}) \\
\text{iff } &\text{rGL} \vdash \varphi^{\sharp\sharp\flat} && (\text{Theorem 6.7}) \\
\text{iff } &\text{rGL} \vdash \varphi && (\text{Lemma 6.6 and MP})
\end{aligned}$$

This proves the equivalence.

(1) A slight modification of the proof of Lemma D.2 shows that $\text{rGL} + G \vdash \psi^{\sharp}$ provided $\text{L}_\mu \vdash \psi$ for a closed L_μ formula ψ . All that needs to be added is that the $\sharp$ -translation of the axioms $[\wedge], \text{K}$ and $[\top]$ are again instances of the same axiom in $\text{rGL} + G$. The same argument as for (2) then applies to show completeness of rGL . $\square$

7 SABOTAGE GAME LOGIC AND MODAL μ

In this section the equiexpressiveness of rGL and modal μ -calculus are exploited to obtain the desired equiexpressiveness of sabotage game logic with the modal μ -calculus as well as completeness of sabotage game logic through translations turning sabotage modalities into named subgames and vice versa.

7.1 Embedding of GL_s into rGL

The difficulty in embedding GL_s into rGL is that the ownership information about previously committed acts of sabotage must be taken into account. This can be done here by coding this information on the claimed atomic games into the nesting structure of the fixpoint variables. To simplify this coding, we use simultaneous fixpoints, which do not add to the expressive power. This is captured by the following proposition, known as Bekić's Theorem,

THEOREM 7.1 ([4, LEMMA 1.4.2]). *For variables $x_1, \dots, x_n$ and rGL games $\alpha_1, \dots, \alpha_n$ there are rGL games $\beta_1, \dots, \beta_n$ such that*

$$\begin{pmatrix} \mathcal{N}[\beta_1]^I \\ \mathcal{N}[\beta_2]^I \\ \vdots \\ \mathcal{N}[\beta_m]^I \end{pmatrix} = \mu \begin{pmatrix} w_1 \\ w_2 \\ \vdots \\ w_n \end{pmatrix} \cdot \begin{pmatrix} \mathcal{N}[\alpha_1]^I[\vec{x} \mapsto \vec{w}] \\ \mathcal{N}[\alpha_2]^I[\vec{x} \mapsto \vec{w}] \\ \vdots \\ \mathcal{N}[\alpha_m]^I[\vec{x} \mapsto \vec{w}] \end{pmatrix}$$

We write $r_i(x_1, \dots, x_n).(\alpha_1, \dots, \alpha_n)$ for this β_i .

The Embedding. We fix for every possible context $c \in C$ a fresh variable y_c . For any formula φ and any game α of sabotage game logic we define a translation α^c depending on the context c . The context allows the translation to depend on the state of ownership of atomic games. Moreover the translation of games will contain free variables y_c . Those mark the end of the game and keep track of the context in which this end has been reached. This allows a compositional definition of the translation. Let

$$\begin{aligned} (P)^c &= ?P; !\perp & (\neg P)^c &= ?\neg P; !\perp \\ (\varphi \vee \psi)^c &= \varphi^c \cup \psi^c & (\varphi \wedge \psi)^c &= \varphi^c \cap \psi^c \\ (a)^c &= \begin{cases} a; y_c & \text{if } c(a) = \emptyset \\ y_c & \text{if } c(a) = \diamond \\ ?\perp & \text{if } c(a) = \square \end{cases} & (a^d)^c &= \begin{cases} a^d; y_c & \text{if } c(a) = \emptyset \\ !\perp & \text{if } c(a) = \diamond \\ y_c & \text{if } c(a) = \square \end{cases} \\ (\sim a)^c &= y_{c \frac{\diamond}{a}} & (\sim a^d)^c &= y_{c \frac{\square}{a}} \\ (? \varphi)^c &= \varphi^c \cap y_c & (! \varphi)^c &= \overline{\varphi^c} \cup y_c \\ (\alpha \cup \beta)^c &= \alpha^c \cup \beta^c & (\alpha \cap \beta)^c &= \alpha^c \cap \beta^c \end{aligned}$$

The translation of atomic games and sabotage games illustrates the importance of translating relative to a context. The translation of formulas $\langle a \rangle \varphi$ and games $\alpha; \beta$ and α^* are slightly more involved. For the first two define

$$\begin{aligned} (\langle a \rangle \varphi)^c &= \alpha^c \frac{\varphi; ?\perp}{y_c} \\ (\alpha; \beta)^c &= \alpha^c \frac{\beta^c}{y_c} \end{aligned}$$

where the notation $\alpha^c \frac{\beta^c}{y_c}$ means that any instance of a variable y_e is replaced by β^e , the translation of β with respect to the context e .

For a fixpoint game α^* list all the atomic games $a_1, \dots, a_m$ such that either $c(a) \neq \emptyset$ or $\sim a$ or $\sim a^d$ appears in α . List $c_1, \dots, c_m$ all

possible contexts that satisfy $c_i(a) = \emptyset$ for all $a \notin \{a_1, \dots, a_m\}$. Define the translation of repetition games

$$\begin{aligned} (\alpha^*)^{c_i} &= r_i(z_{c_1}, \dots, z_{c_n}).(y_{c_1} \cup \alpha^{c_1} \frac{z_{c_1}}{y_{c_1}}, \dots, y_{c_n} \cup \alpha^{c_n} \frac{z_{c_n}}{y_{c_n}}) \\ (\alpha^\times)^{c_i} &= r_i(z_{c_1}, \dots, z_{c_n}).(y_{c_1} \cap \alpha^{c_n} \frac{z_{c_n}}{y_{c_n}}, \dots, y_{c_n} \cap \alpha^{c_1} \frac{z_{c_1}}{y_{c_1}}) \end{aligned}$$

where the z_c are fresh variables. Observe that the translation of any GL_s formula or game is a right-linear rGL game.

The next proposition shows that the translation is correct.

PROPOSITION 7.2. *For any formula φ of sabotage game logic*

$$\mathcal{N}[\varphi]_s \upharpoonright c = \mathcal{N}[\varphi^c](\emptyset).$$

and for any game α of sabotage game logic

$$\mathcal{N}[\alpha]_s(U) \upharpoonright c = \mathcal{N}[\alpha^c]^I(\emptyset)$$

where $I(y_e) = U \upharpoonright e$.

See proof on page 19.

The translation of a sabotage game logic formulas into a formula of right-linear game logic potentially grows very quickly. The upper bound on the length of the translation of a fixpoint game obtained from the proof above is

$$|(\alpha)^c| \leq (C \cdot |\alpha|)^{3^\ell \uparrow k},$$

where k is the fixpoint nesting depth of α and ℓ is the number of atomic games in α .² This comes from the fact that the translation of any game α considers all of the up to 3^ℓ -many relevant contexts. The only known transformation from vectorial fixpoints to non-vectorial nested fixpoints as in Theorem 7.1 grows exponentially in the formula size. Consequently every fixpoint leads to a doubly exponential blow-up in length. In [13] it is shown that reducing vectorial fixpoints to non-vectorial fixpoints is at least as hard as solving parity games, for which the existence of a polynomial time algorithm is a longstanding open question. It was conjectured [12] that a vectorial fixpoint formula can be exponentially smaller than the shortest equivalent non-vectorial formula.

While it is unclear to what extent this upper bound is optimal, it suggests that complex formulas of the modal μ -calculus may in general be expressed much more succinctly in sabotage game logic.

7.2 Embedding of rGL into GL_s

The challenge of the converse translation from rGL to GL_s is that the arbitrarily nested named games of rGL need to be turned into structured repetition games of sabotage game logic. It turns out that using sabotage, players can force the behaviour of nested named games onto structured repetition games. To facilitate this, fix fresh atomic games b_x, c_x for every variable x . By induction define the

²We use Knuth's up-arrow notation for k -fold iterated exponentiation.

translation of a rGL formula and game into GL_s:

$$\begin{aligned}
(P)^{\natural} &= P & (\neg P)^{\natural} &= \neg P \\
(\varphi \vee \psi)^{\natural} &= \varphi^{\natural} \vee \psi^{\natural} & (\varphi \wedge \psi)^{\natural} &= \varphi^{\natural} \wedge \psi^{\natural} \\
(\langle \alpha \rangle \varphi)^{\natural} &= \langle \alpha^{\natural} \rangle \varphi^{\natural} \\
(a)^{\natural} &= a & (a^d)^{\natural} &= a^d \\
(? \varphi)^{\natural} &= ? \varphi^{\natural} & (! \varphi)^{\natural} &= ! \varphi^{\natural} \\
(\alpha \cup \beta)^{\natural} &= \alpha^{\natural} \cup \beta^{\natural} & (\alpha \cap \beta)^{\natural} &= \alpha^{\natural} \cap \beta^{\natural} \\
(\alpha; \beta)^{\natural} &= \alpha^{\natural}; \beta^{\natural} & (x)^{\natural} &= x \\
\\
(rx.\alpha)^{\natural} &= \delta; (c_x; \delta^d; \alpha^{\natural} \frac{\delta}{x})^*; b_x & \text{where } \delta &\equiv \sim b_x^d; \sim c_x \\
(\lambda x.\alpha)^{\natural} &= \delta^d; (c_x^d; \delta; \alpha^{\natural} \frac{\delta^d}{x})^{\times}; b_x^d & \text{where } \delta &\equiv \sim b_x^d; \sim c_x
\end{aligned}$$

While only the translation of closed formulas of right-linear game logic is of interest, the translation is defined more generally for open formulas into the extension of sabotage game logic with free variables.

Intuitively the translation of the fixpoints uses rule changes to remove the non-determinism from the repetition game. In a normal α^* game it is Angel's choice whether to continue playing the game α or not. However in general $rx.\alpha$ games, this choice is made differently. If the variable x is reached the game must be repeated. If the game ends without reaching this variable it must not be repeated. The translation enforces this deterministic behaviour of the repetition game in a $*$ game. Although the $*$ -game theoretically allows Angel to stop prematurely, Demon will have changed the rules ($\sim b_x^d$), so that Angel would automatically lose afterwards. Similarly if Angel does not stop when she ought to, she will also lose immediately, due to a demonic claim on $\sim c_x^d$. This ensures correctness of the translation.

PROPOSITION 7.3 (CORRECT $\natural$). *For any closed, well-named formula φ of right-linear game logic the translation $\varphi^{\natural}$ is a formula of sabotage game logic with*

$$\mathcal{N}[\varphi^{\natural}]_s \mid c = \mathcal{N}[\varphi].$$

See proof on page 20.

THEOREM 7.4 (EQUIEXPRESSIVENESS). *Sabotage game logic, right-linear game logic and the modal μ -calculus are equiexpressive.*

PROOF. By Propositions 7.2 and 7.3 and Corollary 5.5. $\square$

The equiexpressiveness of sabotage game logic and the modal μ -calculus means that sabotage game logic inherits many of the nice properties of the modal μ -calculus for free.

THEOREM 7.5 (META PROPERTIES). *Sabotage game logic has the small model property and the satisfiability problem is decidable.*

PROOF. The modal μ -calculus has these properties [40, 43], and they can be transferred to sabotage game logic by Theorem 7.4. $\square$

7.3 Proof Transformations

This section shows that the translation also respects the proof calculus. Combined with the semantic correctness of the translation this allows the transfer of completeness from rGL to GL_s. The key fact needed about the translation is that the sabotage paraphrasing of the named recursive games, *provably* behaves the same as the extremal fixpoint it denotes.

LEMMA 7.6. *For any formula φ and any game α of right-linear game logic the following hold:*

- (1) $\overline{\varphi^{\natural}} \equiv \overline{\varphi}^{\natural}$
- (2) $GL_s \vdash (\langle \alpha \frac{rx.\alpha}{x} \rangle \varphi \rightarrow \langle rx.\alpha \rangle \varphi)^{\natural}$
- (3) $GL_s \vdash (\langle rx.\alpha \rangle \varphi)^{\natural} \rightarrow \langle \beta \rangle \psi$ if $GL_s \vdash \langle \alpha^{\natural} \frac{\beta; ?\psi; !\perp}{x} \rangle \varphi^{\natural} \rightarrow \langle \beta \rangle \psi$

See proof on page 21.

LEMMA 7.7. *Suppose φ is a formula of sabotage game logic then $GL_s \vdash (\langle \varphi^{c0} \rangle \perp)^{\natural} \rightarrow \varphi$.*

See proof on page 21.

Since the translation $\natural$ provably behaves as intended, rGL proofs can be translated completely to GL_s.

PROPOSITION 7.8. *Suppose φ is a closed right-linear game logic formula and $rGL \vdash \varphi$ then $GL_s \vdash \varphi^{\natural}$.*

See proof on page 23.

THEOREM 7.9 (SABOTAGE GAME LOGIC COMPLETENESS). *Sabotage game logic is sound and complete. That is for all GL_s formulas φ :*

$$GL_s \vdash \varphi \quad \text{iff} \quad \models \varphi$$

See proof on page 23.

7.4 A Completion of Parikh's Calculus

Fix infinite disjoint sets $\mathbb{G}_1 \cup \mathbb{G}_2 \cup \mathbb{G}_3 = \mathbb{G}$. For all $b \in \mathbb{G}_2$ fix a unique $\tilde{b} \in \mathbb{G}_3$. Consider the set $\mathfrak{A}$ consisting of all formulas obtained from instances of the axioms $\sim, \wr, \approx$ and $\mathfrak{N}$ of sabotage game logic which contain atomic games only from $\mathbb{G}_2$ by replacing all $\sim b$ and all $\sim b^d$ by $\tilde{b}$ and $(\tilde{b})^d$ respectively. Taken as axioms these game logic formulas suffice to complete Parikh's proof calculus for game logic. Write $GL + \mathfrak{A} \vdash \varphi$ if there is a proof of φ in Parikh's calculus from these axioms.

THEOREM 7.10 (GAME LOGIC COMPLETENESS). *For any GL formula φ with atomic games from $\mathbb{G}_1$:*

- (1) $\models \varphi$ iff $GL + \mathfrak{A} \vdash \varphi$
- (2) $\models_K \varphi$ iff $GL + G + \mathfrak{A} \vdash \varphi$

PROOF. *Soundness:* Suppose $GL + \mathfrak{A} \vdash \varphi$ and consider a monotone neighbourhood structure $\mathcal{N}$. Replacing every appearance of some $\tilde{b}$ in a proof of $GL + \mathfrak{A} \vdash \varphi$ by $\sim b$ gives a proof for $GL_s \vdash \varphi$. Hence by soundness (Lemma 6.5) conclude that $\models \varphi$.

Completeness: Suppose φ is a valid formula. As in the proof of Theorem 7.9 obtain a proof of $GL_s \vdash \varphi$. By the construction of this proof (Proposition 7.8) it is clear that the proof contains games of the form $\sim a, \sim a^d$ only for atomic games that do not already appear in φ . Without loss of generality assume the proof contains games of the form $\sim b, \sim b^d$ only for $b \in \mathbb{G}_2$.

The GL_s proof of φ can be transformed into a game logic proof by uniformly replacing every game of the form $\sim a$ by $\tilde{b}$ and every $\sim a^d$ by $\tilde{b}^d$. Instances of the axioms $\sim, \wr, \approx$ and $\gg$ in the original proof become instances of axioms in $\mathfrak{A}$. Hence the modified version of the proof is a valid GL proof of φ from the axioms in $\mathfrak{A}$. $\square$

The results in this section pave the way for using the proof calculus for GL_s to address the question of completeness of Parikh's axiomatization for game logic through a proof transformation by eliminating instances of axioms from $\mathfrak{A}$.

8 CONCLUSION

This paper studies how logic, games, and fixpoints meet by introducing two different extensions of Game Logic. The first, sabotage game logic, allows players to sabotage their opponent, while the second, recursive game logic, adds recursive games.

Not only is sabotage game logic (GL_s) well-suited to describe and investigate games with rule changes by logical means, but perhaps surprisingly, sabotage game logic has a number of additional advantages over game logic. Unlike game logic, the extension sabotage game logic allows exactly the right amount of state to increase its expressive power to match the modal μ -calculus, without sacrificing the desirable logical properties of game logic.

The virtue of recursive game logic (RGL) is that it allows the description of games featuring arbitrarily nested recursive games. Unlike ordinary Game Logic, the extended version is significantly more expressive than the modal μ -calculus, although it remains syntactically close to GL. We have identified the fragment of RGL that corresponds exactly to the modal μ -calculus in expressiveness and transferred completeness of modal μ -calculus to obtain a complete and natural proof calculus for this fragment.

Additionally, it was shown that sabotage game logic and the modal μ -calculus are equivalent in expressiveness via a translation going through the right-linear fragment of recursive game logic. Completeness of the natural Hilbert style proof calculus for sabotage game logic GL_s was obtained as a consequence. This is in contrast to game logic GL for which completeness of the natural proof calculus is not known [30]. Completeness of sabotage game logic was used to obtain the completeness of a modest extension of Parikh's proof calculus for game logic GL.

Future Research. The completeness of sabotage game logic suggests an interesting approach to studying proof calculi for game logic. It reduces the problem of the completeness of Parikh's axiomatization of game logic to eliminating instances of the new axioms in a proof. Equiexpressiveness with L_μ indicates that atomic games for sabotage are worth studying further.

The translation from GL_s into L_μ leads to a non-elementary blow-up in formula length. This raises the question whether this increase is necessary and if better algorithms exist that directly target the model checking and satisfiability problems of GL_s .

ACKNOWLEDGMENTS

This project was supported by an Alexander von Humboldt Professorship.

REFERENCES

- [1] S. Abramsky and P.-A. Mellies. 1999. Concurrent games and full completeness. In *Proceedings. 14th Symposium on Logic in Computer Science (Cat. No. PR00158)*. 431–442. <https://doi.org/10.1109/LICS.1999.782638>
- [2] Bahareh Afshari and Graham E. Leigh. 2017. Cut-free completeness for modal μ -calculus. In *32nd Annual ACM/IEEE Symposium on Logic in Computer Science, LICS 2017, Reykjavik, Iceland, June 20-23, 2017*. IEEE Computer Society, 1–12. <https://doi.org/10.1109/LICS.2017.8005088>
- [3] Carlos Areces, Raul Fervari, Guillaume Hoffmann, and Mauricio Martel. 2017. Undecidability of Relation-Changing Modal Logics. In *Dynamic Logic. New Trends and Applications - First International Workshop, DALI 2017, Brasilia, Brazil, September 23-24, 2017, Proceedings (Lecture Notes in Computer Science, Vol. 10669)*, Alexandre Madeira and Mario R. F. Benevides (Eds.). Springer, 1–16. https://doi.org/10.1007/978-3-319-73579-5_1
- [4] André Arnold and Damian Niwinski. 2001. *Rudiments of mu-calculus*. Studies in Logic and the Foundations of Mathematics, Vol. 146. North Holland Publishing Co., Amsterdam.
- [5] Guillaume Aucher, Johan van Benthem, and Davide Grossi. 2015. Sabotage Modal Logic: Some Model and Proof Theoretic Aspects. In *Logic, Rationality, and Interaction - 5th International Workshop, LORI 2015 Taipei, Taiwan, October 28-31, 2015, Proceedings (Lecture Notes in Computer Science, Vol. 9394)*, Wiebe van der Hoek, Wesley H. Holliday, and Wen-Fang Wang (Eds.). Springer, 1–13. https://doi.org/10.1007/978-3-662-48561-3_1
- [6] Guillaume Aucher, Johan van Benthem, and Davide Grossi. 2018. Modal logics of sabotage revisited. *J. Log. Comput.* 28, 2 (2018), 269–303. <https://doi.org/10.1093/LOGCOM/EXX034>
- [7] Dietmar Berwanger. 2003. Game Logic is Strong Enough for Parity Games. *Studia Logica* 75, 2 (2003), 205–219. <https://doi.org/10.1023/A:1027358927272>
- [8] Dietmar Berwanger, Erich Grädel, and Giacomo Lenzi. 2007. The Variable Hierarchy of the μ -calculus is strict. *Theory Comput. Syst.* 40, 4 (2007), 437–466.
- [9] Francesca Zaffora Blando, Krzysztof Mierzewski, and Carlos Areces. 2020. The Modal Logics of the Poison Game. In *Knowledge, Proof and Dynamics*, Fenrong Liu, Hiroakira Ono, and Junhua Yu (Eds.). Springer, 3–23.
- [10] Julian C. Bradfield. 1996. The Modal μ -calculus Alternation Hierarchy is Strict. In *CONCUR '96, Concurrency Theory, 7th International Conference, Pisa, Italy, August 26-29, 1996, Proceedings (LNCS, Vol. 1119)*, Ugo Montanari and Vladimiro Sassone (Eds.). Springer, 233–246. https://doi.org/10.1007/3-540-61604-7_58
- [11] Julian C. Bradfield and Colin Stirling. 2006. Modal μ Calculi. In *Handbook of Modal Logic*, Patrick Blackburn, Johann van Benthem, and Frank Wolter (Eds.). Elsevier, 721–756.
- [12] Julian C. Bradfield and Igor Walukiewicz. 2018. The μ -calculus and Model Checking. In *Handbook of Model Checking*, Edmund M. Clarke, Thomas A. Henzinger, Helmut Veith, and Roderick Bloem (Eds.). Springer, 871–919. https://doi.org/10.1007/978-3-319-10575-8_26
- [13] Florian Bruse, Oliver Friedmann, and Martin Lange. 2015. On guarded transformation in the modal μ -calculus. *Log. J. IGPL* 23, 2 (2015), 194–216. <https://doi.org/10.1093/JIGPAL/JZU030>
- [14] Facundo Carreiro. 2015. *Fragments of fixpoint logics*. Ph.D. Dissertation. Ph.D. thesis, University of Amsterdam.
- [15] Krishnendu Chatterjee, Thomas A. Henzinger, and Nir Piterman. 2007. Strategy Logic. In *CONCUR 2007 - Concurrency Theory, 18th International Conference, CONCUR 2007, Lisbon, Portugal, September 3-8, 2007, Proceedings (Lecture Notes in Computer Science, Vol. 4703)*, Luis Caires and Vasco Thudichum Vasconcelos (Eds.). Springer, 59–73. https://doi.org/10.1007/978-3-540-74407-8_5
- [16] Corina Cirstea, Clemens Kupke, and Dirk Pattinson. 2011. EXPTIME Tableaux for the Coalgebraic μ -Calculus. *Log. Methods Comput. Sci.* 7, 3 (2011). [https://doi.org/10.2168/LMCS-7\(3:3\)2011](https://doi.org/10.2168/LMCS-7(3:3)2011)
- [17] Pierre Clairambault. 2009. Least and Greatest Fixpoints in Game Semantics. In *Foundations of Software Science and Computational Structures, 12th International Conference, FOSSACS 2009, Held as Part of the Joint European Conferences on Theory and Practice of Software, ETAPS 2009, York, UK, March 22-29, 2009. Proceedings (Lecture Notes in Computer Science, Vol. 5504)*, Luca de Alfaro (Ed.). Springer, 16–31. https://doi.org/10.1007/978-3-642-00596-1_3
- [18] Pierre Duchet and Henry Meyniel. 1993. Kernels in directed graphs: a poison game. *Discret. Math.* 115, 1-3 (1993), 273–276. [https://doi.org/10.1016/0012-365X\(93\)90496-G](https://doi.org/10.1016/0012-365X(93)90496-G)
- [19] Andrzej Ehrenfeucht. 1961. An application of games to the completeness problem for formalized theories. *Fundamenta Mathematicae* 49 (1961), 129–141. <https://api.semanticscholar.org/CorpusID:118038695>
- [20] E. Allen Emerson and Charanjit S. Jutla. 1991. Tree Automata, μ -Calculus and Determinacy (Extended Abstract). In *32nd Annual Symposium on Foundations of Computer Science, San Juan, Puerto Rico, 1-4 October 1991*. IEEE Computer Society, 368–377. <https://doi.org/10.1109/SFCS.1991.185392>
- [21] E. Allen Emerson, Charanjit S. Jutla, and A. Prasad Sistla. 2001. On model checking for the μ -calculus and its fragments. *Theor. Comput. Sci.* 258, 1-2 (2001), 491–522.

- [22] Sebastian Enqvist, Helle Hvid Hansen, Clemens Kupke, Johannes Marti, and Yde Venema. 2019. Completeness for Game Logic. In *34th Annual ACM/IEEE Symposium on Logic in Computer Science, LICS 2019, Vancouver, BC, Canada, June 24-27, 2019*. IEEE, 1–13. <https://doi.org/10.1109/LICS.2019.8785676>
- [23] Sebastian Enqvist, Fatemeh Seifan, and Yde Venema. 2018. Completeness for the modal μ -calculus: Separating the combinatorics from the dynamics. *Theor. Comput. Sci.* 727 (2018), 37–100. <https://doi.org/10.1016/j.tcs.2018.03.001>
- [24] Sebastian Enqvist, Fatemeh Seifan, and Yde Venema. 2019. Completeness for μ -calculus: A coalgebraic approach. *Ann. Pure Appl. Log.* 170, 5 (2019), 578–641. <https://doi.org/10.1016/j.apal.2018.12.004>
- [25] Alessandro Facchini, Yde Venema, and Fabio Zanasi. 2013. A Characterization Theorem for the Alternation-Free Fragment of the Modal μ -Calculus. In *Proceedings of the Twenty-Eighth Annual IEEE Symposium on Logic in Computer Science (LICS 2013)* (New Orleans, USA). IEEE Computer Society Press, 478–487.
- [26] Nina Gierasimczuk, Lena Kurzen, and Fernando R. Velázquez-Quesada. 2009. Games for Learning: A Sabotage Approach. In *Proceedings of the Second Multi-Agent Logics, Languages, and Organisations Federated Workshops, Turin, Italy, September 7-10, 2009 (CEUR Workshop Proceedings, Vol. 494)*, Matteo Baldoni, Cristina Baroglio, Jamal Bentahar, Guido Boella, Massimo Cossentino, Mehdi Dastani, Barbara Dunin-Keplicz, Giancarlo Fortino, Marie-Pierre Gleizes, João Leite, Viviana Mascardi, Julian A. Padget, Juan Pavón, Axel Polleres, Amal El Fallah Seghrouchni, Paolo Torroni, and Rineke Verbrugge (Eds.). CEUR-WS.org. <https://ceur-ws.org/Vol-494/famaspaper5.pdf>
- [27] Yuri Gurevich and Leo Harrington. 1982. Trees, Automata, and Games. In *Proceedings of the 14th Annual ACM Symposium on Theory of Computing, May 5-7, 1982, San Francisco, California, USA*, Harry R. Lewis, Barbara B. Simmons, Walter A. Burkhard, and Lawrence H. Landweber (Eds.). ACM, 60–65. <https://doi.org/10.1145/800070.802177>
- [28] Jaakko Hintikka. 1982. Game-theoretical semantics: insights and prospects. *Notre Dame Journal of Formal Logic* 23, 2 (1982), 219–241.
- [29] David Janin and Igor Walukiewicz. 1996. On the Expressive Completeness of the Propositional μ -Calculus with Respect to Monadic Second Order Logic. In *CONCUR '96, Concurrency Theory, 7th International Conference, Pisa, Italy, August 26-29, 1996, Proceedings (LNCS, Vol. 1119)*, Ugo Montanari and Vladimiro Sassone (Eds.). Springer, 263–277. https://doi.org/10.1007/3-540-61604-7_60
- [30] Johannes Kloibhofer. 2023. A note on the incompleteness of Afshari & Leigh’s system Clo. <https://doi.org/10.48550/arXiv.2307.06846> [math.LO]
- [31] Dexter Kozen. 1983. Results on the Propositional μ -Calculus. *Theor. Comput. Sci.* 27, 3 (1983), 333–354. [https://doi.org/10.1016/0304-3975\(82\)90125-6](https://doi.org/10.1016/0304-3975(82)90125-6)
- [32] Christof Löding and Philipp Rohde. 2003. Model Checking and Satisfiability for Sabotage Modal Logic. In *FST TCS 2003: Foundations of Software Technology and Theoretical Computer Science, 23rd Conference, Mumbai, India, December 15-17, 2003, Proceedings (Lecture Notes in Computer Science, Vol. 2914)*, Paritosh K. Pandya and Jaikumar Radhakrishnan (Eds.). Springer, 302–313. https://doi.org/10.1007/978-3-540-24597-1_26
- [33] Markus Müller-Olm. 1999. A Modal Fixpoint Logic with Chop. In *STACS 99, 16th Annual Symposium on Theoretical Aspects of Computer Science, Trier, Germany, March 4-6, 1999, Proceedings (Lecture Notes in Computer Science, Vol. 1563)*, Christoph Meinel and Sophie Tison (Eds.). Springer, 510–520. https://doi.org/10.1007/3-540-49116-3_48
- [34] Sara Negri. 2017. Proof Theory for Non-normal Modal Logics: The Neighbourhood Formalism and Basic Results. *FLAP* 4, 4 (2017). <http://www.collegepublications.co.uk/downloads/ifcolog00013.pdf>
- [35] Damian Niwinski and Igor Walukiewicz. 1996. Games for the μ -Calculus. *Theor. Comput. Sci.* 163, 1&2 (1996), 99–116. [https://doi.org/10.1016/0304-3975\(95\)00136-0](https://doi.org/10.1016/0304-3975(95)00136-0)
- [36] Rohit Parikh. 1983. Propositional Game Logic. In *FOCS*. 195–200. <https://doi.org/10.1109/SFCS.1983.47>
- [37] Marc Pauly and Rohit Parikh. 2003. Game Logic - An Overview. *Stud Logica* 75, 2 (2003), 165–182. <https://doi.org/10.1023/A:1027354826364>
- [38] André Platzer. 2015. Differential Game Logic. *ACM Trans. Comput. Log.* 17, 1 (2015), 1. <https://doi.org/10.1145/2817824>
- [39] Vaughan R. Pratt. 1981. A Decidable μ -Calculus: Preliminary Report. In *FOCS*. IEEE Computer Society, 421–427. <https://doi.org/10.1109/SFCS.1981.4>
- [40] Vaughan R. Pratt. 1981. A Decidable μ -Calculus: Preliminary Report. In *22nd Annual Symposium on Foundations of Computer Science, Nashville, Tennessee, USA, 28-30 October 1981*. IEEE Computer Society, 421–427. <https://doi.org/10.1109/SFCS.1981.4>
- [41] Philipp Rohde. 2006. On the μ -Calculus Augmented with Sabotage. In *Foundations of Software Science and Computation Structures, 9th International Conference, FOSSACS 2006, Held as Part of the Joint European Conferences on Theory and Practice of Software, ETAPS 2006, Vienna, Austria, March 25-31, 2006, Proceedings (Lecture Notes in Computer Science, Vol. 3921)*, Luca Aceto and Anna Ingólfssdóttir (Eds.). Springer, 142–156. https://doi.org/10.1007/11690634_10
- [42] Colin Stirling. 1996. Games and Modal μ -Calculus. In *Tools and Algorithms for Construction and Analysis of Systems, Second International Workshop, TACAS '96, Passau, Germany, March 27-29, 1996, Proceedings (Lecture Notes in Computer Science, Vol. 1055)*, Tiziana Margaria and Bernhard Steffen (Eds.). Springer, 298–312. https://doi.org/10.1007/3-540-61042-1_51
- [43] Robert S. Streett and E. Allen Emerson. 1989. An Automata Theoretic Decision Procedure for the Propositional μ -Calculus. *Inf. Comput.* 81, 3 (1989), 249–264.
- [44] Johan van Benthem. 2005. An Essay on Sabotage and Obstruction. In *Mechanizing Mathematical Reasoning, Essays in Honor of Jörg H. Siekmann on the Occasion of His 60th Birthday (Lecture Notes in Computer Science, Vol. 2605)*, Dieter Hutter and Werner Stephan (Eds.). Springer, 268–276. https://doi.org/10.1007/978-3-540-32254-2_16
- [45] Noah Abou El Wafa and André Platzer. 2022. First-Order Game Logic and Modal μ -Calculus. [arXiv:2201.10012 \[cs.LO\]](https://arxiv.org/abs/2201.10012)
- [46] Igor Walukiewicz. 1995. Completeness of Kozen’s Axiomatisation of the Propositional μ -Calculus. In *Proceedings, 10th Annual IEEE Symposium on Logic in Computer Science, San Diego, California, USA, June 26-29, 1995*. IEEE Computer Society, 14–24. <https://doi.org/10.1109/LICS.1995.523240>

A THE RANK

We define the rank of a recursive game logic formula and a recursive game logic game by (simultaneous) structural induction on the formula

$$\begin{aligned}
\text{rank}(P) &= 0 \\
\text{rank}(\neg P) &= 0 \\
\text{rank}(\varphi \vee \psi) &= \max\{\text{rank}(\varphi), \text{rank}(\psi)\} + 1 \\
\text{rank}(\varphi \wedge \psi) &= \max\{\text{rank}(\varphi), \text{rank}(\psi)\} + 1 \\
\text{rank}(\langle \alpha \rangle \varphi) &= \text{rank}(\alpha) + \text{rank}(\varphi) + 1 \\
\text{rank}(a) &= 0 \\
\text{rank}(a^d) &= 0 \\
\text{rank}(x) &= 0 \\
\text{rank}(\varphi?) &= \text{rank}(\varphi) + 2 \\
\text{rank}(\varphi?) &= \text{rank}(\varphi) + 3 \\
\text{rank}(\alpha \cup \beta) &= \max\{\text{rank}(\alpha), \text{rank}(\beta)\} + 2 \\
\text{rank}(\alpha \cap \beta) &= \max\{\text{rank}(\alpha), \text{rank}(\beta)\} + 2 \\
\text{rank}(\alpha; \beta) &= \text{rank}(\alpha) + \text{rank}(\beta) + 2 \\
\text{rank}(rx.\alpha) &= \text{rank}(\alpha) + 1 \\
\text{rank}(\lambda x.\alpha) &= \text{rank}(\alpha) + 1
\end{aligned}$$

B FIXPOINT LEMMAS

For any valuation I and any $A \in \mathcal{P}(|\mathcal{N}|)$ we let $I[A]$ be the modified valuation with $I[A](x) = I(x)(A)$.

LEMMA B.1. *Suppose φ is a FLC formula without composition and $A \in \mathcal{P}(|\mathcal{N}|)$. Then*

$$\mathcal{N}[\![\varphi]\!]^I(A) = \mathcal{N}[\![\varphi]\!]^{I[A]}(A).$$

Moreover

$$\begin{aligned}
(1) \quad \mathcal{N}[\![\mu x.\varphi]\!]^I(A) &= \mu B.(\mathcal{N}[\![\varphi]\!]^{I[x \mapsto \bar{B}]}(A)) \\
(2) \quad \mathcal{N}[\![\nu x.\varphi]\!]^I(A) &= \nu B.(\mathcal{N}[\![\varphi]\!]^{I[x \mapsto \bar{B}]}(A))
\end{aligned}$$

PROOF. We prove this by induction on φ . The only interesting case is for formulas of the form $\mu x.\varphi$. We define the following monotone functions

$$\Delta_1(w) = \mathcal{N}[\![\varphi]\!]^{I[x \mapsto w]} \quad \Delta_2(w) = \mathcal{N}[\![\varphi]\!]^{I[A][x \mapsto w]}$$

and prove that they satisfy the assumptions of Lemma 2.2. Two applications of the inductive hypothesis yield

$$\begin{aligned}
\Delta_1(w)(B) &= \mathcal{N}[\![\varphi]\!]^{I[x \mapsto w]}(B) = \mathcal{N}[\![\varphi]\!]^{I[x \mapsto w][B]}(B) \\
&= \mathcal{N}[\![\varphi]\!]^{I[x \mapsto \overline{w(B)}][B]}(B) = \mathcal{N}[\![\varphi]\!]^{I[x \mapsto \overline{w(B)}}(B) \\
&= \Delta_1(\overline{w(B)})(B)
\end{aligned}$$

for all $w \in \mathcal{W}(|\mathcal{N}|)$ and all $B \in \mathcal{P}(|\mathcal{N}|)$. The case for Δ_2 is similar.

Note also that $I[A][x \mapsto \bar{B}] = I[x \mapsto \bar{B}][A]$ and hence by the induction hypothesis $\Delta_1(\bar{B})(A) = \Delta_2(\bar{B})(A)$ for all $B \in \mathcal{P}(|\mathcal{N}|)$.

By Lemma 2.2 we compute

$$\begin{aligned}
\mathcal{N}[\![\mu x.\varphi]\!]^I(A) &= \mu w.\Delta_1(w)(A) = \mu B.(\Delta_1(\bar{B})(A)) \\
&= \mu B.(\Delta_2(\bar{B})(A)) = \mathcal{N}[\![\mu x.\varphi]\!]^{I[A]}(A).
\end{aligned}$$

The case for greatest fixpoints is symmetric. $\square$

A similar result holds for right-linear game logic.

LEMMA B.2. *Suppose φ is a formula and α a game of right-linear game logic and $A, B \in \mathcal{P}(|\mathcal{N}|)$. Then*

$$\mathcal{N}[\![\alpha]\!]^I(A) = \mathcal{N}[\![\alpha]\!]^{I[A]}(A).$$

Moreover

$$\begin{aligned}
(1) \quad \mathcal{N}[\![rx.\alpha]\!]^I(A) &= \mu B.\mathcal{N}[\![\alpha]\!]^{I[x \mapsto \bar{B}]}(A) \\
(2) \quad \mathcal{N}[\![\lambda x.\alpha]\!]^I(A) &= \nu B.\mathcal{N}[\![\alpha]\!]^{I[x \mapsto \bar{B}]}(A)
\end{aligned}$$

PROOF. We prove this by structural induction on formulas and games simultaneous. in which the duality operator a^d is only applied to atomic modalities a .

The interesting cases are for where α is a test, a composition or a fixpoint.

(1) If the game is of the form $\varphi?$ then

$$\mathcal{N}[\![\varphi?]\!]^I(A) = \mathcal{N}[\![\varphi]\!]^I \cap A = \mathcal{N}[\![\varphi]\!]^{I[A]} \cap A = \mathcal{N}[\![\varphi]\!]^{I[A]}(A),$$

since φ does not have any free variables by definition of right-linear games.

(2) If the game is of the form $\alpha; \beta$ then

$$\begin{aligned}
\mathcal{N}[\![\alpha; \beta]\!]^I(A) &= \mathcal{N}[\![\alpha]\!]^I(A) \circ \mathcal{N}[\![\beta]\!]^I(A) \\
&= \mathcal{N}[\![\alpha]\!]^{I[A]}(\mathcal{N}[\![\beta]\!]^I(A)) \\
&= \mathcal{N}[\![\alpha]\!]^{I[A]}(\mathcal{N}[\![\beta]\!]^{I[A]}(A)) \\
&= \mathcal{N}[\![\alpha; \beta]\!]^{I[A]}(A)
\end{aligned}$$

where the second equality uses that α does not have any free variables and the third equality is by induction hypothesis.

(3) If the game is of the form $rx.\alpha$. We define the following monotone functions

$$\Delta_1(w) = \mathcal{N}[\![\alpha]\!]^{I[x \mapsto w]} \quad \Delta_2(w) = \mathcal{N}[\![\alpha]\!]^{I[A][x \mapsto w]}$$

and prove that they satisfy the assumptions of Lemma 2.2. Two applications of the inductive hypothesis yield

$$\begin{aligned}
\Delta_1(w)(B) &= \mathcal{N}[\![\alpha]\!]^{I[x \mapsto w]}(B) = \mathcal{N}[\![\alpha]\!]^{I[x \mapsto w][B]}(B) \\
&= \mathcal{N}[\![\alpha]\!]^{I[x \mapsto \overline{w(B)}][B]}(B) = \mathcal{N}[\![\alpha]\!]^{I[x \mapsto \overline{w(B)}}(B) \\
&= \Delta_1(\overline{w(B)})(B)
\end{aligned}$$

for all $w \in \mathcal{W}(|\mathcal{N}|)$ and all $B \in \mathcal{P}(|\mathcal{N}|)$. The case for Δ_2 is similar.

Note also that $I[A][x \mapsto \bar{B}] = I[x \mapsto \bar{B}][A]$ and hence by the induction hypothesis $\Delta_1(\bar{B})(A) = \Delta_2(\bar{B})(A)$ for all $B \in \mathcal{P}(|\mathcal{N}|)$.

By Lemma 2.2 we compute

$$\begin{aligned}
\mathcal{N}[\![\mu x.\alpha]\!]^I(A) &= \mu w.\Delta_1(w)(A) = \mu B.(\Delta_1(\bar{B})(A)) \\
&= \mu B.(\Delta_2(\bar{B})(A)) = \mathcal{N}[\![\mu x.\alpha]\!]^{I[A]}(A).
\end{aligned}$$

The case for greatest fixpoints is analogous. $\square$

LEMMA B.3. *Suppose ψ is an FLC formula with no free variables other than x and in which x is not bound. Then $\psi \equiv \psi \stackrel{\text{id}}{x} \circ x$.*

PROOF. By induction on the formula ψ . The only interesting case is for formulas of the form $\mu y.\psi$. By Lemma B.1

$$\begin{aligned} \mathcal{N}[(\mu y.\psi \frac{\text{id}}{x}) \circ x]^I(A) &= \mu B.(\mathcal{N}[\psi \frac{\text{id}}{x}]^{I[y \mapsto \overline{B}]}(\mathcal{N}[x]^I(A))) \\ &= \mu B.(\mathcal{N}[\varphi]^{I[y \mapsto \overline{B}]}(A)) \\ &= \mathcal{N}[\mu y.\psi]^I(A) \end{aligned}$$

where the second equality is by the induction hypothesis. $\square$

C DERIVED AXIOMS FOR GL_s

To facilitate proofs in GL_s we use some derived axioms, which are immediate consequences of the original axioms of GL_s for convenience.

$$\begin{aligned} (\text{W}) \quad & \langle \sim a; a^d \rangle \perp \\ (\text{C}) \quad & \langle \sim a \rangle \langle a \rangle \varphi \leftrightarrow \langle \sim a \rangle \varphi \\ (\vee) \quad & \langle \sim a \rangle (\varphi \vee \psi) \leftrightarrow \langle \sim a \rangle \varphi \vee \langle \sim a \rangle \psi \\ (\wedge) \quad & \langle \sim a \rangle (\varphi \wedge \psi) \leftrightarrow \langle \sim a \rangle \varphi \wedge \langle \sim a \rangle \psi \\ (\Theta) \quad & \langle \sim a \rangle \varphi \leftrightarrow \varphi \quad (\text{if } a, a^d \text{ not in } \varphi) \\ (\text{P}) \quad & \langle \alpha; \sim a \rangle \varphi \leftrightarrow \langle \sim a; \alpha \rangle \varphi \quad (\{a, a^d, \sim a, \sim a^d\} \text{ not in } \alpha) \\ (\approx) \quad & \langle \sim a; \sim a \rangle \varphi \leftrightarrow \langle \sim a \rangle \varphi \\ (\Re) \quad & \langle \sim a; \sim a^d \rangle \varphi \leftrightarrow \langle \sim a^d \rangle \varphi \end{aligned}$$

PROOF. Axiom W is derived from the $\sim$ instance

$$\langle \sim a \rangle \langle x \frac{a^d}{x} \rangle \perp \leftrightarrow \langle x \frac{1 \perp}{x} \rangle \perp.$$

Axiom C is exactly the instance

$$\langle \sim a \rangle \langle x \frac{a}{x} \rangle \varphi \leftrightarrow \langle x \frac{\sim a}{x} \rangle \varphi$$

of $\sim$. For $\vee$ use the $\sim$ -instance

$$\langle \sim a \rangle \langle (x \cup y) \frac{? \varphi \ ? \psi}{x \ y} \rangle \top \leftrightarrow \langle (x \cup y) \frac{\sim a; ? \varphi \ \sim a; ? \psi}{x \ y} \rangle \top$$

and similarly for $\wedge$. Axiom Θ is derived from the instance

$$\langle \sim a \rangle \langle ? \top \rangle \varphi \leftrightarrow \langle ? \top \rangle \varphi$$

or $\sim$. For axiom P use the instance

$$\langle \sim a \rangle \langle (\alpha; x) \frac{? \varphi}{x} \rangle \top \leftrightarrow \langle (\alpha; x) \frac{\sim a; ? \varphi}{x} \rangle \top$$

of $\sim$. Axiom $\approx$ and $\Re$ are instances of $\simeq$ with $\alpha \equiv x; y$. $\square$

D PROOFS

PROOF OF LEMMA 2.2. We prove the case for the least fixpoint. Let $w = \mu u.\Delta(u)$. For the $\subseteq$ inclusion we pick any $B \in \mathcal{P}(X)$ with $\Delta(\overline{B})(r) \subseteq B$ and show that $w(A) \subseteq B$. Define $u \in \mathcal{W}(X)$

$$u(C) = \begin{cases} B & \text{if } C = A \\ X & \text{otherwise.} \end{cases}$$

Note that $\Delta(u) \subseteq u$ because

$$\Delta(u)(A) = \Delta(\overline{u(A)})(A) = \Delta(\overline{B})(A) = B = u(A).$$

Hence $w \subseteq u$ and in particular $w(A) \subseteq u(A) = B$.

For the $\supseteq$ inclusion note that $\Delta(w) \subseteq w$. Hence

$$\Delta(\overline{w(A)})(A) = \Delta(w)(A) \subseteq w(A)$$

for all A . $\square$

PROOF OF LEMMA 3.4. For games of the form $\sim a$ observe

$$\begin{aligned} \mathcal{N}[(\sim a)^d]_s(A^C) &= \{(\omega, c) : (\omega, c \frac{\square}{a}) \in A^C\} \\ &= \{(\omega, c) : (\omega, \overline{c \frac{\square}{a}}) \notin A\} \\ &= \{(\omega, c) : (\omega, c \frac{\square}{a}) \in A\}^C = (\mathcal{N}[\sim a]_s(A))^C \end{aligned}$$

The general claim is proved by mutual induction on the definition of formulas and games. We do the interesting case for games of the form α^* . Observe

$$\begin{aligned} \mathcal{N}[\alpha^*]_s^D(A) &= \mathcal{N}[\alpha^*]_s(A^C)^C = (\mathcal{N}[\alpha]_s(\mathcal{N}[\alpha^*]_s(A^C)) \cup A^C)^C \\ &= \mathcal{N}[\alpha^d]_s(\mathcal{N}[\alpha^*]_s(A^C)^C) \cap A \\ &= \mathcal{N}[\alpha^d]_s(\mathcal{N}[\alpha^*]_s^D(A)) \cap A \end{aligned}$$

Hence by maximality $\mathcal{N}[\alpha^*]_s^D(A) \subseteq \mathcal{N}[\alpha^d]_s(A)$.

For the reverse inclusion note

$$\begin{aligned} \mathcal{N}[\alpha^d]_s(A)^C &= (\mathcal{N}[\alpha^d]_s(\mathcal{N}[\alpha^d]_s(A)) \cap A)^C \\ &= \mathcal{N}[\alpha]_s(\mathcal{N}[\alpha^d]_s(A^C)^C) \cup A^C \\ &= \mathcal{N}[\alpha]_s(\mathcal{N}[\alpha^d]_s^D(A)) \cup A^C \end{aligned}$$

Hence by minimality $\mathcal{N}[\alpha]_s(A^C) \subseteq \mathcal{N}[\alpha^d]_s(A)^C$. The required $\mathcal{N}[\alpha^*]_s^D(A) \supseteq \mathcal{N}[\alpha^d]_s(A)$ follows by taking C . $\square$

PROOF OF LEMMA 3.5. We do the cases for tests $? \varphi$ and fixpoints $\text{rx}.\alpha$ in a mutual induction on formulas and games explicitly. For a test

$$\begin{aligned} (\mathcal{N}[? \varphi]^{I^d})(A) &= |\mathcal{N}| \setminus (\mathcal{N}[? \varphi]^{I^d}(|\mathcal{N}| \setminus A)) \\ &= |\mathcal{N}| \setminus (\mathcal{N}[\varphi]^{I^d} \cap |\mathcal{N}| \setminus A) \\ &= |\mathcal{N}| \setminus (\mathcal{N}[\varphi]^{I^d}) \cup A \\ &= \mathcal{N}[\overline{? \varphi}]^I(A) \end{aligned}$$

The last equality holds, since any formula appearing in a test does not have free variables by definition.

Consider games of the form $\text{rx}.\alpha$.

$$\begin{aligned} (\mathcal{N}[\text{rx}.\alpha]^{I^d})^d &= (\mu w.\mathcal{N}[\alpha]^{I^d[x \mapsto w]})^d \\ &= v w.(\mathcal{N}[\alpha]^{I[x \mapsto w]^d})^d \\ &= v w.\mathcal{N}[\alpha^d]^{I[x \mapsto w]} \\ &= \mathcal{N}[(\text{rx}.\alpha)^d]^{I[x \mapsto w]} \end{aligned} \quad \square$$

PROOF OF PROPOSITION 5.2. This is shown by a mutual induction on formulas φ and games α of RGL. Most cases of the induction are straightforward. We do the interesting cases.

First we consider formulas of the form $\langle \alpha \rangle \psi$.

$$\begin{aligned} \mathcal{N}[\langle \alpha \rangle \psi]^I(A) &= \mathcal{N}[\alpha^b \frac{\psi^b}{u,v}]^I(A) = \mathcal{N}[\alpha^b]^I[u,v \mapsto \mathcal{N}[\psi^b]^I](A) \\ &= \mathcal{N}[\alpha]^I(\mathcal{N}[\psi^b]^I(A)) = \mathcal{N}[\alpha]^I(\mathcal{N}[\varphi]^I) \\ &= \mathcal{N}[\langle \alpha \rangle \psi]^I \end{aligned}$$

The well-namedness assumption is used to ensure that the substitution does not capture free variables. The case for the games of the form $\alpha; \beta$ is very similar, as it is a similar composition.

For a game of the form x observe

$$\mathcal{N}[x]^I \circ w = I(x) \circ w = \mathcal{N}[x \circ v]^I[u,v \mapsto w] = \mathcal{N}[x^b]^I[u,v \mapsto w].$$

Finally we also consider the case of games $rx.\alpha$. with the induction hypothesis we compute

$$\begin{aligned} \mathcal{N}[(rx.\alpha)^b]^I[u,v \mapsto w] &= \mathcal{N}[rx.\alpha \frac{id^b}{u,v} \circ u]^I[u,v \mapsto w] \\ &= (\mu w. \mathcal{N}[\varphi^b \frac{id^b}{u,v}]^I[u,v \mapsto id][x \mapsto w]) \circ w \\ &= (\mu w. \mathcal{N}[\varphi^b]^I[u,v \mapsto id][x \mapsto w]) \circ w \\ &= (\mu w. \mathcal{N}[\varphi]^I[x \mapsto w] \circ id) \circ w \\ &= \mathcal{N}[rx.\alpha]^I \circ w \quad \square \end{aligned}$$

PROOF OF PROPOSITION 5.4. For the purposes of this proof we call a valuation I *constant* if $I(x)$ is constant for all x except u and v . We prove by structural induction on *all* well-named normal form rGL formula φ all well-named rGL games α that

$$\mathcal{N}[\varphi^d]^I = \mathcal{N}[\varphi^b]^I \quad \text{and} \quad \mathcal{N}[\alpha^d]^I = \mathcal{N}[\alpha^b]^I.$$

for all constant valuations I .

Most cases of the induction are straightforward. For formulas of the form $\langle \alpha \rangle \varphi$

$$\begin{aligned} \mathcal{N}[\langle \alpha \rangle \varphi]^d]^I &= \mathcal{N}[\alpha^d \frac{\varphi^d}{u}]^I = \mathcal{N}[\alpha^d]^I[u,v \mapsto \mathcal{N}[\varphi^d]^I] \\ &= \mathcal{N}[\alpha^b]^I[u,v \mapsto \mathcal{N}[\varphi^b]^I] = \mathcal{N}[\langle \alpha \rangle \varphi]^b]^I \end{aligned}$$

The well-namedness property of the formula ensures that α does not bind a variable that is free in φ , so that the substitution above does not capture variables. For variables x note that

$$\mathcal{N}[x^d]^I = I(x) = I(x) \circ I(v) = \mathcal{N}[x \circ v]^I = \mathcal{N}[x^b]^I$$

because I is constant. The argument for games of the form $\alpha; \beta$ is similar.

Finally we also consider the case of games $vx.\alpha$. Using Lemmas B.1 and B.2 we can compute the fixpoint pointwise

$$\begin{aligned} \mathcal{N}[(vx.\alpha)^d]^I(A) &= \mathcal{N}[vx.\alpha^d]^I(A) = \mu B.(\mathcal{N}[\alpha^d]^I[x \mapsto \bar{B}](A)) \\ &= \mu B.(\mathcal{N}[\alpha^b]^I[x \mapsto \bar{B}](A)) \\ &= \mu B.(\mathcal{N}[\alpha]^I[x \mapsto \bar{B}](I(u)(A))) \\ &= \mathcal{N}[vx.\alpha]^I(I(u)(A)) \\ &= \mathcal{N}[(vx.\alpha)^b]^I(A) \end{aligned}$$

The third equality is by the induction hypothesis. The fourth and the sixth equalities are by (2) of Proposition 5.2.

The moreover follows with Proposition 5.2. $\square$

PROOF OF LEMMA 5.6.

(1) The only way the $\#$ translation can give rise to a non-GL formula is through translation of fixpoints. The translation of the $*$ fixpoints in L_* formulas are clearly into repetition games.

(2) First observe that if φ is separable then $\varphi \frac{\psi}{u}$ is also separable if the free variables of ψ are never bound in φ and u is not bound in φ .

We prove by simultaneous structural induction that for any formula φ and any game α of GL the translations φ^d and α^d are separable L_μ formulas.

Most cases are straightforward for formulas $\langle \alpha \rangle \varphi$ and games $\alpha; \beta$ we use the induction hypothesis and the observation above, which applies by the assumption that φ is well-named. The most interesting case is for games α is of the form $\alpha^* = rx.(\top \cup \alpha; x)$, where x is not in α . The translation is

$$(\alpha^*)^d = \mu x.(u \vee \alpha^d \frac{x}{u}).$$

Because α has no free variables as a GL game only u is free in α^d . Hence the translation is separable.

(3) We prove by structural induction on a formula of the modal μ -calculus that provided it is separable, there is an equivalent L_* formula. The only interesting case is for fixpoint formulas. Consider a separable formula φ . Pick separable formulas ψ, ρ such that $\varphi \leftrightarrow \mu x.(\rho \vee \psi)$ where x is not free in ρ , only x is free in ψ . By renaming we ensure that x is not bound in ψ . Note that by Lemma B.3 semantically $\psi \equiv \psi \frac{id}{x} \circ x$.

By the induction hypothesis pick L_* formulas ψ', ρ' equivalent to ψ, ρ respectively. We claim that φ is equivalent to the L_* formula $(\psi' \frac{id}{x})^* \circ \rho'$. Semantically

$$\begin{aligned} (\psi' \frac{id}{x})^* \circ \rho' &\equiv \mu x.(id \vee \psi' \frac{id}{x} \circ x) \circ \rho \\ &\equiv \mu x.(id \vee \psi' \frac{id}{x} \circ x) \circ \rho \\ &\equiv \mu x.(id \vee \psi) \circ \rho \end{aligned}$$

Because $id \vee \psi$ is a modal μ -calculus formula without composition by Lemma B.3 we compute

$$\begin{aligned} &\mathcal{N}[\mu x.(id \vee \psi) \circ \rho]^I(A) \\ &= \mu B.(\mathcal{N}[id \vee \psi]^I[x \mapsto \bar{B}](\mathcal{N}[\rho]^I(A))) \\ &= \mu B.(\mathcal{N}[\rho]^I(A) \cup \mathcal{N}[\psi \frac{id}{x} \circ x]^I[x \mapsto \bar{B}](\mathcal{N}[\rho]^I(A))) \\ &= \mu B.(\mathcal{N}[\rho]^I[x \mapsto \bar{B}](A) \cup \mathcal{N}[\psi \frac{id}{x} \circ x]^I[x \mapsto \bar{B}](A)) \\ &= \mu B.(\mathcal{N}[\rho \vee \psi]^I[x \mapsto \bar{B}](A)) \\ &= \mathcal{N}[\mu x.(\rho \vee \psi)]^I(A) \end{aligned}$$

The third equality holds because x is not free in ρ and by the fact that $\mathcal{N}[x]^I[x \mapsto \bar{B}](C) = B$ is constant. $\square$

PROOF OF LEMMA 6.5. Soundness of the common part of the proof calculus goes through exactly as for game logic [36].

We say a set $A \subseteq |\mathcal{N}| \times C$ is a -invariant if

$$\mathcal{N}[\sim a]^s(A) = \mathcal{N}[\sim a^d]^s(A) = A.$$

Observe that $\mathcal{N}[\varphi]_s$ is $\sim a$ -invariant as it does not mention a , and a^d . Hence it suffices to show

$$\mathcal{N}[\sim a]_s(\mathcal{N}[\alpha \frac{a;\beta}{x} \frac{a^d;\tilde{\gamma}}{y} \frac{\tilde{\delta}}{z}]_s(A)) = \mathcal{N}[\alpha \frac{a;\beta}{x} \frac{1;\perp}{y} \frac{\sim a;\tilde{\delta}}{z}]_s(A)$$

for all $\sim a$ -invariant sets A and all α satisfying $\dagger$. We prove this by induction on α . Most cases of the induction are routine. If α is a loop this is immediate by $\sim a$ -invariance, since α does not contain variables.

For $\cup$ let

$$E = |\mathcal{N}| \times \bigcup_{i=1}^n \{c \in C : c(a_i) = \diamond, \forall k \neq i \ c(a_k) = \square, \\ \forall k \leq j \ c(b_k) = \diamond, \forall k > j \ c(b_k) = \square\}$$

where $\beta_i \equiv \sim b_1^i; \dots \sim b_j^i; \sim b_{j+1}^i; \dots \sim b_m^i$. It is easy to see by induction on α that

$$\begin{aligned} & \mathcal{N}[\alpha \frac{(a_1;\beta_1 \cup \dots \cup a_n;\beta_n);\beta}{w}]_s(U) \cap E \\ &= \mathcal{N}[\alpha \frac{(a_1;\beta_1 \cup \dots \cup a_n;\beta_n);\beta}{w}]_s(U \cap E) \cap E \\ &= \mathcal{N}[\alpha \frac{\beta}{w}]_s(U \cap E) \cap E \\ &= \mathcal{N}[\alpha \frac{\beta}{w}]_s(U) \cap E \end{aligned}$$

for all U .

For $\approx$ observe by induction on α that $\mathcal{N}[\alpha \frac{\tilde{\gamma};\tilde{\delta};\beta}{x}]_s$ is a function that maps a -invariant sets to a invariant sets. Again by induction on α now observe that it equals $\mathcal{N}[\alpha \frac{2\top;\tilde{\delta};\beta}{x}]_s$ on a -invariant sets.

For $\sim_1$ let $E = \{U \subseteq |\mathcal{N}| \times C : \text{if } (\omega, c) \in U \text{ and } c(a) = \square \text{ then } (\omega, c \frac{\diamond}{b} \in U) \Leftrightarrow (\omega, c \frac{\square}{b} \in U)\}$. By induction on α simultaneously prove that

$$\mathcal{N}[\alpha \frac{a;\beta}{x} \frac{\sim a^d}{y}]_s = \mathcal{N}[\alpha \frac{a;\beta}{x} \frac{\sim a^d;\eta}{y}]_s$$

as functions from E to itself.

For $\approx$ let $E = |\mathcal{N}| \times \{c \in C : c(a) \neq \emptyset\}$. By induction on α it is straightforward to prove the following equalities

$$\begin{aligned} \mathcal{N}[\alpha \frac{a}{x}]_s(U) \cap E &= \mathcal{N}[\alpha \frac{a}{x}]_s(U \cap E) \cap E \\ &= \mathcal{N}[\alpha \frac{a \sim a}{x}]_s(U \cap E) \cap E \\ &= \mathcal{N}[\alpha \frac{a \sim a}{x}]_s(U) \cap E \end{aligned}$$

for all U . $\square$

PROOF OF LEMMA 6.6. For (1) we prove only the backward implication. The forward direction will later follow from Theorem 6.8 and corollary 5.3 and unlike the backward implication is not required for the proof of Theorem 6.8. We prove the implication for all ρ by induction on the rank of φ . We distinguish based on the shape of the formula.

If φ is a proposition constant of the form P then we can prove

$$\text{rGL} \vdash \langle ?P; !\perp \rangle \rightarrow P$$

by $;$, $?$ and $!$. Similarly for $\neg P$. If φ is a conjunction the equivalence is an instance of $\cup$. If φ is of the form $\langle \alpha \rangle \psi$ we distinguish on the shape of α :

(1) Case a . The desired implication

$$\text{rGL} \vdash \langle a \rangle \psi^{\dagger\#} \rightarrow \langle a \rangle \psi$$

is derivable with an application of M_G from the induction hypothesis $\text{rGL} \vdash \psi^{\dagger\#} \rightarrow \psi$. Similarly for a^d and variables x .

(2) Case $\beta; \gamma$. Then by induction hypothesis since $\langle \beta \rangle \langle \gamma \rangle \psi$ is of lower rank

$$\text{rGL} \vdash (\langle \beta \rangle \langle \gamma \rangle \psi)^{\dagger\#} \rightarrow \langle \beta \rangle \langle \gamma \rangle \psi.$$

Since $(\langle \beta \rangle \langle \gamma \rangle \psi)^{\dagger} \equiv (\langle \beta; \gamma \rangle \psi)^{\dagger}$ the implication is derivable with a use of $;$.

(3) Case $?\rho$. Observe that $(\langle ?\rho \rangle \psi)^{\dagger} = (\rho \wedge \psi)^{\dagger}$ and $\text{rGL} \vdash \langle ?\rho \rangle \psi \leftrightarrow \rho \wedge \psi$. Hence the implication follows by induction hypothesis on the lower rank formula $\rho \wedge \psi$. Analogously for $!\rho$.

(4) Case $\beta \cup \gamma$. The induction hypothesis on the lower rank formula $\langle \alpha \rangle \psi \wedge \langle \beta \rangle \psi$ can be used by $\cup$ because $(\langle \alpha \rangle \psi \wedge \langle \beta \rangle \psi)^{\dagger} \leftrightarrow (\langle \alpha \cup \beta \rangle \psi)^{\dagger}$. The case $\beta \cap \gamma$ is similar.

(5) Case $rx.\alpha$. Note that by the inductive hypothesis applied to the lower rank formula $\langle \alpha \rangle \psi$

$$\text{rGL} \vdash \langle \alpha^{\dagger\#} \frac{\psi^{\dagger\#}}{u} \rangle_{\perp} \rightarrow \langle \alpha \rangle \psi$$

By substituting x by $rx.\alpha; ?\psi; !\perp$

$$\text{rGL} \vdash \langle \alpha^{\dagger\#} \frac{\psi^{\dagger\#}}{u} \frac{rx.\alpha; ?\psi; !\perp}{x} \rangle_{\perp} \rightarrow \langle \alpha \frac{rx.\alpha; ?\psi; !\perp}{x} \rangle \psi$$

and an application of RL and fp_G yields

$$\text{rGL} \vdash \langle \alpha^{\dagger\#} \frac{\psi^{\dagger\#}}{u} \frac{rx.\alpha; ?\psi; !\perp}{x} \rangle_{\perp} \rightarrow \langle rx.\alpha \rangle \psi$$

The implication follows with an application of μ_G , since the translation of $\langle rx.\alpha \rangle \psi$ is

$$(\langle rx.\alpha \rangle \psi)^{\dagger\#} \equiv \langle \mu x.\alpha^{\dagger\#} \frac{\psi^{\dagger\#}}{u} \rangle_{\perp}.$$

(6) Case $\lambda x.\alpha$. By definition of the translations

$$\text{rGL} \vdash (\langle \lambda x.\alpha \rangle \psi)^{\dagger\#} \rightarrow \langle \alpha^{\dagger\#} \frac{\psi^{\dagger\#}}{u} \frac{(\langle \lambda x.\alpha \rangle \psi)^{\dagger\#}}{x} \rangle_{\perp}$$

is a consequence of rfp_G and consequently

$$\text{rGL} \vdash (\langle \lambda x.\alpha \rangle \psi)^{\dagger\#} \rightarrow \langle \alpha^{\dagger\#} \frac{\psi^{\dagger\#}}{u} \frac{(\langle \lambda x.\alpha \rangle \psi)^{\dagger\#}}{x} \frac{?!\perp; !\perp}{x} \rangle_{\perp}$$

by RL. By the induction hypothesis applied to α and substituting x by $(\langle \lambda x.\alpha \rangle \psi)^{\dagger\#}; ?!\perp; !\perp$ this implies

$$\text{rGL} \vdash (\langle \lambda x.\alpha \rangle \psi)^{\dagger\#} \rightarrow \langle \alpha \frac{(\langle \lambda x.\alpha \rangle \psi)^{\dagger\#}}{x} \rangle \psi$$

Hence by ν_G the implication follows.

The equivalence in (2) follows immediately from Proposition 6.1 and Corollary 5.3 $\square$

LEMMA D.1. $\text{rGL} \vdash \overline{\varphi^{\dagger\#}} \leftrightarrow \overline{\varphi^{\dagger\#}}$ for φ a closed L_{μ} formula.

PROOF OF LEMMA D.1. We prove the equivalence by induction on the formula φ .

- (1) If the formula is of the form id , then we need to prove the equivalence

$$\text{rGL} \vdash \langle ?\top; !\perp \rangle \perp \leftrightarrow \langle !\perp \rangle \top$$

This is immediate from $;$, $\cup$, $\cap$

- (2) If the formula is of the form P we prove the equivalence

$$\text{rGL} \vdash \langle ?\neg P; !\perp \rangle \perp \leftrightarrow \langle !P; ?\perp \rangle \top$$

by $;$, $?$ and $!$.

- (3) If the formula is of the form $\varphi \vee \psi$ the equivalence follows from the induction hypothesis with $\cap$. Similarly for formulas of the form $\varphi \wedge \psi$
- (4) If the formula is of the form $\langle a \rangle \varphi$ the equivalence follows from the induction hypothesis with $;$ and an application of M_G .
- (5) If the formula is of the form $\mu x. \psi$. The equivalence we need to prove is $\text{rGL} \vdash \langle \lambda x. \bar{\psi}^\# \rangle \perp \leftrightarrow \langle \lambda x. \psi^\# \rangle \top$.

For the forward direction $\text{rGL} \vdash \langle \lambda x. \bar{\psi}^\# \rangle \perp \rightarrow \langle \bar{\psi}^\# \frac{\lambda x. \bar{\psi}^\#}{x} \rangle \perp$ is an instance of rfp_G . By induction hypothesis

$$\text{rGL} \vdash \langle \lambda x. \bar{\psi}^\# \rangle \perp \rightarrow \langle \bar{\psi}^\# \frac{\lambda x. \bar{\psi}^\#}{x} \rangle \top$$

An application of RL allows to deduce the desired implication with ν_G . The reverse implication is analogous. $\square$

LEMMA D.2. For closed L_μ formulas φ , if $\text{mL}_\mu \vdash \varphi$ then $\text{rGL} \vdash \varphi^\#$.

PROOF OF LEMMA D.2. First note that

$$\text{rGL} \vdash (\varphi \vee \psi)^\# \leftrightarrow \varphi^\# \vee \psi^\# \quad \text{rGL} \vdash (\varphi \wedge \psi)^\# \leftrightarrow \varphi^\# \wedge \psi^\#$$

Indeed these are instances of $\cup$ and $\cap$ respectively. By Lemma D.1 $\text{rGL} \vdash \neg P^\# \leftrightarrow (\neg P)^\#$. Thus if φ is a propositional tautology $\text{rGL} \vdash \varphi^\#$ is provably reducible to a propositional tautology and therefore provable in $\text{rGL} \vdash$.

Note also by Lemma D.1

$$(\#) (\varphi \rightarrow \psi)^\# \leftrightarrow (\varphi^\# \rightarrow \psi^\#)$$

is an instance of $\cup$ after expanding the abbreviations.

From $\#$ it follows that the $\#$ -translation of any instance of fp is provably equivalent to an instance of fp_G . Hence the translations of all axioms of the monotone modal μ -calculus are provable in right-linear game logic. We now proceed to show the claim by induction on the length of the proof.

If the last step of the proof of $\text{mL}_\mu \vdash \psi$ is an instance of MP of the kind

$$M_a \frac{\varphi \rightarrow \psi}{\psi} \quad \text{then} \quad M_G \frac{\varphi^\# \rightarrow \psi^\#}{\psi^\#}$$

is an instance of MP_G . Because $\#$ provably distributes over implications we can derive

$$M_G \frac{\varphi^\# \quad \# \text{MP}_G \frac{\varphi \rightarrow \psi}{\varphi^\# \rightarrow \psi^\#}}{\psi^\#}$$

If the last step of a proof of $\text{mL}_\mu \vdash \langle a \rangle \varphi \rightarrow \langle a \rangle \psi$ is an instance of rule M_a of the kind

$$M_a \frac{\varphi \rightarrow \psi}{\langle a \rangle \varphi \rightarrow \langle a \rangle \psi} \quad \text{then} \quad M_G \frac{\varphi^\# \rightarrow \psi^\#}{\langle a \rangle \varphi^\# \rightarrow \langle a \rangle \psi^\#}$$

is an instance of M_G . We can derive

$$\begin{array}{c} \# \text{MP}_G \frac{(\varphi \rightarrow \psi)^\#}{\varphi^\# \rightarrow \psi^\#} \\ M_G \frac{\varphi^\# \rightarrow \psi^\#}{\langle a \rangle \varphi^\# \rightarrow \langle a \rangle \psi^\#} \\ \# \text{MP}_G \frac{\langle a \rangle \varphi^\# \rightarrow \langle a \rangle \psi^\#}{(\langle a \rangle \varphi)^\# \rightarrow (\langle a \rangle \psi)^\#} \\ \# \text{MP}_G \frac{(\langle a \rangle \varphi)^\# \rightarrow (\langle a \rangle \psi)^\#}{(\langle a \rangle \varphi \rightarrow \langle a \rangle \psi)^\#} \end{array}$$

The induction hypothesis yields a proof for $\text{rGL} \vdash (\langle a \rangle \varphi \rightarrow \langle a \rangle \psi)^\#$. If the last step of the proof is an instance of μ of the kind

$$\mu \frac{\varphi \frac{\psi}{x} \rightarrow \psi}{\mu x. \varphi \rightarrow \psi}$$

then by the induction hypothesis and $\#$ observe $\text{rGL} \vdash \langle \varphi^\# \frac{\psi^\#}{x} \rangle \perp \rightarrow \langle \psi^\# \rangle \perp$. By an application of RL deduce $\text{rGL} \vdash \langle \varphi^\# \frac{\psi^\#; ?\perp; !\perp}{x} \rangle \perp \rightarrow \langle \psi^\# \rangle \perp$. The implication $\text{rGL} \vdash (\mu x. \varphi \rightarrow \psi)^\#$ follows by μ_G with $\#$. $\square$

LEMMA D.3 (SUBSTITUTION). If $\alpha, \beta_1, \dots, \beta_n$ are rGL games such that β_i does not mention freely any variable x such that x_i appears in α in a context where x is bound, then

$$\mathcal{N}[\![\alpha]\!]^I \frac{\mathcal{N}[\![\beta_i]\!]}{x_i} = \mathcal{N}[\![\alpha \frac{\beta_i}{x_i}]\!]^I.$$

PROOF OF LEMMA D.3. By a straightforward induction on α . $\square$

PROOF OF PROPOSITION 7.2. By simultaneous induction on formulas and games of sabotage game logic. For tests this uses the observation that the translation φ^c of a formula does not have free variables.

We consider the interesting cases. For formulas of the form $\langle \alpha \rangle \varphi$ note that by the induction hypothesis

$$\mathcal{N}[\![\langle \alpha \rangle \varphi]\!]_s \upharpoonright c = \mathcal{N}[\![\alpha^c]\!]^I(\emptyset)$$

where $I(y_e) = \mathcal{N}[\![\varphi^c]\!](\emptyset) = \mathcal{N}[\![\varphi^c; ?\perp]\!]$. Hence $\mathcal{N}[\![\langle \alpha \rangle \varphi]\!]_s \upharpoonright c = \mathcal{N}[\![\alpha^c \frac{\varphi^c; ?\perp}{y_e}]\!](\emptyset)$ by Lemma D.3 as required. The case for games of the kind $\alpha; \beta$ is similar.

We also explicitly treat games of the form α^* . Let $a_1, \dots, a_m$ be the list of atomic games in α such that either $c(a) \neq \emptyset$ or $\sim a$ or $\sim a^d$ appears in α . List $c_1, \dots, c_m$ all possible contexts that satisfy $c_i(a) = \emptyset$ for all $a \notin \{a_1, \dots, a_m\}$.

Consider first the $\subseteq$ inclusion. Fix $B_i = \mathcal{N}[\![\alpha^*]^{c_i}]\!]^I(\emptyset)$. By definition of the translation, Lemma D.3 the inductive hypothesis

$$\begin{aligned} B_i &= \mathcal{N}[\![y_i \cup \alpha^{c_i} \frac{(\alpha^*)^{c_i}}{y_i}]\!]^I(\emptyset) \\ &= U \upharpoonright c_i \cup \mathcal{N}[\![\alpha^{c_i}]\!]^I \frac{B_i}{y_i}(\emptyset) \\ &= U \upharpoonright c_i \cup \mathcal{N}[\![\alpha]\!]_s(B) \upharpoonright c_i \end{aligned}$$

where $B = \bigcup_{i=1}^m B_i \times \{c_i\}$. Hence $\mathcal{N}[\![\alpha^*]\!]_s(U) \upharpoonright c_i \subseteq B_i$ follows by pointwise minimality.

Consider next the $\supseteq$ inclusion. Define $A_i = \mathcal{N}[\alpha^*]_s(U) \upharpoonright c_i$. By induction hypothesis

$$\begin{aligned} A_i &= U \upharpoonright c_i \cup \mathcal{N}[\alpha]_s(\mathcal{N}[\alpha^*]_s(U)) \upharpoonright c_i \\ &= U \upharpoonright c_i \cup \mathcal{N}[\alpha]^{I_i}(\emptyset) \end{aligned}$$

where $I_i(y_{c_i}) = \mathcal{N}[\alpha^*]_s(U) \upharpoonright c_i$ for all i . The $\supseteq$ inclusion follows from Theorem 7.1 by minimality. $\square$

PROOF OF PROPOSITION 7.3. Note that the proof of Lemma 7.6 does not rely on this lemma. For convenience we therefore freely use Lemma 7.6 in this proof.

The identity is proved by structural induction on the rank of a closed formula. Most cases are straightforward. The interesting cases are for least and greatest fixpoint formulas. So consider a formula $\text{rx}.\alpha$. Observe that for $Z = \mathcal{N}[\text{rx}.\alpha^h]_s$ by (2) of Lemma 7.6 and soundness

$$\mathcal{N}[\alpha^h]_s^{I[x \mapsto Z]} \subseteq Z$$

Hence by induction hypothesis and minimality $\mathcal{N}[\text{rx}.\alpha] \subseteq Z$. Analogously compute

$$\mathcal{N}[\alpha^h]_s^{I[x \mapsto W]} = \mathcal{N}[\alpha]^{I[x \mapsto W]} \subseteq W$$

where $W = \mathcal{N}[\text{rx}.\alpha] \times C$. The reverse inclusion $\mathcal{N}[\varphi] \supseteq Z$ follows by minimality of Z .

The case for greatest fixpoints follows from the least fixpoint with (1) of Lemma 7.6. $\square$

LEMMA D.4. If α is a game of right-linear game logic, $\delta \equiv \sim b_x^d; \sim c_x$ are fresh and β is a GL_s game, then

$$\text{GL}_s \vdash \langle \alpha^h \frac{\delta; \beta}{x} \rangle \varphi \leftrightarrow \langle \delta^d; \alpha^h \frac{\delta}{x}; (c_x; \delta; \beta \cup b_x) \rangle \varphi$$

PROOF OF LEMMA D.4. For the purposes of this proof we call a GL_s game α with free variables *controlled* if it is right-linear and it contains iteration games only in the form

$$\eta; (c; \eta^d; \alpha \frac{\eta}{x})^*; b \quad \text{and} \quad \eta^d; (c^d; \eta; \alpha \frac{\eta^d}{x})^\times; b^d$$

for $\eta \equiv \sim b^d; \sim c$, where neither $b, c, \sim b, \sim c$ nor their duals appear in α .

Fix a fresh variable y and first define for every controlled GL_s game α a modification α_U with respect to a set $U \subseteq \mathbb{V}$ by induction on α as follows

$$\begin{aligned} w_U &= \begin{cases} x & \text{if } x \in U \\ w; y & \text{otherwise} \end{cases} & (\alpha; \beta)_U &= \alpha; (\beta)_U \\ a_U &= a; y & (a^d)_U &= a^d; y \\ (\sim a)_U &= \sim a; y & (\sim a^d)_U &= \sim a^d; y \\ (? \varphi)_U &= ? \varphi; y & (! \varphi)_U &= ! \varphi; y \\ (\alpha \cup \beta)_U &= \alpha_U \cup \beta_U & (\alpha \cap \beta)_U &= \alpha_U \cap \beta_U \end{aligned}$$

and for controlled loop games:

$$\begin{aligned} (\eta; (c; \eta^d; \alpha \frac{\eta}{x})^*; b)_U &= \eta; (c; \eta^d; \alpha_{U \cup \{x\}} \frac{\eta}{x})^*; b \\ (\eta^d; (c^d; \eta; \alpha \frac{\eta^d}{x})^\times; b^d)_U &= \eta^d; (c^d; \eta; \alpha_{U \cup \{x\}} \frac{\eta^d}{x})^\times; b^d \end{aligned}$$

where $\eta \equiv \sim b^d; \sim c$.

By induction on the definition observe that

$$\text{GL}_s \vdash \langle \alpha_U \frac{\sigma; \gamma}{u} \frac{\gamma}{y} \rangle \varphi \leftrightarrow \langle \alpha \frac{\sigma; \gamma}{u}; \gamma \rangle \varphi \quad (1)$$

for all right-linear GL_s σ, γ , where u ranges over all elements of U .

The only interesting case is for controlled loop games. Before we give a proof we show that the following pairwise equivalences are provable for any ψ

$$\begin{aligned} &\langle \eta^d; \alpha_{U \cup \{x\}} \frac{\eta}{x} \frac{\sigma; \gamma}{u} \frac{\gamma}{y} \rangle \psi \\ \text{iff} \quad &\langle \eta^d; \alpha_{U \cup \{x\}} \frac{\eta}{x} \frac{\eta^d; \sigma; \gamma}{u} \frac{\eta^d; \gamma}{y} \rangle \psi \\ \text{iff} \quad &\langle \eta^d; \alpha_{U \cup \{x\}} \frac{\eta; \tilde{\gamma}}{x} \frac{\eta^d; \sigma; \tilde{\gamma}}{u} \frac{\eta^d; \tilde{\gamma}}{y} \rangle \psi \\ \text{iff} \quad &\langle \eta^d; \alpha \frac{\eta}{x} \frac{\eta^d; \sigma}{u}; \tilde{\gamma} \rangle \psi \end{aligned}$$

where $\tilde{\gamma} \equiv (c \cup b; \gamma)$. The first two equivalence use $\sim$ and $\wr$ and the third is by induction hypothesis and using Θ since η sets fresh variables not in σ . We now turn to the proof of the case of (1) for controlled repetition. This amounts to proving

$$\text{GL}_s \vdash \langle \eta; (c; \eta^d; \alpha_{U \cup \{x\}} \frac{\eta}{x} \frac{\sigma; \gamma}{u} \frac{\gamma}{y})^*; b \rangle \varphi \leftrightarrow \langle \eta; (c; \eta^d; \alpha \frac{\eta}{x} \frac{\sigma}{u})^*; b; \gamma \rangle \varphi \quad (2)$$

We first prove the forward implication. For this purpose let $\rho = \langle c; (c; \eta^d; \alpha \frac{\eta}{x} \frac{\sigma}{u})^*; b; \gamma \rangle \varphi \vee \langle b \rangle \varphi$. The following chain of equivalences is provable:

$$\begin{aligned} &\langle c; c; \eta^d; \alpha_{U \cup \{x\}} \frac{\eta}{x} \frac{\sigma; \gamma}{u} \frac{\gamma}{y} \rangle \rho \\ \text{then} \quad &\langle c; c; \eta^d; \alpha \frac{\eta}{x} \frac{\eta^d; \sigma}{u}; \tilde{\gamma} \rangle \rho \\ \text{then} \quad &\langle c; c; \eta^d; \alpha_{U \cup \{x\}} \frac{\eta; \tilde{\gamma}}{x} \frac{\eta^d; \sigma; \tilde{\gamma}}{u} \frac{\eta^d; \tilde{\gamma}}{y} \rangle \rho \\ \text{then} \quad &\langle c; c; \eta^d; \alpha_{U \cup \{x\}} \frac{\eta; \tilde{\rho}_0}{x} \frac{\eta^d; \sigma; \tilde{\rho}_0}{u} \frac{\tilde{\rho}_0}{y} \rangle \rho \\ \text{then} \quad &\langle c; c; \eta^d; \alpha \frac{\eta}{x} \frac{\eta^d; \sigma}{u}; \tilde{\rho}_0 \rangle \rho \\ \text{then} \quad &\langle c; \tilde{\rho}_0 \rangle \rho \\ \text{then} \quad &\rho \end{aligned}$$

where $\tilde{\rho} \equiv (c; \tilde{\rho}_0) \cup b$ and $\tilde{\rho}_0 \equiv (c; \eta^d; \alpha \frac{\eta}{x} \frac{\sigma}{u})^*; b; \gamma$. Hence we also get:

$$\text{GL}_s \vdash \langle b \rangle \varphi \vee \langle c; c; \eta^d; \alpha_{U \cup \{x\}} \frac{\eta}{x} \frac{\sigma; \gamma}{u} \frac{\gamma}{y} \rangle \rho \rightarrow \rho$$

By applying μ_G^* this yields

$$\text{GL}_s \vdash \langle (c; c; \eta^d; \alpha_{U \cup \{x\}} \frac{\eta}{x} \frac{\sigma; \gamma}{u} \frac{\gamma}{y})^*; b \rangle \varphi \rightarrow \rho.$$

By applying M_G we obtain

$$\text{GL}_s \vdash \langle \eta; (c; c; \eta^d; \alpha_{U \cup \{x\}} \frac{\eta}{x} \frac{\sigma; \gamma}{u} \frac{\gamma}{y})^*; b \rangle \varphi \rightarrow \langle \eta \rangle \rho.$$

Applying $\cong$ we can remove the additional c and obtain:

$$\text{GL}_s \vdash \langle \eta; (c; \eta^d; \alpha_{U \cup \{x\}} \frac{\eta}{x} \frac{\sigma; \gamma}{u} \frac{\gamma}{y})^*; b \rangle \varphi \rightarrow \langle \eta \rangle \rho.$$

Finally using $\forall$, $\wr$ and $\sim$ on the succedent of the implication we get the forward implication of (2).

We next prove the backward implication (2). For this purpose let $\chi = \langle c; (c; \eta^d; \alpha_{U \cup \{x\}} \frac{\eta}{x} \frac{\sigma; \gamma}{u} \frac{\gamma}{y})^*; b \rangle \varphi \vee \langle b; \gamma \rangle \varphi$. The following chain

of equivalences is provable:

$$\begin{aligned}
& \langle c; c; \eta^d; \alpha \frac{\eta}{x} \frac{\sigma}{u} \rangle \chi \\
\text{then} \quad & \langle c; c; \eta^d; \alpha \frac{\eta}{x} \frac{\eta^d; \sigma}{u} \rangle \chi \\
\text{then} \quad & \langle c; c; \eta^d; \alpha_{U \cup \{x\}} \frac{\eta; \tilde{\chi}}{x} \frac{\eta^d; \sigma; \tilde{\chi}}{u} \frac{\eta^d; \tilde{\chi}}{y} \rangle \varphi \\
\text{then} \quad & \langle c; c; \eta^d; \alpha_{U \cup \{x\}} \frac{\eta; \tilde{\chi}_0}{x} \frac{\eta^d; \sigma; \tilde{\chi}_0}{u} \frac{\eta^d; \tilde{\chi}_0}{y} \rangle \varphi \\
\text{then} \quad & \langle c; c; \eta^d; \alpha \frac{\eta}{x} \frac{\eta^d; \sigma}{u}; \tilde{\chi}; \tilde{\chi}_0 \rangle \varphi \\
\text{then} \quad & \langle c; c; \eta^d; \alpha_{U \cup \{x\}} \frac{\eta; \tilde{\chi}}{x} \frac{\eta^d; \sigma; \tilde{\chi}}{u} \frac{\eta^d; \tilde{\chi}}{y}; \tilde{\chi}_0 \rangle \varphi \\
\text{then} \quad & \langle c; c; \eta^d; \alpha_{U \cup \{x\}} \frac{\eta}{x} \frac{\sigma; \tilde{\chi}}{u} \frac{\tilde{\chi}}{y}; \tilde{\chi}_0 \rangle \varphi \\
\text{then} \quad & \langle c; \tilde{\chi}_0 \rangle \varphi \\
\text{then} \quad & \chi
\end{aligned}$$

where $\tilde{\chi} \equiv (c; \tilde{\chi}_0) \cup (b; \gamma)$ and $\tilde{\chi}_0 \equiv (c; \eta^d; \alpha_{U \cup \{x\}} \frac{\eta}{x} \frac{\sigma; \tilde{\chi}}{u} \frac{\tilde{\chi}}{y})^*; b$. Hence

$$\text{GL}_s \vdash \langle b; \gamma \rangle \varphi \vee \langle c; c; \eta^d; \alpha \frac{\eta}{x} \frac{\sigma}{u} \rangle \chi \rightarrow \chi$$

The remainder of the proof of the backward implication of (2) is analogous to the proof of the forward implication.

Finally the pairwise equivalences of the following formulas is provable in GL_s :

$$\begin{aligned}
& \langle \delta^d; \alpha^h \frac{\delta}{x}; (c; \delta; \beta \cup b_x) \rangle \varphi \\
\text{iff} \quad & \langle \delta^d; (\alpha^h)_{\{x\}} \frac{\delta; (c; \delta; \beta \cup b_x)}{x} \frac{(c; \delta; \beta \cup b_x)}{y} \rangle \varphi \\
\text{iff} \quad & \langle (\alpha^h)_{\{x\}} \frac{\delta; (c; \delta; \beta \cup b_x)}{x} \frac{\delta^d (c; \delta; \beta \cup b_x)}{y} \rangle \varphi \\
\text{iff} \quad & \langle (\alpha^h)_{\{x\}} \frac{\delta; \beta}{x} \frac{? \top}{y} \rangle \varphi \\
\text{iff} \quad & \langle \alpha^h \frac{\delta; \beta}{x} \rangle \varphi
\end{aligned}$$

where the first and the fourth equivalence are by (1). The second and third equivalences use $\sim$, $\wr$ and $\approx$. $\square$

PROOF OF LEMMA 7.6. (1) is a straightforward structural induction.

(2) The translation of $\langle \alpha \frac{rx; \alpha}{x} \rangle \varphi$ with $\mathfrak{h}$ is

$$\langle \alpha^h \frac{\delta; (c; \delta; \delta^d; \alpha^h \frac{\delta}{x}); b_x}{x} \rangle \varphi^h$$

where $\delta \equiv \sim b_x^d; \sim c_x$. For readability we drop the variable subscript. By Lemma D.4 this provably implies

$$\langle \delta^d; \alpha^h \frac{\delta}{x}; (c; \delta; (c; \delta^d; \alpha^h \frac{\delta}{x})^*; b \cup b) \rangle \varphi^h$$

By $\cup$, $;$, $?$ and M_G this provably implies

$$\langle \delta^d; \alpha^h \frac{\delta}{x}; (c; \delta; (c; \delta^d; \alpha^h \frac{\delta}{x})^* \cup ? \top); b \rangle \varphi^h$$

By rfp_G and M_G this provably implies

$$\langle \delta^d; \alpha^h \frac{\delta}{x}; (c; \delta; c; \delta^d; \alpha^h \frac{\delta}{x}; (c; \delta^d; \alpha^h \frac{\delta}{x})^* \cup ? \top); b \rangle \varphi^h$$

By P , C and $\approx$:

$$\langle \delta^d; \alpha^h \frac{\delta}{x}; (c; \delta^d; \alpha^h \frac{\delta}{x}; (c; \delta^d; \alpha^h \frac{\delta}{x})^* \cup ? \top); b \rangle \varphi^h$$

Hence by $*_G$ and M_G

$$\langle \delta^d; \alpha^h \frac{\delta}{x}; (c; \delta^d; \alpha^h \frac{\delta}{x})^*; b \rangle \varphi^h$$

By $\approx$ and C this in turn provably implies the translation

$$\langle \delta; c; \delta^d; \alpha^h \frac{\delta}{x}; (c; \delta^d; \alpha^h \frac{\delta}{x})^*; b \rangle \varphi^h$$

By M_G and $*_G$ this in turn provably implies the translation

$$(\langle rx; \alpha \rangle \varphi)^h \equiv \langle \delta; (c; \delta^d; \alpha^h \frac{\delta}{x})^*; b \rangle \varphi^h$$

as required.

(3) Let $\rho \equiv \langle c; \delta^d; \beta \rangle \psi \vee \langle b \rangle \varphi^h$. Assume $\text{GL}_s \vdash \langle \alpha^h \frac{\beta; ? \psi; ! \perp}{x} \rangle \varphi^h \rightarrow \langle \beta \rangle \psi$. By M_G and some propositional reasoning deduce

$$\text{GL}_s \vdash \langle b \rangle \varphi^h \vee \langle c; \delta^d; \alpha^h \frac{\delta; \beta; ? \psi; ! \perp}{x} \rangle \varphi^h \rightarrow \rho$$

By Lemma D.4 (and $\approx$) this implies (by M_G and Θ)

$$\text{GL}_s \vdash \langle b \rangle \varphi^h \vee \langle c; \delta^d; \alpha^h \frac{\delta}{x}; (c; \delta; \beta; ? \psi; ! \perp \cup b) \rangle \varphi^h \rightarrow \rho$$

Since b, c do not appear in β, ψ this implies

$$\text{GL}_s \vdash \langle b \rangle \varphi^h \vee \langle c; \delta^d; \alpha^h \frac{\delta}{x} \rangle \rho \rightarrow \rho$$

By applying μ_G^* it follows that

$$\text{GL}_s \vdash \langle (c; \delta^d; \alpha^h \frac{\delta}{x})^*; b \rangle \varphi^h \rightarrow \rho$$

is provable. By M_G it follows that

$$\text{GL}_s \vdash \langle \delta; (c; \delta^d; \alpha^h \frac{\delta}{x})^*; b \rangle \varphi^h \rightarrow \langle \delta \rangle (\langle c; \delta^d; \beta \rangle \psi \vee \langle b \rangle \varphi^h)$$

By C , $\approx$, Θ , $\forall$ and W this provably implies

$$\text{GL}_s \vdash \langle \delta; (c; \delta^d; \alpha^h \frac{\delta}{x})^*; b \rangle \varphi^h \rightarrow \langle \beta \rangle \psi$$

as required. $\square$

PROOF OF LEMMA 7.7. We can without loss of generality fix a finite set U of atomic games and prove the lemma only for formulas restricted to atomic games from U . We fix for every atomic game $a \in U$ three fresh atomic games $a_\emptyset$, $a_\diamond$ and $a_\square$. The role of these games is to capture the context syntactically, that is we maintain the invariant $c(a) = i$ iff $c(a_i) = \diamond$. We say c is a U -context if $c(a) = \emptyset$ for all $a \notin U$. Let C_U be the set of all U -contexts.

Fix some context c and list all elements of U as

$$a_1, \dots, a_n, b_1, \dots, b_m, c_1, \dots, c_k$$

such that $c(a_i) = \diamond$, $c(b_i) = \square$ and $c(c_i) = \emptyset$ for all i . Define the games

$$\begin{aligned}
\eta_c & \equiv \sim a_1; \dots; \sim a_n; \sim b_1^d; \dots; \sim b_m^d \\
\zeta_c & \equiv \sim (a_1)_\emptyset^d; \sim (a_1)_\diamond; \sim (a_1)_\square^d; \dots; \sim (a_n)_\emptyset^d; \sim (a_n)_\diamond; \sim (a_n)_\square^d \\
& \quad \sim (c_1)_\emptyset^d; \sim (c_1)_\diamond; \sim (c_1)_\square; \dots; \sim (c_k)_\emptyset^d; \sim (c_k)_\diamond; \sim (c_k)_\square \\
& \quad \sim (c_1)_\emptyset; \sim (c_1)_\diamond^d; \sim (c_1)_\square^d; \dots; \sim (c_k)_\emptyset; \sim (c_k)_\diamond; \sim (c_k)_\square^d \\
\xi_c & \equiv (a_1)_\diamond; \dots; (a_n)_\diamond; (b_1)_\square; \dots; (b_m)_\square; (c_1)_\emptyset; \dots; (c_k)_\emptyset
\end{aligned}$$

For a C_U -indexed family α_c of GL_s games we define the guarded version

$$\check{\alpha} = \bigcup_{c \in C_U} (\xi_c; \zeta_c; \alpha_c)$$

For a single game α we mean by $\check{\alpha} = \check{\beta}$ where $\beta_c = \alpha$ for all $c \in C_U$.

We define for every formula φ and every game α of GL_s we define modified versions $\widehat{\varphi}$ and $\widehat{\alpha}$ by induction on the definition as follows

$$\begin{array}{ll} \widehat{P} \equiv P & \widehat{\neg P} \equiv \neg P \\ \widehat{\varphi \vee \psi} \equiv \widehat{\varphi} \vee \widehat{\psi} & \widehat{\varphi \wedge \psi} \equiv \widehat{\varphi} \wedge \widehat{\psi} \\ \widehat{\langle \alpha \rangle \varphi} \equiv \langle \widehat{\alpha} \rangle \widehat{\varphi} & \widehat{\alpha; \beta} \equiv \widehat{\alpha}; \widehat{\beta} \\ \widehat{a} \equiv a_0; a \cup a_\diamond & \widehat{a^d} \equiv a_0; a^d \cup a_\diamond; !\perp \cup a_\square \\ \widehat{\sim a} \equiv \sim a_0^d; \sim a_\diamond; \sim a_\square^d & \widehat{\sim a^d} \equiv \sim a_0^d; \sim a_\diamond^d; \sim a_\square^d \\ \widehat{? \varphi} \equiv ? \widehat{\varphi} & \widehat{! \varphi} \equiv ! \widehat{\varphi} \\ \widehat{\alpha \cup \beta} \equiv \widehat{\alpha} \cup \widehat{\beta} & \widehat{\alpha \cap \beta} \equiv \widehat{\alpha} \cap \widehat{\beta} \\ \widehat{\alpha^*} \equiv (\widehat{\alpha})^* & \widehat{\alpha^\times} \equiv (\widehat{\alpha})^\times \end{array}$$

Before we prove the lemma we make the following observations

- (1) $GL_s \vdash \langle \zeta_c \rangle \widehat{\varphi} \rightarrow \langle \eta_c \rangle \varphi$
- (2) $GL_s \vdash \langle \zeta_c \rangle \langle \widehat{\alpha}; \widehat{\beta} \rangle \top \rightarrow \langle \zeta_c \rangle \langle \widehat{\alpha}; \beta \rangle \perp$
- (3) $GL_s \vdash \langle \zeta_c \rangle (\langle \widehat{\beta} \rangle \varphi \leftrightarrow \langle \beta_c \rangle \varphi)$

for all formulas φ and all games α .

For observation (1) define first a second modification $(\cdot)^E$ exactly like $\widehat{\cdot}$ except that

$$(\sim a)^E \equiv \sim a_0^d; \sim a_\diamond; \sim a_\square^d; \sim a \quad (\sim a^d)^E \equiv \sim a_0^d; \sim a_\diamond^d; \sim a_\square; \sim a^d.$$

Then by $\sim_1$ observe $GL_s \vdash \langle \zeta_c \rangle (\widehat{\varphi} \leftrightarrow \varphi^E)$. Define a further modification $(\cdot)^F$ exactly like $(\cdot)^E$ except that

$$(a)^F \equiv (a_0 \cup a_\diamond \cup a_\square); a \quad (a^d)^F \equiv (a_0 \cup a_\diamond \cup a_\square); a^d$$

Using $\cup$ it is not hard to deduce that $GL_s \vdash \langle \zeta_c \rangle (\varphi^E \leftrightarrow \varphi^F)$. Again applying $\cup$ yields that $GL_s \vdash \langle \zeta_c \rangle (\varphi^F \leftrightarrow \varphi^G)$, where $(\cdot)^G$ exactly like $(\cdot)^E$ except that

$$(a^*)^G \equiv (\alpha^F)^* \quad (\alpha^\times)^G \equiv (\alpha^F)^\times.$$

Finally the required implication follows with $\simeq$ since no $a_0, a_\diamond, a_\square$ appears. Observation (2) can also be proved with $\cup$ and (3) is immediate.

We finally prove the main claim of the lemma. By simultaneous induction on formulas φ and games α of sabotage game logic we prove

- (1) $GL_s \vdash \langle \varphi^{c\sharp} \rangle \perp \rightarrow \langle \zeta_c \rangle \widehat{\varphi}$
- (2) $GL_s \vdash \langle \alpha^{c\sharp} \frac{\zeta_c; \beta_c}{y_c} \rangle \perp \rightarrow \langle \zeta_c \rangle \langle \widehat{\alpha}; \widehat{\beta} \rangle \perp$

for all contexts c and all games β_c . The lemma is easily deduced from the case for $c = c_0$ with observation 1.

- (1) If the formula is of the form P the implication to show is $GL_s \vdash \langle ?P; !\perp \rangle \perp \rightarrow \langle \zeta_c \rangle P$. The implication is provably by $\emptyset, ;$ and $?$. The case for $\neg P$ is analogous.
- (2) If the formula is of the form $\varphi \vee \psi$ or $\varphi \wedge \psi$ then the implication is by the induction hypothesis, $\cup, \cap, \forall$ and $\wedge$.
- (3) If the formula is of the form $\langle \alpha \rangle \psi$ we first apply the induction hypothesis to ψ and obtain

$$GL_s \vdash \langle \alpha^{c\sharp} \frac{\varphi^{c\sharp}; ?\perp}{y_c} \rangle \perp \rightarrow \langle \alpha^{c\sharp} \frac{\zeta_c; ?\widehat{\varphi}; !\perp}{y_c} \rangle \perp.$$

Now applying the induction hypothesis for α combined with (2) we get

$$GL_s \vdash \langle \alpha^{c\sharp} \frac{\zeta_c; ?\widehat{\varphi}; !\perp}{y_c} \rangle \perp \rightarrow \langle \zeta_c \rangle \langle \widehat{\alpha}; ?\widehat{\varphi}; !\perp \rangle \perp.$$

Putting this together we get

$$GL_s \vdash \langle \langle \langle \alpha \rangle \varphi \rangle^{c\sharp} \rangle \perp \rightarrow \langle \zeta_c \rangle \langle \widehat{\alpha} \rangle \widehat{\varphi}.$$

as required.

- (4) If the game is atomic of the form a there are three cases. If $c(a) = \emptyset$ we need to show

$$GL_s \vdash \langle a; \zeta_c; \beta_c \rangle \perp \rightarrow \langle \zeta_c \rangle \langle \widehat{a}; \widehat{\beta} \rangle \perp$$

This is easy to see using P, (3) and $\sim$.

If $c(a) = \diamond$ note that the required

$$GL_s \vdash \langle \zeta_c; \beta_c \rangle \perp \rightarrow \langle \zeta_c \rangle \langle \widehat{a}; \widehat{\beta} \rangle \perp$$

is easily provable with $\sim$ since $\sim a_\diamond$ appears in ζ_c and using observation (3).

If $c(a) = \square$ then $\langle \langle a^c \rangle \perp \rangle^\sharp$ is provably false, so the implication holds vacuously.

- (5) If the game is of the form a^d there are again three cases. The case for $c(a) = \emptyset$ is analogous to the previous case. If $c(a) = \diamond$ the antecedent is valid, so we need to show $GL_s \vdash \langle \zeta_c \rangle \langle \widehat{a^d} \rangle \langle \widehat{\beta} \rangle \perp$. This follows with $\sim$ since ζ_c contains $\sim a_\diamond$.

If $c(a) = \square$ the proof is analogous to the case of $c(a) = \diamond$ for games a .

- (6) If the game is of the form $\sim a$ the claim is

$$GL_s \vdash \langle \zeta_c \frac{\diamond}{a}; \beta_c \frac{\diamond}{a} \rangle \perp \rightarrow \langle \zeta_c \rangle \langle \widehat{\sim a}; \widehat{\beta} \rangle \perp.$$

This is easy to see by rearranging the sabotages with P and observation (3).

The case for $\sim a^d$ is analogous.

- (7) If the game is of the form $? \varphi, ! \varphi, \beta \cup \gamma$ or $\beta \cap \gamma$ the claim is easy to derive from the induction hypothesis using $\forall$ and $\wedge$.
- (8) If the game is of the form $\alpha; \gamma$ note that by induction hypothesis on α :

$$GL_s \vdash \langle \langle \alpha; \gamma \rangle^{c\sharp} \frac{\zeta_c; \beta_c}{y_c} \rangle \perp \rightarrow \langle \zeta_c \rangle \langle \widehat{\alpha}; \bigcup_{e \in C_U} (\xi_e; \zeta_e; \gamma^{e\sharp} \frac{\zeta_e; \beta_e}{y_e}) \rangle \perp$$

Applying the induction hypothesis on γ this shows:

$$GL_s \vdash \langle \langle \alpha; \gamma \rangle^{c\sharp} \frac{\zeta_c; \beta_c}{y_c} \rangle \perp \rightarrow \langle \zeta_c \rangle \langle \widehat{\alpha}; \bigcup_{e \in C_U} (\xi_e; \zeta_e; \zeta_e; \widehat{\gamma}; \widehat{\beta}_e) \rangle \perp$$

Now $\approx$ and (3) yield the required implication.

- (9) If the game is of the form α^* list the contexts $c_1, \dots, c_n$ as in the definition of α^{*c} .

We need to show:

$$GL_s \vdash \langle \langle r_i(z_{c_1}, \dots, z_{c_n}). (y_{c_1} \cup \alpha^{c_1} \frac{z_{c_1}}{y_{c_1}}, \dots, y_{c_n} \cup \alpha^{c_n} \frac{z_{c_n}}{y_{c_n}}) \rangle^\sharp \frac{\zeta_c; \beta_c}{y_c} \rangle \perp \rightarrow \langle \zeta_{c_i} \rangle \langle \widehat{\alpha^*}; \widehat{\beta}_c \rangle \perp$$

By the inductive generalization (3) of Lemma 7.6 to the vectorial fixpoints from Theorem 7.1 this reduces to proving (for all $i = 1, \dots, n$):

$$GL_s \vdash \langle \zeta_{c_i}; \beta_{c_i} \cup (\alpha^{c_i})^\sharp \frac{\zeta_c; \widehat{\alpha^*}; \widehat{\beta}_c; ?\perp}{y_c} \rangle \perp \rightarrow \langle \zeta_{c_i} \rangle \langle \widehat{\alpha^*}; \widehat{\beta}_c \rangle \perp$$

By the induction hypothesis and (3) on α this reduces to

$$\text{GL}_s \vdash \langle \zeta_{c_i}; \beta_{c_i} \cup \zeta_{c_i}; \widehat{\alpha}; \check{\alpha}^*; \check{\beta}_c; ? \perp \rangle \perp \rightarrow \langle \zeta_{c_i}; (\widehat{(\alpha^*)}; \check{\beta}_c) \rangle \perp$$

Using (3) this again reduces to

$$\text{GL}_s \vdash \langle \zeta_{c_i}; (\check{\beta}_c \cup \check{\alpha}^*; \check{\beta}_c; ? \perp) \rangle \perp \rightarrow \langle \check{\alpha}^*; \check{\beta}_c \rangle \perp$$

This now easily follows from an instance of $*_G$ with M_G .

- (10) If the game is of the form $\alpha^\times$ list the contexts $c_1, \dots, c_n$ as in the definition of $\alpha^\times$. Let

$$\gamma_{c_i} \equiv (\iota_i(z_{c_1}, \dots, z_{c_n}) \cdot (y_{c_1} \cap \alpha^{c_1} \frac{w_{c_1}; z_{c_1}}{y_{c_1}}, \dots, y_{c_n} \cap \alpha^{c_n} \frac{w_{c_n}; z_{c_n}}{y_{c_n}})) \frac{\zeta_{c_i}; \beta_{c_i}}{y_{c_i}}$$

And let $\delta_i = \gamma_{c_i} \frac{\zeta_{c_i}}{w_{c_i}}$. By (2) of Lemma 7.6 and $\neg$

$$\text{GL}_s \vdash \langle \eta_{c_i} \rangle \perp \rightarrow \langle \zeta_{c_i}; \beta_{c_i} \cap (\alpha^{c_i}) \frac{\zeta_{c_i}; \delta_{c_i}}{y_{c_i}} \rangle \perp. \quad (3)$$

(Although it is written as a vectorial fixpoint, formally it is a single variable fixpoint as defined by Theorem 7.1 and Lemma 7.6 applies.) By applying M_G

$$\text{GL}_s \vdash \langle \zeta_{c_i}; \delta_{c_i} \rangle \perp \rightarrow \langle \zeta_{c_i}; \beta_{c_i} \cap \zeta_{c_i}; (\alpha^{c_i}) \frac{\zeta_{c_i}; \delta_{c_i}}{y_{c_i}} \rangle \perp.$$

By the induction hypothesis on α and $\approx$ we obtain

$$\text{GL}_s \vdash \langle \zeta_{c_i}; \delta_{c_i} \rangle \perp \rightarrow \langle \zeta_{c_i}; \beta_{c_i} \cap \zeta_{c_i}; \widehat{\alpha}; \check{\delta}_c \rangle \perp.$$

Combining these for all i and using M_G

$$\text{GL}_s \vdash \langle \check{\delta} \rangle \perp \rightarrow \langle \check{\beta}_c \cap \bigcup_{e \in C_U} (\zeta_e; \zeta_e; \widehat{\alpha}); \check{\delta}_c \rangle \perp.$$

By $*_G$ this implies

$$\text{GL}_s \vdash \langle \check{\delta} \rangle \perp \rightarrow \langle \widehat{(\alpha^*)}; \check{\beta}_c \rangle \perp.$$

Using M_G and observation (3)

$$\text{GL}_s \vdash \langle \zeta_{c_i}; \delta_{c_i} \rangle \perp \rightarrow \langle \zeta_{c_i}; \widehat{(\alpha^*)}; \check{\beta}_c \rangle \perp.$$

By $\simeq$ the game $\zeta_{c_i}; \delta_{c_i}$ is equivalent to $\alpha^{c_i} \frac{\zeta_{c_i}; \beta_{c_i}}{y_{c_i}}$. $\square$

PROOF OF PROPOSITION 7.8. Let π be a proof of φ in rlGL . By replacing every free variable in any formula in π by a fixed atomic game a we turn π into a rlGL proof of φ which does not mention formulas which have free variables. Next replace every formula φ in this proof by $\varphi^{\natural}$ and call the resulting proof π' . This proof can now be transformed into a GL_s proof for $\text{GL}_s \vdash \varphi^{\natural}$.

By definition and (1) of Lemma 7.6 it is clear that the translation of most proof rules and axioms remain proof rules and axioms respectively. For the fp_G axiom we use (2) of Lemma 7.6 and include a derivation of this axiom. For the rule μ_G we can extend the proof appropriately by (3) of Lemma 7.6. $\square$

PROOF OF THEOREM 7.9. Soundness holds by Lemma 6.5. If φ is a valid formula of GL_s , the translation $\langle \varphi^{c^0} \rangle \perp$ is valid by Proposition 7.2. Moreover by Theorem 6.8 and Proposition 7.8 the formula $\text{GL}_s \vdash (\langle \varphi^{c^0} \rangle \perp)^{\natural}$ is provable and therefore by Lemma 7.7 also $\text{GL}_s \vdash \varphi$. $\square$