

Model Checking Nondeterministic and Randomly Timed Systems

Martin R. Neuhäuser

Graduation committee:

Prof. Dr. Ir. A. J. Mouthaan (chairman)	University of Twente, The Netherlands
Prof. Dr. Ir. Joost-Pieter Katoen (promotor)	RWTH Aachen / University of Twente, Germany / The Netherlands
Dr. Mariëlle I. A. Stoelinga (referent)	University of Twente, The Netherlands
Prof. Dr. Jos C. M. Baeten	Eindhoven University of Technology, The Netherlands
Prof. Dr. Ir. Boudewijn R. Haverkort	University of Twente, The Netherlands
Prof. Dr.-Ing. Holger Hermanns	Saarland University, Germany
Prof. Dr. Jaco C. van de Pol	University of Twente, The Netherlands
Prof. Dr. Roberto Segala	University of Verona, Italy



IPA Dissertation Series 2010-02.
CTIT Ph.D.-Thesis Series No. 09-165, ISSN 1381-3617.
ISBN: 978-90-365-2975-4.

The research reported in this dissertation has been carried out under the auspices of the Institute for Programming Research and Algorithmics (IPA) and within the context of the Center for Telematics and Information Technology (CTIT). The research funding was provided by the NWO Grant through the project: Verifying Quantitative Properties of Embedded Software (QUPES).

Translation of the abstract: Viet Yen Nguyen (MSc).
Typeset in L^AT_EX.
Cover design: Anja Balsfulland
Publisher: Wöhrmann Print Service - <http://www.wps.nl>.

Copyright © 2010 by Martin R. Neuhäuser, Aachen, Germany.

MODEL CHECKING NONDETERMINISTIC AND RANDOMLY TIMED SYSTEMS

Dissertation

to obtain the doctor's degree
at the University of Twente, on the authority of
the rector magnificus, Prof. Dr. H. Brinksma,
on account of the decision of the graduation committee
to be publicly defended
on Friday, January 22, 2010 at 13:15

by

Martin Richard Neuhäuser
born on 01 September 1979
in Kulmbach, Germany

The dissertation has been approved by the promotor:

Prof. Dr. Ir. Joost-Pieter Katoen

Model Checking Nondeterministic and Randomly Timed Systems

Von der Fakultät für Mathematik, Informatik und
Naturwissenschaften der Rheinisch-Westfälischen Technischen
Hochschule Aachen zur Erlangung des akademischen Grades
eines Doktors der Naturwissenschaften genehmigte Dissertation

vorgelegt von

Diplom-Informatiker
Martin Richard Neuhäuser

aus
Kulmbach

Berichter: Prof. Dr. Ir. Joost-Pieter Katoen
Prof. Dr. Franck van Breugel

Tag der mündlichen Prüfung: 25. Januar 2010

Abstract

Formal methods initially focused on the mathematically precise specification, design and analysis of functional aspects of software and hardware systems. In this context, model checking has proved to be tremendously successful in analyzing qualitative properties of distributed systems. This observation has encouraged people in the field of performance and dependability evaluation to extend existing model checking techniques to also account for quantitative measures. As a result, nowadays, the automatic analysis of Markovian models has become an indispensable tool for the design and evaluation of safety and performance critical systems.

Markovian models are classified according to their underlying notion of time, being either discrete or continuous. In the discrete-time setting, Markov decision processes are a nondeterministic model which is widely known in mathematics, computer science and operations research. Moreover, efficient algorithms are available for their analysis. This stands in sharp contrast to the continuous-time setting, where no techniques exist to analyze models that combine stochastic timing and nondeterminism. In the present thesis, we bridge this gap and propose quantifiably precise model checking algorithms for a variety of nondeterministic and stochastic models.

We first consider continuous-time Markov decision processes (CTMDPs). To uniquely determine the quantitative properties of a CTMDP, all its nondeterministic choices must be resolved according to some strategy. Therefore, we propose a hierarchy of scheduler classes and investigate their impact on the achievable performance and dependability measures. In this context, we identify late schedulers, which resolve the nondeterminism as neatly as possible. Apart from their interesting theoretical properties, they facilitate the analysis of locally uniform CTMDPs considerably. In a locally uniform CTMDP, the timing in a state is independent of the scheduler. This observation culminates in an efficient and quantifiably precise *approximation algorithm for locally uniform CTMDPs*.

In contrast to CTMDPs which closely entangle nondeterminism and stochastic time, interactive Markov chains (IMCs) are a highly versatile model that strictly uncouples the two aspects. Due to this separation of concerns, IMCs are locally uniform by definition. This allows us to apply analysis techniques which are similar to those that we developed for locally uniform CTMDPs, also to IMCs. In this way, we solve the open problem of *model checking arbitrary IMCs*.

In the next step, we return to CTMDPs and prove that they can be transformed into alternating IMCs in a measure preserving way. As our proof does not rely on local uniformity, it enables the analysis of quantitative measures on arbitrary CTMDPs by model checking their induced IMCs. However, the underlying scheduler class slightly differs

from the late schedulers that we used initially. In fact, it coincides with the time- and history dependent schedulers that are proposed in the literature. Thus, our result for IMCs also solves the long standing problem of *model checking arbitrary CTMDPs*.

However, the applicability of model checking is limited by the infamous state space explosion problem: Even systems of moderate size often yield models with an exponentially larger state space that foils their analysis. To tackle this problem, many techniques have been developed that minimize the state space while preserving important properties of the model. In process algebras, *bisimulation minimization* identifies processes with the same quantitative behavior and replaces equivalent ones by a single representative. Depending on the redundancy in the model, this can lead to enormous reductions in the size of the state space. As IMCs have a process algebraic background, it is not surprising that bisimulation minimization is readily available for them. However, this is not the case for CTMDPs. That is why we introduce bisimulation minimization for CTMDPs and prove that it preserves all quantitative measures.

Finally, we apply the achieved results and propose an alternative semantics for *generalized stochastic Petri nets* (GSPN), which avoids the shortcomings of earlier definitions that were needed to rule out nondeterministic choices. More precisely, we transform a GSPN model into an equivalent IMC which can be model checked.

To show the applicability of our approach, we analyze *the dependability of a workstation cluster* which is modeled by a nondeterministic GSPN. The comparison of our results with those that are available in the literature is illuminating: When the latter were published, no analysis technique for nondeterministic and randomly timed systems was available. Therefore, the nondeterministic choices in the GSPN model were replaced by static probability distributions.

For measures that are mostly independent of the scheduling policy, our results coincide with those in the literature. However, for other measures, choosing antagonistic schedulers mitigates the inferred dependability characteristic of the system that we study by up to 18%. These false positives in the earlier analyses clearly prove the necessity of nondeterministic modeling in the field of performance and dependability analysis.

Samenvatting

Formele methoden worden van oudsher toegepast met een wiskundig rigoureuze benadering van specificatie, ontwerp en analyse van functionele aspecten in hard- en software. Met name model checking bleek enorm succesvol te zijn om kwalitatieve eigenschappen van gedistribueerde systemen te analyseren. Dit moedigde onderzoekers in performance evaluatie en betrouwbaarheidsanalyse aan om diezelfde technieken te benutten voor kwantitatieve analyses. Als gevolg daarvan is de automatische analyse van Markov modellen een onmisbaar middel geworden voor het ontwerp en evaluatie van betrouwbare systemen.

Markov modellen worden doorgaans geclassificeerd aan de hand van hun onderliggende interpretatie van tijd, hetzij discreet of continu. Betreffende het eerstgenoemde, zijn Markov decision processes wijdverspreid in de wiskunde, informatica en operationele research. Er zijn efficiënte algoritmen beschikbaar om deze modellen te analyseren. Dit staat in scherp contrast met haar continue-tijdstegenhanger. Er waren tot heden nog geen technieken ontwikkeld voor modellen met stochastische timing en non-determinisme. In dit proefschrift overbruggen we deze tekortkoming met onze behandeling van kwantitatief precieze model checking algoritmes voor een scala van non-deterministische en stochastische modellen.

We behandelen eerst Continuous-Time Markov Decision Processes (CTMDPs). Om de kwantitatieve eigenschappen van een non-deterministisch model te bepalen moeten alle non-deterministische keuzes vastgelegd worden volgens een strategie. Om die reden presenteren wij een hiërarchie van scheduler klassen en onderzoeken wij hun impact op performance en betrouwbaarheidsmaten. In deze context identificeren we de klasse van "late schedulers". Naast hun interessante theoretische eigenschappen, faciliteren zij de analyse van lokaal uniform CTMDPs. Voor deze schedulers en modellen presenteren we namelijk een precies benaderingsalgoritme.

In tegenstelling tot CTMDPs, waarbij non-determinisme en stochastische tijd sterk verstrengeld zijn, zijn Interactive Markov Chains (IMCs) een extreem veelzijdig formalisme waarin deze twee aspecten zijn ontkoppeld. Door deze ontkoppeling zijn IMCs per definitie lokaal uniform. De technieken die we hebben ontwikkeld voor lokaal uniform CTMDPs zijn conceptueel vergelijkbaar met die voor IMCs. Op deze wijze hebben we het openstaande model checking probleem van IMCs opgelost.

Vervolgens laten we zien hoe CTMDPs afbeeldbaar zijn op alternerende IMCs waarbij de maten behouden blijven. Ons bewijs van dit resultaat vereist niet dat de CTMDP lokaal uniform is. Dit maakt kwantitatieve analyses mogelijk voor algemene CTMDPs door hun geïnduceerde IMCs te analyseren. De scheduler klasse die hierbij nodig is wijkt

enigszins af van die we gebruikten om lokaal uniform CTMDPs te analyseren. Sterker nog, die afwijkende klasse valt samen met de tijds- en historie afhankelijke schedulers die bekend zijn in de literatuur. De resultaten lossen derhalve een langdurig openstaand probleem op, namelijk het model checken van arbitraire CTMDPs.

De toepassing van model checking is echter gelimiteerd door de fameuze explosie van de toestandsruimte. Zelfs systemen van gemiddelde complexiteit leiden vaak tot een exponentieel groeiende toestandsruimte wat het model checken bemoeilijkt. Om dit probleem aan te pakken zijn er vele technieken ontwikkeld die de toestandsruimte minimaliseren terwijl haar eigenschappen intact blijven. In proces algebra's identificeert bisimulatie minimalisatie de processen die eenzelfde kwantitatief gedrag vertonen en vervangt deze door een enkel representatief gedrag. Afhankelijk van de redundantie in het model kan de toestandsruimte aanzienlijk reduceren. Aangezien IMCs als basis dienen voor stochastische proces algebra's is het niet verwonderlijk dat er reeds bisimulatie minimalisatie technieken voor IMCs bestaan. Dit is echter niet het geval voor CTMDPs. Daarom onderzochten wij tevens bisimulatie minimalisatie voor CTMDPs en bewijzen dat die alle kwantitatieve maten intact houdt.

Ten slotte passen we onze resultaten toe en presenteren we een alternatieve semantiek voor generalized stochastic Petri nets (GSPNs). Deze vermijdt de tekortkomingen van voorgaande definities in de literatuur die nodig waren om non-deterministische keuzes te omzeilen. Hiertoe beelden we een GSPN model af op haar equivalente IMC model die vervolgens met onze technieken gemodelcheckt kan worden.

Ter demonstratie van onze aanpak, analyseren wij de betrouwbaarheid van een workstation cluster die gemodelleerd is als een niet-deterministische GSPN. Een vergelijking van onze resultaten met die uit de literatuur levert enkele interessante bevindingen op. Hier dient vermeld te worden dat de eerder gepubliceerde resultaten verkregen zijn door niet-deterministische keuzemomenten door uniforme kansverdelingen te vervangen.

Voor maten die grotendeels onafhankelijk zijn van de scheduling tactiek, komen onze resultaten overeen met de bestaande. Echter, voor andere maten leidt de keuze van antagonistische schedulers tot een verslechtering van de verkregen betrouwbaarheidskarakteristieken met maar liefst 18%. Deze uitkomsten tonen de noodzaak van het meenemen van niet-deterministische keuzes in de prestatie- en betrouwbaarheidsanalyse onomstootelijk aan.

Zusammenfassung

In der Informatik beschäftigt sich das Gebiet der formalen Methoden ursprünglich mit der Spezifikation, dem Design und der Analyse funktionaler Aspekte von Hard- und Software. Vor diesem Hintergrund hat sich Model Checking als äußerst nützlich beim Analysieren quantitativer Eigenschaften verteilter Systeme erwiesen. Daraufhin wurde im Bereich der Leistungs- und Verlässlichkeitsbewertung begonnen, die existierenden Model Checking Verfahren auf quantitative Eigenschaften zu erweitern. Heute ist die Analyse der entsprechenden Markovmodelle ein unabdingbarer Bestandteil beim Design und der Evaluierung der Sicherheit und Leistung kritischer Systeme.

Es werden entsprechend dem zugrunde liegenden Zeitbegriff diskrete und kontinuierliche Markovmodelle unterschieden. Im zeitdiskreten Fall sind Markov-Entscheidungsprozesse (MDPs) ein weit verbreitetes nichtdeterministisches Modell in der Mathematik und der Informatik. Für die Analyse von MDPs stehen effiziente Algorithmen zur Verfügung. Dagegen sind für den zeitkontinuierlichen Fall bisher keine Methoden für die automatische Analyse von Modellen bekannt, die stochastisch quantifiziertes Zeitverhalten und Nichtdeterminismus verbinden. Die vorliegende Dissertation schließt diese Lücke und führt präzise und quantifizierbar korrekte Model Checking Algorithmen für eine Vielzahl von nichtdeterministischen und stochastischen Modellen ein.

Anfangs betrachten wir sogenannte zeitkontinuierliche Markov-Entscheidungsprozesse (CTMDPs). Um die quantitativen Eigenschaften einer CTMDP eindeutig zu bestimmen, müssen zunächst alle in ihr vorkommenden nichtdeterministischen Wahlmöglichkeiten anhand einer Strategie aufgelöst werden. Dazu führen wir eine Hierarchie von Schedulerklassen ein und untersuchen ihren Einfluss auf die erzielbaren Leistungs- und Verlässlichkeitsanforderungen. In diesem Zusammenhang beschreiben wir sogenannte verzögerte Scheduler, die den Nichtdeterminismus bestmöglich auflösen. Neben ihren interessanten theoretischen Eigenschaften erleichtern sie die Analyse von lokal uniformen CTMDPs erheblich. Dabei bilden lokal uniforme CTMDPs eine Teilklasse, in der das Zeitverhalten der Zustände unabhängig vom Scheduler ist. Diese Beobachtung ist Grundlage für einen effizienten und quantifizierbar korrekten Approximationsalgorithmus für lokal uniforme CTMDPs.

Im Gegensatz zu CTMDPs, die Nichtdeterminismen und stochastisches Zeitverhalten eng miteinander verbinden, sind interaktive Markovketten (IMCs) ein Modell, das diese beiden Aspekte strikt trennt. Aus diesem Grund sind IMCs per Definition bereits lokal uniform. Das ermöglicht es, Analysetechniken, die denen für lokal uniforme CTMDPs ähneln, auch auf IMCs anzuwenden. Auf diese Weise lösen wir die offene Frage nach einem Model Checking Algorithmus für IMCs.

Im nächsten Schritt kehren wir zu CTMDPs zurück und beweisen, dass sie auf maßerhaltende Art und Weise in alternierende IMCs transformiert werden können. Da unser Beweis nicht auf lokale Uniformität angewiesen ist, ermöglicht er die Analyse quantitativer Eigenschaften von allgemeinen CTMDPs anhand ihrer induzierten IMCs. Jedoch unterscheiden sich die zugrunde liegenden Schedulerklassen leicht von den bisher betrachteten verzögerten Schemulern. Tatsächlich stimmen sie mit den zeit- und verlaufsabhängigen Schemulern, die in der Literatur bekannt sind, überein. Damit lösen unsere Resultate auch das seit langem offene Problem der Analyse allgemeiner CTMDPs.

Im Allgemeinen wird die Anwendbarkeit von Model Checking durch das exponentielle Anwachsen der Zustandsräume begrenzt. Viele Techniken sind entwickelt worden, um den Zustandsraum unter Beibehaltung wichtiger Eigenschaften zu minimieren. Im Bereich der Prozessalgebren fasst Bisimulation Zustände zusammen, die die gleichen Eigenschaften haben. Abhängig von der im Modell enthaltenen Redundanz führt das oft zu einer erheblichen Reduktion des Zustandsraums. Da IMCs aus Prozessalgebren hervorgehen, ist es nicht verwunderlich, dass Bisimulationsminimierung für sie bereits untersucht wurde. Das trifft jedoch nicht auf CTMDPs zu. Daher führen wir Bisimulation auf CTMDPs ein und weisen nach, dass durch sie alle quantitativen Maße erhalten bleiben.

Abschließend wenden wir die erzielten Resultate an und entwickeln eine alternative Semantik für GSPNs, die die Nachteile früherer Ansätze hinsichtlich der Berücksichtigung von Nichtdeterminismen umgeht. Dazu transformieren wir GSPN Modelle in äquivalente IMCs, die anschließend analysiert werden.

Um die Anwendbarkeit unseres Ansatzes zu zeigen, analysieren wir so die Verlässlichkeit eines Workstation-Clusters, der als nichtdeterministisches GSPN modelliert wird. Interessant ist dabei besonders der Vergleich unserer Ergebnisse mit früher veröffentlichten Resultaten. Letztere wurden publiziert, als noch keine Analysetechniken für nichtdeterministische Systeme mit stochastischem Zeitverhalten verfügbar waren. Daher wurden die im GSPN-Modell auftretenden Nichtdeterminismen auf festgelegte Art und Weise durch Wahrscheinlichkeitsverteilungen ersetzt.

Für Maße, die kaum von den Wahlmöglichkeiten des Schemulers abhängen, stimmen unsere Resultate mit denen aus der Literatur überein. Für andere Maße jedoch liegen die ableitbaren Verlässlichkeitscharakteristika des Systems für antagonistische Schemuler um bis zu 18% unter den Vorhersagen früherer Modelle. Diese falsch positiven früheren Analysen verdeutlichen die Notwendigkeit nichtdeterministischer Modellierung im Bereich der Leistungs- und Verlässlichkeitsbewertung.

Acknowledgments

Writing a dissertation has been a big challenge for me. I would not have completed the present work without the many people I met during the last four years.

First of all, I thank my promotor Joost-Pieter Katoen for all his support and encouragement. With his guidance, the many fruitful discussion that we had and with his patience, he laid the solid base that I relied on during all my research.

Most of the results presented in this thesis are a product of joint work with my colleagues. Without David Jansen's mathematical rigor and his patience, I would never have been able to appreciate measure theory. Further, I thank Mariëlle Stoelinga and Lijun Zhang for our pleasant and fruitful cooperation. It is great fun to write papers with you!

During the last four years, the colleagues at Joost-Pieter Katoen's MOVES group in Aachen became close friends. I will always remember our skiing vacations, the daily chats in Stefan's and Carsten's office and the summer schools and conference dinners that we attended. Without Alexandru, Arnd, Carsten, Daniel, Elke, Haidi, Henrik, Jonathan, Stefan, Thomas, Tingting and Viet Yen, my PhD life would not have been half that enjoyable!

Last but not least, I would like to thank Alena and my parents for their unconditional love, support and advice. Without their encouragement and patience, I would not have reached that far.

Contents

1	Introduction	3
1.1	System validation	3
1.2	The quantitative analysis of stochastic models	5
1.3	The contribution of the thesis	7
1.4	Outline of the thesis	8
1.5	Origins of the chapters and credits	9
2	Basics of measure & probability theory	11
2.1	Basics of measure theory	12
2.2	The Borel σ -field and the Lebesgue measure	24
2.3	A set that is not Lebesgue measurable	30
2.4	The Lebesgue integral	33
2.5	Product σ -fields	41
2.6	Concluding remarks	52
3	An overview of stochastic models	55
3.1	Stochastic processes	55
3.2	Markov chains	56
3.3	Nondeterminism in stochastic models	69
3.4	Conclusion	84
4	Schedulers in CTMDPs	85
4.1	A hierarchy of scheduler classes	86
4.2	Local uniformization	90
4.3	Preservation results for local uniformization	103
4.4	Delaying nondeterministic choices	108
4.5	Conclusion	111
5	The analysis of late CTMDPs	113
5.1	Locally uniform CTMDPs	114
5.2	A fixed point characterization for time-bounded reachability	118
5.3	Computing time-bounded reachability probabilities	130
5.4	A case study: The stochastic job scheduling problem	141
5.5	Conclusion and related work	143

6	Model Checking Interactive Markov Chains	145
6.1	Interactive Markov chains	147
6.2	Interval bounded reachability probability	154
6.3	A discretization that reduces IMCs to IPCs	162
6.4	Solving the problem on the reduced IPC	184
6.5	Model checking the continuous stochastic logic	189
6.6	Experimental results	194
6.7	Interval bounded reachability in early CTMDPs	194
6.8	Comparison of different scheduler classes	200
6.9	Related work and conclusions	200
7	Equivalences and logics for CTMDPs	203
7.1	Strong bisimilarity	204
7.2	Continuous Stochastic Logic	209
7.3	Strong bisimilarity preserves CSL	212
7.4	Conclusion	217
8	Model checking generalized stochastic Petri nets	219
8.1	Preliminaries	221
8.2	The syntax of GSPNs	221
8.3	A new semantics for GSPNs	223
8.4	Dependability analysis of a workstation cluster	226
8.5	Conclusion	232
9	Conclusion	233
	Bibliography	235

Summary of Notation

We indicate here the basic notational conventions that are used throughout the thesis. We use $\square$ and $\diamond$ to denote the end of proofs and examples, respectively.

Numbers

We use $\mathbb{R}_{\geq 0}$, $\mathbb{R}_{> 0}$ and $\mathbb{R}$ to denote the sets of nonnegative, positive and the set of all real numbers; similarly, the sets $\mathbb{Q}_{\geq 0}$, $\mathbb{Q}_{> 0}$ and $\mathbb{Q}$ refer to the nonnegative, positive and all rational numbers. Moreover, $\mathbb{N} = \{0, 1, 2, \dots\}$ denotes the set of natural numbers. If $T \subseteq \mathbb{R}_{\geq 0}$ and $t \in \mathbb{R}_{\geq 0}$, we define

$$T \oplus t = \{x + t \mid x \in T\}, \text{ and} \\ T \ominus t = \{x - t \mid x \in T, x \geq t\}.$$

Sets

Let $\mathcal{Z}$ be a set with subsets A and B . If $A \cap B = \emptyset$, we use $A \cup B$ to denote the disjoint union of the sets A and B . The indicator for a subset A of $\mathcal{Z}$ is defined as the function

$$\mathbf{I}_A : \mathcal{Z} \rightarrow \{0, 1\} : x \mapsto \begin{cases} 1 & \text{if } x \in A \\ 0 & \text{otherwise.} \end{cases}$$

If $A_1 \subseteq A_2 \subseteq \dots$ is an increasing sequence of subsets of $\mathcal{Z}$ and $\lim_{n \rightarrow \infty} A_n = A$, we write $A_n \uparrow A$. Similarly, $A_n \downarrow A$ denotes a decreasing sequence with limit set A .

Functions

If $f : \mathcal{Z}_1 \times \mathcal{Z}_2 \times \dots \times \mathcal{Z}_n \rightarrow \mathcal{Z}$ is an n -ary function, we use $f(z_1, z_2, \dots, z_{i-1}, \cdot, z_{i+1}, \dots, z_{n-1}, z_n)$ and, depending on the context, also $f(z_1, z_2, \dots, z_{i-1}, [\cdot], z_{i+1}, \dots, z_{n-1}, z_n)$ to denote the function $z_i \mapsto f(z_1, z_2, \dots, z_{i-1}, z_i, z_{i+1}, \dots, z_{n-1}, z_n)$.

Probability distributions

Let $\mathcal{X} = \{x_0, x_1, x_2, \dots, x_n\}$ be a finite set. Probability distributions on $\mathcal{X}$ are functions $\mu : \mathcal{X} \rightarrow [0, 1]$ with $\sum_{x \in \mathcal{X}} \mu(x) = 1$. Moreover, we write $\mu = \{x_0 \mapsto p_0, x_1 \mapsto p_1, \dots, x_n \mapsto p_n\}$ to denote the probability distribution μ where $\mu(x_i) = p_i$. If $\mu(x) = 1$ for some $x \in \mathcal{X}$, we write $\mu = \{x \mapsto 1\}$ and identify μ and x . The set of all probability distributions over $\mathcal{X}$ is denoted $\text{Distr}(\mathcal{X})$. If $\mu \in \text{Distr}(\mathcal{X})$ and $A \subseteq \mathcal{X}$, then $\mu(A) = \sum_{x \in A} \mu(x)$.

1 Introduction

It is fair to state, that in this digital era correct systems for information processing are more valuable than gold.

(Henk Barendregt)

When you woke up today, the first thing that you perceived was probably the microcontroller-driven bell of your alarm clock. On the way to your office, you rely on the software that schedules your metro train while optimizing the metro system's signal headway. At work, you expect the operating system of your workstation to store and manipulate your data correctly. And if you happen to be involved in an accident on your way back home, you depend on an operational mobile phone network to call an ambulance that takes you to the hospital. But even there, you are confronted with software and hardware systems that monitor your pulse, provide oxygen to your lungs or compute the X-Ray dose necessary for radiation therapy.

Today, the ubiquitous use of embedded systems in our daily lives makes us highly dependent on their correctness. The consequences of failures range from just getting up too late to social and economic disasters. However, accompanied by the unmatched advancements that have been achieved in the design of integrated circuits since the late 1960's, the realizable software and hardware systems have become evermore complex. Today, this growing complexity leads to serious errors in safety critical systems [Baa08] as witnessed by prominent examples, such as the erroneous flight control unit which destroyed the Ariane-5 rocket, or the Therac-25 radiation therapy machine which killed at least three patients due to a race condition in its control software, which led to a lethal overdose of X-Rays. Hence, it is fair to state that methodologies which assure the correctness of safety critical systems are of vital importance.

1.1 System validation

In computer science, the field of *formal methods* focuses on techniques for the mathematically precise design, modeling and verification of functional aspects of safety critical systems. Accordingly, the aim of system validation is to guarantee that the physical system fulfills its intended purpose.

In this context, *model checking* refers to the automatic verification of a system model

against a specification that is usually given as a logic formula. As depicted in Fig. 1.1, the model checking approach relies on at least three ingredients: the model, the property specification and the verification algorithm that checks the validity of the property in the model. We discuss each of them shortly.

Model checking can only guarantee that a mathematical model of the actual system — where the model is usually given by a Kripke structure — conforms to the specification. Obviously, all results are void if the model does not accurately reflect the behavior of the system. Thus, a fundamental requirement for formal validation is to derive a mathematically precise model so that the verification results that are obtained on the model carry over to its actual implementation.

If software engineers used a formal modeling language during the design phase, the system model could be inferred automatically. However, in today's practice, mostly semi-formal approaches like the UML [BR04] or even informal natural language specifications are used. This lack of mathematical rigor leads to ambiguities in the design and impedes a formal validation of the system. Therefore, most people in the formal methods community favor the use of completely formal specification languages like Statecharts [Har87, Jan03], queueing networks [CG89], Petri nets [Rei85] or process algebras [Mil82, Hoa85, BW90, Mil99]. In this way, the system specification automatically translates into a precise *system model*, which allows us to formally validate the system.

Having a formal model at hand, the next step is to identify the properties that need to be checked. Usually, logics like LTL [Pnu77] and CTL [CES86] are used for the *property specification*. They permit to express functional aspects of the model such as “Two trains never collide in the metro system” or “The routing algorithm stabilizes eventually after a router has failed”.

Finally, given the model T of the system and a formula Φ which specifies the desired property, a model checking tool like Spin [Hol04] or NuSMV [CCGR00] automatically *verifies* whether the model satisfies the property. A positive outcome allows us to conclude that the system satisfies the corresponding property. Moreover, if the result is negative, model checking offers diagnostic feedback by identifying the faulty behaviors.

In this way, classical model checking verifies *qualitative* system properties by providing a definite yes-or-no answer. However, it is often impossible to completely prove the correctness of realistic systems, as they are embedded in an environment and therefore subject to random phenomena. For example, a detailed model of a distributed system should reflect the probability that messages get lost or become garbled during transmission. Although this closely reflects the physical behavior of the system, it is hard to guarantee its correctness by providing a definite yes-or-no answer. Therefore, we strive for a less stringent notion of correctness, which enables us to quantify the degree at which the model meets its specification. For example, proving that the probability of a system failure is less than 0.1% might convince us to rely on that system despite the unlikely event that it might fail.

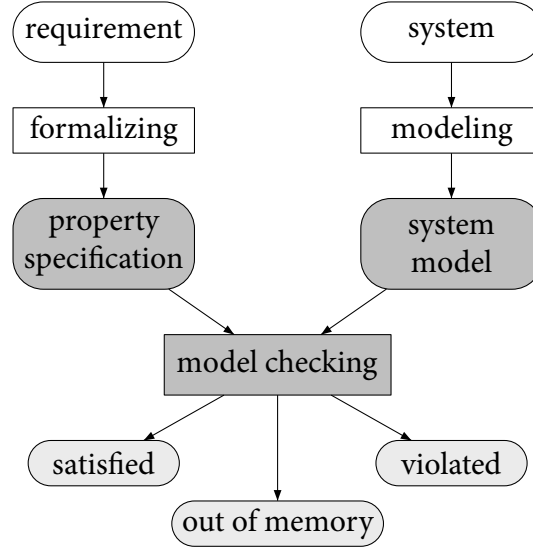


Figure 1.1: Verifying system correctness by model checking [BK08].

1.2 The quantitative analysis of stochastic models

Applying model checking to analyze quantitative properties allows us to infer a variety of performance and dependability measures automatically. Typical examples are the average throughput of a router, the expected round trip time of an IP-packet or the mean time between failures of a hard disk drive. In all these scenarios, we do not expect a rigid yes-or-no answer, but need to find quantitative measures that describe the system.

A plethora of models has been proposed that incorporate probability distributions into the classical transition system formalism; thereby, they permit to specify the quantitative behavior of the underlying system. In the context of this thesis, we classify quantitative models along two dimensions:

1. *Discrete vs. continuous.* Time can be measured either in discrete entities or continuously: In *probabilistic* models, time is represented by a sequence of discrete steps which are usually identified with the natural numbers. Hence, the transitions in a probabilistic model occur synchronously with its discrete time ticks. The randomness of the system is determined by discrete probability distributions over successor states that specify the likelihood to move from one state to another and by a probability distribution over initial states.

Unlike discrete-time models, *stochastic models* adopt a continuous notion of time. In this setting, transitions are delayed by a random amount of time which is governed by a continuous probability distribution. Hence, time points are drawn from the set of nonnegative real numbers. A continuous-time model moves from one state to another according to the transition which executes first. In this way, probabilistic and timed behaviors are closely entangled in stochastic models.

2. *Deterministic vs. nondeterministic*: The behavior of a *deterministic* model is completely specified by its (discrete or continuous) probability distributions. Note that we use the term *deterministic*, although the system behavior is only determined quantitatively.

Accordingly, we call a system *nondeterministic*, if its probabilistic or stochastic behavior is not decided completely. This situation can arise intentionally, for example, if the modeler does not have enough information to estimate the probability distribution that governs the system's behavior in a specific state and therefore decides to leave it unspecified. Apart from the deliberate use of *underspecifications*, another implicit source of nondeterminism is the scheduling freedom that occurs in randomized distributed systems, where the order of executing is only partly specified. Moreover, nondeterminism occurs naturally in open systems that communicate with other components in their environment.

We summarize the models that are used in the thesis in Table 1.1. The most fundamental ones are discrete- and continuous-time Markov chains [KS76, Kul95]. Discrete-time Markov chains (DTMC) were used as a dependability model for the first time in the seminal work of Hansson and Jonsson [HJ94]. Due to their discrete notion of time, DTMCs can be used to model randomized algorithms or hardware circuits which obey a global clock pulse.

The work in [Var85, HJ94] led to further research towards model checking of continuous-time Markov chains [Kul95, ASSB96] (CTMC), which had already been widely accepted in the area of performance evaluation [Hav98]. However, an automatic analysis technique for CTMC only became available with the corresponding model checking algorithm in [BHHK03]. Nowadays, model checking tools like PRISM [KNP02, HKNP06] and MRMC [Zap08, KZH⁺09] enable an efficient analysis of CTMC models. They have been successfully adopted for the performance evaluation of queueing systems and QoS constraints, to name a few.

However, neither DTMCs nor CTMC are appropriate to model nondeterminism. In effect, this shortcoming prevents the analysis of distributed systems, which is the traditional realm of model checking.

In the discrete-time setting, *Markov decision processes* (MDPs) [Put94] are a widely known formalism in mathematics and discrete optimization which incorporates nondeterminism into DTMCs. In computer science, several extensions of MDPs like probabilistic automata [SL95, Seg95], ACP-style process algebras [And02] and interactive probabilistic chains [CHLS09] have been considered. They all support nondeterminism and have successfully been applied to study quantitative measures of randomized distributed algorithms [Seg97, SV99].

In this thesis, we focus on the bottom right corner of Table 1.1: Whereas DTMCs have successfully been extended to MDPs to account for nondeterministic choices, the corresponding continuous-time model has received scant attention in computer science. Continuous-time Markov decision processes have been studied in mathematics [Mil68b,

	discrete-time	continuous-time
deterministic	DTMC, Def. 3.5	CTMCs, Def. 3.7
non-deterministic	MDPs, Def. 3.8 IPCs, Def. 6.5	CTMDPs, Def. 3.11 IMCs, Def. 6.1

Table 1.1: The basic stochastic models used in this thesis.

Mil68a] and are mentioned shortly in [Put94, Chapter 11]. In [BHKH05], the authors develop a first model checking algorithm that works on a narrow subclass of CTMDPs; it has received quite some attention and was extended in [Joh07] to analyze interactive Markov chains [HHK02], which are another prominent model for nondeterministic and randomly timed systems. However, these approaches are severely restricted, as they assume that all states of the system have the same timed behavior.

1.3 The contribution of the thesis

Apart from the subclass of globally uniform CTMDPs, no model checking algorithms exist for nondeterministic and randomly timed systems. The aim of this thesis is to fill this gap in the theory of formal methods.

First, we investigate a hierarchy of scheduler classes which differ in the information that they can use to resolve nondeterministic choices. We compare their impact on the achievable quantitative measures and introduce the new class of *late* schedulers, which strictly improve upon those that are known from the literature.

Further, we introduce bisimulation minimization on CTMDPs and prove that all quantitative measures are preserved in the quotient. As a consequence, we are able to minimize the state space of CTMDPs prior to their analysis.

However, the main contribution of this thesis are precise and efficient model checking algorithms for a variety of nondeterministic and randomly timed systems:

- We develop a quantifiably precise model checking algorithm for locally uniform CTMDPs and late schedulers. Compared to the earlier result [BHKH05], this enlarges the class of analyzable CTMDPs considerably, as we only require that the timing in each state is independent on the resolution of the nondeterminism in that state.
- We extend the previous result to interactive Markov chains and obtain an efficient model checking algorithm. Most notably, our extension does no longer depend on any kind of uniformity. To the best of our knowledge, this is the first time that a model checking algorithm is available for arbitrary IMCs.
- By applying our results for IMCs, we succeed in model checking *arbitrary* CTMDPs. This is achieved by transforming a given CTMDP into an equivalent IMC

which we can analyse. However, compared to our native results on locally uniform CTMDPs, we have to impose mild restrictions on the scheduler class: In fact, the CTMDP model checking algorithm that we obtain computes the optimal quantitative measures with respect to the classical definition of time- and history dependent schedulers.

- Finally, we introduce a new semantics for generalized stochastic Petri nets (GSPNs), which overcomes the shortcomings in the support of nondeterminism in the previous definitions. More precisely, we transform a nondeterministic GSPN into an IMC which is subject to our analysis. In a case study, we compare the new GSPN semantics to the previous one and show the necessity of nondeterministic modeling.

All algorithms are implemented in a prototypical model checker which has been used to obtain the quantitative measures that can be found throughout the thesis.

1.4 Outline of the thesis

- In **Chapter 2**, we summarize the definitions and measure theoretic results that are necessary for a deeper understanding of the forthcoming chapters. In fact, Chapter 2 is a computer scientist's summary of the excellent, but mathematically dense textbook [ADD00].
- In **Chapter 3**, we formally introduce the probabilistic and stochastic models that form the basis of this thesis. Further, we introduce the notation that is used in the later chapters.
- In **Chapter 4**, we investigate a hierarchy of scheduler classes for CTMDPs and propose a technique to achieve local uniformity. We prove that local uniformization preserves quantitative measures for important scheduler classes. Moreover, we introduce the new class of *late* schedulers, which outperforms all previous scheduler definitions on locally uniform CTMDPs.
- In **Chapter 5**, we apply those results and derive an approximation algorithm for time-bounded reachability probabilities in locally uniform CTMDPs. Most notably, our algorithm is quantifiably precise, that is, we prove that the computed results meet an a priori specified precision. We show the applicability of our approach by analyzing a stochastic job scheduling problem.
- In **Chapter 6**, we build upon the time-bounded reachability algorithm for locally uniform CTMDPs and develop a model checking algorithm that verifies formulas in the continuous stochastic logic [BHHK03] on IMCs. Again, the obtained analysis technique is quantifiably precise. In the last part of Chapter 6, we establish the result that CTMDPs can be transformed into alternating IMCs.

- In **Chapter 7**, we introduce bisimulation for CTMDPs and extend the continuous stochastic logic (CSL) to CTMDPs. Moreover, we prove that all measures are preserved when considering the quotient. This result justifies to use bisimulation minimization to reduce the size of the state space before applying the model checking algorithm.
- In **Chapter 8**, we propose a new semantics for GSPNs which allows for nondeterministic choices and conservatively extends stochastic activity networks. By applying our definition, we can transform GSPNs into IMCs, thereby making their analysis feasible. In the second part of Chapter 8, we show the applicability of this approach and study dependability characteristics of a workstation cluster. Moreover, we compare our results to those that are available in the literature.
- In **Chapter 9**, we mention some directions for further research and conclude.

1.5 Origins of the chapters and credits

The results presented in Chapters 6, 5, 4 and 7 are based on the following work (in that order):

- Lijun Zhang and Martin R. Neuhäuser. *Model Checking Interactive Markov Chains*. Accepted at the 16th International Conference on Tools and Algorithms for the Construction and Analysis of Systems (TACAS) 2010.
- Martin R. Neuhäuser and Lijun Zhang. *Time-Bounded Reachability in Continuous-Time Markov Decision Processes*. Technical Report, RWTH Aachen University, 2009. To be submitted.
- Martin R. Neuhäuser, Mariëlle I. A. Stoelinga and Joost-Pieter Katoen. *Delayed Nondeterminism in Continuous-Time Markov Decision Processes*. In Proceedings of the 12th International Conference on Foundations of Software Science and Computation Structures (FoSSaCS) 2009. Lecture Notes in Computer Science. Vol. 5504. 364–379. Springer Verlag.
- Martin R. Neuhäuser and Joost-Pieter Katoen. *Bisimulation and Logical Preservation for Continuous-Time Markov Decision Processes*. In Proceedings of the 18th International Conference on Concurrency Theory (CONCUR) 2007. Lecture Notes in Computer Science. Vol. 4703. 412–427. Springer Verlag.

Further publications not included in this thesis are

- Joost-Pieter Katoen, Daniel Klink and Martin R. Neuhäuser. *Compositional Abstraction for Stochastic Systems*. In Proceedings of the 7th International Conference on Formal Modeling and Analysis of Timed Systems (FORMATS) 2009. Lecture Notes in Computer Science. Vol. 5813. 195–211. Springer Verlag.

- Martin R. Neuhäuser and Thomas Noll. *Abstraction and Model Checking of Core Erlang Programs in Maude*. In Proceedings of the 6th International Workshop on Rewriting Logic and its Applications (WRLA) 2007. Electronic Notes in Theoretical Computer Science. Vol. 176. 147–163. Elsevier.

The results in Chapter 8 are new and not published yet.

2 Basics of measure & probability theory

The Axiom of Choice is
obviously true, the
well-ordering principle
obviously false, and who can
tell about Zorn's lemma?

(Prof. Jerry Lloyd Bona)

The focus of this thesis is on the analysis of stochastic systems that evolve in continuous time, which is usually modeled by the nonnegative real numbers. In the later chapters, we reason about the probability that an event occurs in a certain period of time; for example, we could be interested in the probability to leave a certain state within the next 1.5 time units.

The advantage of modeling time in a continuous domain is pretty clear, as it allows us to formalize phenomena that are best described by continuous probability distributions. Examples include the probability that a failure occurs within a certain amount of time (which usually is exponentially distributed) or the probability that a measurement error deviates by a certain percentage from its average value (which can often be described by the normal distribution).

However, we pay for this greater generality by a more complex mathematical framework: Whereas for discrete probabilistic systems (like MDPs and DTMCs), it suffices to restrict to discrete probability theory, in our continuous setting, we need the concepts of modern probability theory with its measure-theoretic background.

Therefore, this chapter provides an overview of the measure theoretic concepts which are used throughout the thesis.

In Sec. 2.1, we give an abstract introduction to measure theory. In a journey of step-wise extensions, we start with an abstract, uncountable set Ω and a measure on a class of subsets of Ω which have a simple structure. By several extensions, we subsequently increase the complexity of the sets that we are able to measure.

Section 2.2 applies the previously obtained results: Starting with the natural notion of the length of a (time) interval, we arrive at a measure on the large class of so-called Borel measurable sets.

To point out the limits of measure theory, Sec. 2.3 explains Vitali sets, which turn out to be neither Borel nor Lebesgue measurable. Hence, they provide a barrier that we may

not overcome in our extensions.

Section 2.4 explains the details of the Lebesgue integral, which allow us to integrate Borel measurable functions over sets different from the ordinary real numbers. Moreover, it is much more versatile, as it mitigates many of the restrictions of the Riemann integral.

Finally, the finite- and infinite-dimensional product spaces that we discuss in Sec. 2.5 allow us to measure the probability of sets of (finite and infinite) paths that describe the trajectories in our system models.

Most of the results presented here are taken from the excellent textbook “Probability & Measure Theory” by Robert B. Ash and Catherine A. Doléans-Dade [ADD00]. Therefore, many of the concepts explained in this section are a reproduction of those that can be found in [ADD00]. However, in contrast to Ash, we suppose a computer scientist’s background on probability theory; therefore, we strive for a compromise between the full complexity of some of the intricate measure theoretic constructions and an easier to read introductory text, where we emphasize those aspects that are useful for an understanding of the subsequent chapters. Another introduction to measure and probability theory can be found in [Bil95].

2.1 Basics of measure theory

A measure is a generalization of the concepts of “size”, “length” or “volume” which are intuitively known from Euclidean space. The aim in measure theory is to define a measure, that is, a function that assigns to each subset A of a given set Ω a value which corresponds to the size of A .

However, a measure has to satisfy certain constraints: Obviously, if $A, B \subseteq \Omega$ are subsets of Ω which do not have any element of Ω in common and if $\mu(A)$ and $\mu(B)$ denote their respective sizes, we naturally require their disjoint union $A \cup B \subseteq \Omega$ to have size $\mu(A \cup B) = \mu(A) + \mu(B)$.

Another requirement for a general definition of a measure is that if we know the size of $A \subseteq \Omega$, we should also define the size of its complement, i.e. of $A^c = \Omega \setminus A$.

Finally, it is a natural assumption to assume that the empty set should have size 0, as it does not contain any element of Ω .

As long as Ω is a finite or countably infinite set, no measure theoretic arguments are necessary. It suffices to define the size of each element $\omega \in \Omega$ and to extend this to subsets A of Ω by simply adding the elements’ sizes. Any measure defined in this way satisfies the above mentioned properties.

However, if Ω is an uncountable set, the existence of a measure that satisfies the above properties for all subsets of Ω is not guaranteed. For example, it is impossible to construct such a measure on all subsets of the real numbers. The proof and the necessary constructions can be found in Sec. 2.3.

Definition 2.1 (Field, σ -field). Let Ω be a set and $\mathfrak{F} \subseteq 2^\Omega$ a class of subsets of Ω . Then $\mathfrak{F}$ is a field iff $\mathfrak{F}$ satisfies the following conditions:

- (a) $\Omega \in \mathfrak{F}$,
- (b) $A \in \mathfrak{F} \Rightarrow A^c \in \mathfrak{F}$ and
- (c) $A_1, A_2, \dots, A_n \in \mathfrak{F} \Rightarrow \bigcup_{i=1}^n A_i \in \mathfrak{F}$.

$\mathfrak{F}$ is a σ -field iff $\mathfrak{F}$ satisfies Cond. (a) and (b) and instead of Cond. (c) it holds

- (d) $A_1, A_2, A_3, \dots \in \mathfrak{F} \Rightarrow \bigcup_{i=1}^\infty A_i \in \mathfrak{F}$.

Hence, a field $\mathfrak{F}$ is a σ -field iff for every countable family $A_1, A_2, A_3, \dots \in \mathfrak{F}$ it holds that $\bigcup_{i=1}^\infty A_i \in \mathfrak{F}$. If $\mathfrak{F} \subseteq 2^\Omega$ is a σ -field of subsets of Ω , then the tuple $(\Omega, \mathfrak{F})$ is called a *measurable space*.

Example 2.1. Let Ω be a set. According to Def. 2.1, the smallest σ -field of subsets of Ω is the set $\mathfrak{F} = \{\emptyset, \Omega\}$; the largest σ -field is the set $\mathfrak{F} = 2^\Omega$. $\diamond$

The link between measure and probability theory is established as follows: In probability theory, the set Ω is called the *sample space* and interpreted as the set of all possible outcomes (called samples) of a random experiment. Accordingly, the aim in probability theory is to measure the probability of *events*, where an event is understood as a subset of Ω which belongs to Ω 's associated σ -field $\mathfrak{F}$. Hence, measuring an event $A \in \mathfrak{F}$ yields the probability of A . In the context of probability theory, the closure properties that Def. 2.1 requires for a class of subsets of Ω to be a field, have the following informal justification: By Conditions (b) and (d), they permit to reason about the probability of the negation (A^c) and (finite and countably infinite) conjunction ($A \cup B$) of events. The sample space Ω is understood as the set of all possible outcomes of the random experiment; accordingly, the probability that the outcome of a random experiment falls within Ω is 1. Therefore, Ω is the *certain event* and included in $\mathfrak{F}$. As $\mathfrak{F}$ is closed under complement, the set $\Omega^c = \emptyset$ is in $\mathfrak{F}$ as well; it is the *impossible event*, which is assigned probability 0.

Example 2.2. Let Ω be a countably infinite set and define $\mathfrak{F}_0$ as the smallest class of subsets of Ω such that for all $A \subseteq \Omega$:

$$|A| < +\infty \Rightarrow A \in \mathfrak{F}_0 \quad \text{and} \quad A \in \mathfrak{F}_0 \Rightarrow A^c \in \mathfrak{F}_0.$$

Note that the definition is non-trivial, i.e. in general $\mathfrak{F}_0 \subsetneq 2^\Omega$: For example, if $\Omega = \mathbb{N}$, then the set $\{2n \mid n \in \mathbb{N}\}$ of even numbers is not in $\mathfrak{F}_0$, as both $\{2n \mid n \in \mathbb{N}\}$ and $\{2n+1 \mid n \in \mathbb{N}\}$ are countably infinite sets.

In order to show that $\mathfrak{F}_0$ is a field, we check the properties required by Def. 2.1: By definition, $\mathfrak{F}_0$ is closed under complement; hence, Cond. (b) is satisfied. For Cond. (a), note

that $|\emptyset| = 0 < +\infty$ implies $\emptyset \in \mathfrak{F}_0$. As $\mathfrak{F}_0$ is closed under complement, $\emptyset \in \mathfrak{F}_0$ implies $\emptyset^c = \Omega \in \mathfrak{F}_0$; hence $\mathfrak{F}_0$ satisfies Cond. (a). For Cond. (c), let $A, B \in \mathfrak{F}_0$. If both $|A| < +\infty$ and $|B| < +\infty$, then $|A \cup B| < +\infty$ and $A \cup B \in \mathfrak{F}_0$. For the other cases, assume w.l.o.g. that $|A| = +\infty$. By definition of $\mathfrak{F}_0$, $|A| = +\infty$ implies $|A^c| < +\infty$ (otherwise, $A \notin \mathfrak{F}_0$). Therefore $|A^c \cap B^c| < +\infty$ and $(A^c \cap B^c) \in \mathfrak{F}_0$. As $\mathfrak{F}_0$ is closed under complement, this implies that $(A^c \cap B^c)^c \in \mathfrak{F}_0$ and by De Morgan's law, we conclude that $(A^c \cap B^c)^c = (A \cup B) \in \mathfrak{F}_0$. Hence, $\mathfrak{F}_0$ is closed under finite union.

Lemma 2.1 (Generated σ -field). Let $\mathcal{J} \subseteq 2^\Omega$ be a class of subsets of some set Ω and define

$$\sigma(\mathcal{J}) = \bigcap \{ \mathfrak{F} \subseteq 2^\Omega \mid \mathfrak{F} \text{ is a } \sigma\text{-field, } \mathcal{J} \subseteq \mathfrak{F} \}.$$

Then $\sigma(\mathcal{J})$ is the smallest σ -field which contains $\mathcal{J}$. It is called the smallest σ -field generated by $\mathcal{J}$.

Proof. Let $\mathfrak{J} = \{ \mathfrak{F} \subseteq 2^\Omega \mid \mathfrak{F} \text{ is a } \sigma\text{-field, } \mathcal{J} \subseteq \mathfrak{F} \}$.

First, we prove that $\sigma(\mathcal{J})$ is a field: Therefore, we check Conditions (a), (b) and (d) of Def. 2.1: For Cond. (a), note that $\Omega \in \mathfrak{F}$ for all $\mathfrak{F} \in \mathfrak{J}$; hence, $\Omega \in \sigma(\mathcal{J})$. For Cond. (b), let $A \in \sigma(\mathcal{J})$. Then $A \in \mathfrak{F}$ for all $\mathfrak{F} \in \mathfrak{J}$, implying $A^c \in \mathfrak{F}$ for all $\mathfrak{F} \in \mathfrak{J}$. Hence, $A^c \in \sigma(\mathcal{J})$. Finally, $\sigma(\mathcal{J})$ satisfies Cond. (d): If $A_1, A_2, \dots \in \mathfrak{J}$, then $A_1, A_2, \dots \in \mathfrak{F}$ for all $\mathfrak{F} \in \mathfrak{J}$; as each $\mathfrak{F}$ is a σ -field, it holds that $\bigcup_{i=1}^\infty A_i \in \mathfrak{F}$ for all $\mathfrak{F} \in \mathfrak{J}$. Therefore $\bigcup_{i=1}^\infty A_i \in \sigma(\mathcal{J})$. Thus, $\sigma(\mathcal{J})$ is a σ -field.

By definition, $\mathcal{J} \subseteq 2^\Omega$. Further, 2^Ω is a σ -field. This implies that $2^\Omega \in \mathfrak{J}$ so that $\mathfrak{J}$ is nonempty. Furthermore, $\mathcal{J} \subseteq \mathfrak{F}$ for all $\mathfrak{F} \in \mathfrak{J}$. Hence $\mathcal{J} \in \sigma(\mathcal{J})$.

Finally, if $\mathfrak{F}'$ is a σ -field of subsets of Ω with $\mathcal{J} \subseteq \mathfrak{F}'$, then $\mathfrak{F}' \in \mathfrak{J}$ and $\sigma(\mathcal{J}) \subseteq \mathfrak{F}'$. Hence, $\sigma(\mathcal{J})$ is the *smallest* σ -field that contains $\mathcal{J}$. $\square$

Definition 2.2 (Measure, probability measure). A measure μ on a measurable space $(\Omega, \mathfrak{F})$ is a function $\mu : \mathfrak{F} \rightarrow \mathbb{R}_{\geq 0}^\infty$ such that for all finite or countably infinite families $\{A_i\}_{i \in I}$ of pairwise disjoint sets $A_i \in \mathfrak{F}$ (where $I \subseteq \mathbb{N}$), it holds that

$$\mu \left(\bigcup_{i \in I} A_i \right) = \sum_{i \in I} \mu(A_i). \quad (2.1)$$

If $\mu(\Omega) = 1$, then μ is a probability measure.

Any measurable space $(\Omega, \mathfrak{F})$ together with a measure μ forms a *measure space*, denoted by the triple $(\Omega, \mathfrak{F}, \mu)$. If μ is a probability measure, the measurable space $(\Omega, \mathfrak{F}, \mu)$ is a *probability space*.

For what follows, we generalize the notion of a measure to also account for *fields* (instead of σ -fields as required in Def. 2.2): Therefore, let Ω be a set and $\mathfrak{F}_0$ a field of subsets of Ω . A set function $\mu : \mathfrak{F}_0 \rightarrow \mathbb{R}^\infty$ on $\mathfrak{F}_0$ is *countably additive* on $\mathfrak{F}_0$ iff $\mu(\bigcup_{i \in I} A_i) = \sum_{i \in I} \mu(A_i)$ for all finite or countably infinite families $\{A_i\}_{i \in I}$ of pairwise disjoint sets $A_i \in \mathfrak{F}_0$ (where $I \subseteq \mathbb{N}$) that satisfy $\bigcup_{i \in I} A_i \in \mathfrak{F}_0$. Observe the intricate point in this definition: For μ to be countably additive on a field, it suffices to consider only those countably infinite collections of disjoint sets, whose union actually belongs to $\mathfrak{F}_0$: As $\mathfrak{F}_0$ is only a field (and not a σ -field), there may exist countably infinite collections $A_1, A_2, \dots$ of disjoint sets $A_i \in \mathfrak{F}_0$ such that $\bigcup_{i=1}^\infty A_i \notin \mathfrak{F}_0$.

Accordingly, we extend Def. 2.2 and call a set function $\mu : \mathfrak{F}_0 \rightarrow \mathbb{R}^\infty$ on a field $\mathfrak{F}_0$ a *measure on the field $\mathfrak{F}_0$* iff μ is countably additive on $\mathfrak{F}_0$ and $\mu(A) \geq 0$ for all $A \in \mathfrak{F}_0$. Further, if $\mu(\Omega) = 1$, μ is called a *probability measure on the field $\mathfrak{F}_0$* . Note that if $\mathfrak{F}_0$ is not only a field but also a σ -field and μ is countably additive and nonnegative, then μ is a measure according to Def. 2.2.

Naturally, finite additivity is a weaker condition than countable additivity: We say that a set function $\mu : \mathfrak{F}_0 \rightarrow \mathbb{R}^\infty$ is *finitely additive* iff $\mu(\bigcup_{i=1}^n A_i) = \sum_{i=1}^n \mu(A_i)$ for all finite collections $A_1, A_2, \dots, A_n$ of pairwise disjoint sets $A_i \in \mathfrak{F}_0$.

Further, a set function $\mu : \mathfrak{F}_0 \rightarrow \mathbb{R}_{\geq 0}^\infty$ is σ -finite on a field $\mathfrak{F}_0$ iff there exists a collection $A_1, A_2, \dots \in \mathfrak{F}_0$ such that $\Omega = \bigcup_{i=1}^\infty A_i$ and $\mu(A_i) < +\infty$ for all $i \in \mathbb{N}$. Thus, if μ is σ -finite, we can build Ω from an at most countably infinite collection of sets in $\mathfrak{F}_0$ that all have a finite measure.

Example 2.3. Reconsider the field $\mathfrak{F}_0$ from Ex. 2.2 and define the set function μ on $\mathfrak{F}_0$ such that $\mu(A) = 0$ if $|A| < +\infty$ and $\mu(A) = 1$, otherwise. Then μ is finitely additive, but not countably additive: Let $A_1, A_2, \dots, A_n$ be pairwise disjoint sets in $\mathfrak{F}_0$. To show finite additivity, we consider two cases:

First, assume that $|A_k| = +\infty$ for at least one $k \in \{1, 2, \dots, n\}$. Then $\mu(\bigcup_{i=1}^n A_i) = 1$. To show that $\sum_{i=1}^n \mu(A_i) = 1$ holds as well, recall that by definition of $\mathfrak{F}_0$, it holds that $|A_k| = +\infty$ implies $|A_k^c| < +\infty$. As $A_i \subseteq A_k^c$ for all $i \neq k$, we derive $|A_i| < +\infty$; thus $\mu(A_i) = 0$ for all $i \neq k$ by definition of μ and $\mathfrak{F}_0$. Hence, $\sum_{i=1}^n \mu(A_i) = \mu(A_k) = 1$ and therefore $\mu(\bigcup_{i=1}^n A_i) = \sum_{i=1}^n \mu(A_i)$.

For the second case, assume that $|A_i| < +\infty$ for all $i \in \{1, 2, \dots, n\}$. Then $\mu(\bigcup_{i=1}^n A_i) = 0 = \sum_{i=1}^n \mu(A_i)$. Thus μ is finitely additive.

On the other hand, it is easy to see that μ is not countably additive: Let $\omega_1, \omega_2, \dots$ be an enumeration of the elements in Ω and define $A_i = \{\omega_i\}$. Then $\sum_{i=1}^\infty \mu(A_i) = 0$, but $\mu(\bigcup_{i=1}^\infty A_i) = \mu(\Omega) = 1$. $\diamond$

By definition, any σ -field $\mathfrak{F}$ is closed under countable union; hence, if $A_1 \subseteq A_2 \subseteq \dots$ is an increasing sequence of sets $A_i \in \mathfrak{F}$, its limit $\lim_{i \rightarrow \infty} A_i = \bigcup_{i=1}^\infty A_i$ is an element of $\mathfrak{F}$. Therefore, σ -fields are closed under increasing sequences. Moreover, σ -fields are also closed under decreasing sequences, i.e. if $A_1 \supseteq A_2 \supseteq \dots$ are elements in $\mathfrak{F}$, then $\bigcap_{i=1}^\infty A_i \in \mathfrak{F}$. To see this, note that any σ -field $\mathfrak{F}$ is closed under complement and countable union. Hence, it is also closed under countable intersection and $\bigcap_{i=1}^\infty A_i \in \mathfrak{F}$.

The obvious next question is whether measures, or more generally, countably additive set functions agree with these closure properties of σ -fields:

Lemma 2.2 (Continuity of countably additive set functions). *Let $\mathfrak{F}$ be a σ -field of subsets of some set Ω and let $\mu : \mathfrak{F} \rightarrow \mathbb{R}^\infty$ be a countably additive set function on $\mathfrak{F}$.*

- (a) *If $A_1 \subseteq A_2 \subseteq A_3 \subseteq \dots \in \mathfrak{F}$ and $A_i \uparrow A$, then $\lim_{i \rightarrow \infty} \mu(A_i) = \mu(A)$.*
- (b) *If $A_1 \supseteq A_2 \supseteq A_3 \supseteq \dots \in \mathfrak{F}$ such that $A_i \downarrow A$ and $-\infty < \mu(A_i) < +\infty$ for all $i \in \mathbb{N}$, then $\lim_{i \rightarrow \infty} \mu(A_i) = \mu(A)$.*

Proof. For a proof, see [ADD00, Th. 1.2.7]. □

Although Lemma 2.2 is stated in full generality, note that any measure μ on $(\Omega, \mathfrak{F})$ is a nonnegative, countable additive set function. Hence, the statements (a) and (b) in Lemma 2.2 hold for any measure.

2.1.1 Extension from $\mathfrak{F}_0$ to $\sigma(\mathfrak{F}_0)$

In general, if Ω is an uncountable set like the set of real numbers, and we are to define a measure μ on all subsets of Ω , it turns out that this is impossible (see Sec. 2.3). More precisely, if we insist on the natural assumption that a measure should be countably additive (cf. Def. 2.2(2.1)), we cannot define a measure on the σ -field 2^Ω : This is due to the fact, that in general (for example, on $2^\mathbb{R}$) there exist subsets of Ω such that no countably additive set function can be defined on 2^Ω .

As a consequence, if Ω is countably infinite, we are forced to restrict ourselves to the subclass of measurable subsets of Ω . This can be achieved as follows: First, we identify those subsets of Ω that we need to measure. In a second step, we need to find a field $\mathfrak{F}_0$ which contains those desirable sets and allows us to define the corresponding measure on $\mathfrak{F}_0$. Note that due to the simple structure of a field, this is usually an easy task.

However, there are important properties (like the measure of the limit of in- or decreasing sequences) that require to extend μ from the field $\mathfrak{F}_0$ to the smallest σ -field $\sigma(\mathfrak{F}_0)$ that is generated by $\mathfrak{F}_0$. This is a nontrivial task, as it turns out that the structure of the elements in the σ -field $\sigma(\mathfrak{F}_0)$ is much more complex than the structure of the elements of its underlying field $\mathfrak{F}_0$.

Therefore, this section introduces the measure theoretic results that guarantee the existence (and uniqueness) of the extension of μ from $\mathfrak{F}_0$ to $\sigma(\mathfrak{F}_0)$. In what follows, we obtain an easier description if we assume that μ is a finite measure, that is, $\mu(A) < +\infty$ for all $A \in \mathfrak{F}_0$. As we shall see later, this restriction is too strict; in fact, we already obtain a unique extension of μ from $\mathfrak{F}_0$ to $\sigma(\mathfrak{F}_0)$ if we assume that μ is σ -finite on $\mathfrak{F}_0$; however, this result is easily established later, so that we do not lose anything if we restrict to finite measures first.

In the following, we proceed stepwise and extend μ to more and more complex classes of subsets of Ω , until we arrive at $\sigma(\mathfrak{F}_0)$. The first step is to extend μ to the class $\mathcal{G}$ of all countable unions of elements in $\mathfrak{F}_0$. Note that in contrast to the first impression, $\mathcal{G}$ is a strict subset of $\sigma(\mathfrak{F}_0)$ and should not be confused with the latter!

Extension to countable unions of elements in $\mathfrak{F}_0$.

To begin with, consider the class $\mathcal{G} \subseteq 2^\Omega$ of subsets of Ω which is defined such that

$$A \in \mathcal{G} \iff \exists A_1, A_2, \dots \in \mathfrak{F}_0. A_i \uparrow A.$$

Thus, $\mathcal{G}$ is the set of all limits of increasing sequences of elements in $\mathfrak{F}_0$; further, $\mathfrak{F}_0 \subseteq \mathcal{G}$, as for any set $A \in \mathfrak{F}_0$, the sequence which is obtained by defining $A_i = A$ for all $i \in \mathbb{N}$ increases to A .

Note that $\mathcal{G}$ is also the class of all countable unions of elements in $\mathfrak{F}_0$: To see this, let $A_1, A_2, \dots \in \mathfrak{F}_0$ and define the sets $B_k = \bigcup_{i=1}^k A_i$ and $A = \bigcup_{i=1}^\infty A_i$. Each B_k is a finite union of elements in $\mathfrak{F}_0$ and therefore, $B_k \in \mathfrak{F}_0$. Moreover, $B_k \uparrow A$ by construction. Thus, by definition of $\mathcal{G}$ it holds that $A \in \mathcal{G}$. Hence, $\mathcal{G}$ contains all countable unions of elements in $\mathfrak{F}_0$. To show that $\mathcal{G}$ does not contain more, consider the reverse direction: If $A \in \mathcal{G}$, then there exists an increasing sequence $A_1, A_2, \dots \in \mathfrak{F}_0$ such that $A_i \uparrow A$. But then $A = \bigcup_{i=1}^\infty A_i$ is a countable union of elements in $\mathfrak{F}_0$.

Now that we have defined the class $\mathcal{G}$ of subsets of Ω , we extend the measure μ from the field $\mathfrak{F}_0$ to $\mathcal{G}$:

Lemma 2.3 (Extension of μ to $\mathcal{G}$). *Let $\mathfrak{F}_0$ be a field and μ a finite measure on $\mathfrak{F}_0$. Further, let $\mathcal{G}$ be the class of all countable unions of elements in $\mathfrak{F}_0$. Then $\mu' : \mathcal{G} \rightarrow \mathbb{R}_{\geq 0}$ denotes the extension of μ from $\mathfrak{F}_0$ to $\mathcal{G}$. For $A \in \mathcal{G}$, we define*

$$\mu'(A) = \lim_{n \rightarrow \infty} \mu(A_n),$$

where $A_1, A_2, \dots \in \mathfrak{F}_0$ are such that $A_n \uparrow A$. Then it holds:

(a) $\mu'(A) = \mu(A)$ for all $A \in \mathfrak{F}_0$.

(b) If $G_1, G_2, (G_1 \cup G_2), (G_1 \cap G_2) \in \mathcal{G}$, then

$$\mu'(G_1 \cup G_2) + \mu'(G_1 \cap G_2) = \mu'(G_1) + \mu'(G_2).$$

(c) If $G_1, G_2 \in \mathcal{G}$ and $G_1 \subseteq G_2$, then $\mu'(G_1) \leq \mu'(G_2)$.

(d) If $G_1, G_2, \dots \in \mathcal{G}$ and $G_n \uparrow G$, then $G \in \mathcal{G}$ and $\lim_{n \rightarrow \infty} \mu'(G_n) = \mu'(G)$.

Proof. A proof can be found in [ADD00, Lemma 1.3.2]. □

First, note that by definition of $\mathcal{G}$, there exists a sequence $A_1, A_2, \dots \in \mathfrak{F}_0$ that increases to A ; further, if $A'_1, A'_2, \dots \in \mathfrak{F}_0$ is another sequence with $A'_n \uparrow A$, it can be shown that $\lim_{n \rightarrow \infty} \mu(A_n) = \lim_{n \rightarrow \infty} \mu(A'_n)$ [ADD00, Lemma 1.3.1]. Hence, μ' is well-defined.

Observe that μ' satisfies the requirements that we expect from a measure, i.e. by (a) it coincides with the original measure μ on $\mathfrak{F}_0$, by (d) it preserves limits, by (b) it works as expected for (not necessarily disjoint) set union and finally, by (c) it obeys the ordering on the measures of sets according to set inclusion.

However, at this stage the extension is not complete, as $\mathcal{G}$ is not a σ -field yet. Hence, there are still sets in $\sigma(\mathfrak{F}_0) \setminus \mathcal{G}$ that μ' is unable to measure. As an example, note that the class $\mathcal{G}$ is not closed under complement: We derive $\mathcal{G}$ by extending $\mathfrak{F}_0$ to the class of all countable unions of elements in $\mathfrak{F}_0$; however, $\mathcal{G}$ is closed under complement only with respect to elements in $\mathfrak{F}_0$. More precisely, if $A = \bigcup_{i=1}^{\infty} A_i$ with $A_i \in \mathfrak{F}_0$ is a countable union that does not belong to $\mathfrak{F}_0$, then $A \in \mathcal{G}$ still holds by definition of $\mathcal{G}$. However, this does not imply that $A^c \in \mathcal{G}$. To see this, note that the set A^c cannot always be represented as a countable union of elements in $\mathfrak{F}_0$. Therefore, in general, $A^c \notin \mathcal{G}$ so that $\mathcal{G}$ is not closed under complement. We postpone the construction of a concrete counterexample and refer the reader to Ex. 2.5 on page 26 for further details.

Therefore, although Lemma 2.3 considerably extends the domain of μ , we still do not cover all desirable subsets of Ω . This problem is overcome (only partly, as we will see) in the next step:

Extension to an outer measure.

With $\mu' : \mathcal{G} \rightarrow \mathbb{R}_{\geq 0}$ and the class $\mathcal{G}$, we have extended the measure μ on $\mathfrak{F}_0$ to a larger class of subsets of Ω . Now we aim at an extension of μ' to an *outer measure* which is defined on the entire power set 2^Ω :

Definition 2.3 (Outer measure). An outer measure on a set Ω is a set function $\lambda : 2^\Omega \rightarrow \mathbb{R}_{\geq 0}^\infty$ that satisfies

- (a) $\lambda(\emptyset) = 0$,
- (b) if $A, B \subseteq \Omega$ and $A \subseteq B$, then $\lambda(A) \leq \lambda(B)$ and
- (c) if $A_1, A_2, \dots \subseteq \Omega$, then $\lambda\left(\bigcup_{n=1}^{\infty} A_n\right) \leq \sum_{n=1}^{\infty} \lambda(A_n)$.

It is important to note that Cond. (c) (which is also called countable subadditivity) does neither require the sets A_n to be disjoint, nor does it state that $\lambda\left(\bigcup_{n=1}^{\infty} A_n\right) = \sum_{n=1}^{\infty} \lambda(A_n)$ holds if they happen to be pairwise disjoint (which is required in Def. 2.2 for λ to be a measure)! Hence, we could suspect already here that something is wrong with extending μ' to a measure on 2^Ω .

In fact, albeit its name, an *outer measure* is not a measure in general. In our case, it will turn out that by extending μ' to 2^Ω , the extension loses important properties of a measure. Before we address this issue, let us define how to extend μ' to an outer measure on all subsets of Ω :

Lemma 2.4 (Extension to an outer measure). *Let $\mathfrak{F}_0$ be a field of subsets of some set Ω , $\mathcal{G}$ the class of all countable unions of elements in $\mathfrak{F}_0$ and μ' the extension of a finite measure μ on $\mathfrak{F}_0$ to $\mathcal{G}$. Define the set function*

$$\mu^* : 2^\Omega \rightarrow \mathbb{R}_{\geq 0}^\infty : A \mapsto \inf \{ \mu'(B) \mid B \supseteq A \wedge B \in \mathcal{G} \}.$$

Then μ^ is an outer measure on Ω with the additional properties that*

- (a) $\mu^*(A) = \mu'(A)$ for all $A \in \mathcal{G}$,
- (b) $\mu^*(A \cup B) + \mu^*(A \cap B) \leq \mu^*(A) + \mu^*(B)$ for all $A, B \subseteq \Omega$ and
- (c) if $A_1, A_2, \dots \subseteq \Omega$ with $A_n \uparrow A$, then $\lim_{n \rightarrow \infty} \mu^*(A_n) = \mu^*(A)$.

Proof. The proof can be found in, e.g. [ADD00, p.16ff]. □

This definition of μ^* provides an extension of μ' to the whole power set of Ω . Note however, that countable additivity which is required for μ^* to be a measure on 2^Ω (cf. Eq. (2.1) of Def. 2.2) is replaced by the weaker property of subadditivity in Def. 2.3(c). In fact, it turns out that in general, μ^* is not countably additive on all subsets of Ω , that is, there exist sequences $A_1, A_2, \dots \subseteq \Omega$ of pairwise disjoint sets A_n such that $\mu^*(\bigcup_{n=1}^\infty A_n) < \sum_{n=1}^\infty \mu^*(A_n)$.

By the above argument, extending μ' to the whole power set 2^Ω is too ambitious. Therefore, to still obtain a measure, we have to exclude certain elements in 2^Ω and restrict to a σ -field smaller than 2^Ω . In the following, we identify a large (but proper) subset of 2^Ω that is a σ -field and allows an extension of μ that is countably additive:

Lemma 2.5 (Extension of finite measures). *Let $\mathfrak{F}_0$ be a field of subsets of a set Ω , μ a finite measure on $\mathfrak{F}_0$ and $\mathcal{G}$ the class of all countable unions of elements in $\mathfrak{F}_0$. For the outer measure μ^* defined as above, let*

$$\mathcal{H} = \{ H \subseteq \Omega \mid \mu^*(H) + \mu^*(H^c) = \mu(\Omega) \}.$$

Then $\mathcal{H}$ is a σ -field and μ^ is a measure on $\mathcal{H}$.*

Proof. The proof can be found in [ADD00, Thm. 1.3.5]. □

To see that the class $\mathcal{H}$ indeed extends $\mathcal{G}$, let $A \in \mathcal{G}$. By definition of $\mathcal{G}$, there exists an increasing sequence $A_1, A_2, \dots \in \mathfrak{F}_0$ such that $A_n \uparrow A$, implying that $A^c \subseteq A_n^c$ for all $n \in \mathbb{N}$. As μ^* is an outer measure, it holds by Def. 2.3(b) that $\mu^*(A^c) \leq \mu^*(A_n^c)$. Further, recall that μ^* agrees with μ' on $\mathcal{G}$ and with μ on $\mathfrak{F}_0$; hence

$$\mu(A_n) + \mu^*(A^c) \leq \mu(A_n) + \mu(A_n^c) = \mu(\Omega). \quad (2.2)$$

Further, $\lim_{n \rightarrow \infty} \mu'(A_n) = \mu'(A)$ by Lemma 2.3(d). Hence, taking the limit for $n \rightarrow \infty$ on both sides of Eq. (2.2) yields $\mu^*(A) + \mu^*(A^c) \leq \mu(\Omega)$.

On the other hand, Lemma 2.4(b) implies that $\mu^*(A \cup A^c) + \mu^*(A \cap A^c) \leq \mu'(A) + \mu^*(A^c)$; as $\mu^*(A \cup A^c) = \mu(\Omega)$ and $\mu^*(A \cap A^c) = \mu(\emptyset) = 0$, we obtain $\mu(\Omega) \leq \mu'(A) + \mu^*(A^c)$. Further, $\mu'(A) = \mu^*(A)$ by Lemma 2.4(a). Hence, $\mu^*(A) + \mu^*(A^c) \geq \mu(\Omega)$.

Therefore we have established that $\mu^*(A) + \mu^*(A^c) = \mu(\Omega)$ and $A \in \mathcal{H}$. As this applies to all $A \in \mathcal{G}$, this proves that $\mathcal{G} \subseteq \mathcal{H}$.

The class $\mathcal{H}$ has another important property: By transitivity of set inclusion, we conclude from the fact that $\mathcal{G} \subseteq \mathcal{H}$ and $\mathfrak{F}_0 \subseteq \mathcal{G}$, that $\mathfrak{F}_0 \subseteq \mathcal{H}$. Moreover, by Lemma 2.5 we know that $\mathcal{H}$ is a σ -field of subsets of Ω . But by definition, $\sigma(\mathfrak{F}_0)$ is the *smallest* σ -field that contains $\mathfrak{F}_0$. Hence, $\sigma(\mathfrak{F}_0) \subseteq \mathcal{H}$.

To summarize the different steps in extending μ from $\mathfrak{F}_0$ to $\sigma(\mathfrak{F}_0)$, Table 2.1 depicts the complete chain of inclusions (from left to right) as well as the corresponding extensions of μ and their properties.

As we have seen, $\sigma(\mathfrak{F}_0)$ and $\mathcal{H}$ are both σ -fields that contain the field $\mathfrak{F}_0$; further, we are able to extend μ to a measure on $\sigma(\mathfrak{F}_0)$ and $\mathcal{H}$. Hence $\sigma(\mathfrak{F}_0)$ and $\mathcal{H}$ seem to be related closely. In fact, it turns out that they differ only in sets of measure zero. More precisely, it can be shown (see [ADD00, Thm. 1.3.8]) that any element $A \in \mathcal{H}$ can be decomposed such that $A = B \cup M$, where $B \in \sigma(\mathfrak{F}_0)$ and $M \subseteq N$ is a subset of some set $N \in \sigma(\mathfrak{F}_0)$ which has measure zero, i.e. $\mu^*(N) = 0$. Therefore, we say that $\mathcal{H}$ is the completion of $\sigma(\mathfrak{F}_0)$ with respect to μ^* and sets of measure zero:

Definition 2.4 (Completion of a measure space). Let $(\Omega, \mathfrak{F}, \mu)$ be a measure space. Then

$$\mathfrak{F}^\mu = \{A \cup M \mid A \in \mathfrak{F}, M \subseteq N, N \in \mathfrak{F}, \mu(N) = 0\}$$

is the completion of $\mathfrak{F}$ with respect to the measure μ . Further, a measure space $(\Omega, \mathfrak{F}, \mu)$ is complete iff for all $N \in \mathfrak{F}$, $\mu(N) = 0$ implies that $M \in \mathfrak{F}$ for all $M \subseteq N$.

Therefore, we complete a measure space $(\Omega, \mathfrak{F}, \mu)$ by extending any set $A \in \mathfrak{F}$ with all subsets of sets of measure zero which are in $\mathfrak{F}$. Further, it directly follows from Def. 2.4 that the completion of a measure space is indeed complete.

Using the construction outlined above (i.e. from $\mathfrak{F}_0$ over $\mathcal{G}$ to 2^Ω and back via $\mathcal{H}$ to $\sigma(\mathfrak{F}_0)$), we are now able to state the first important result regarding the extension of a finite measure μ on $\mathfrak{F}_0$ to the smallest σ -field generated by $\mathfrak{F}_0$:

$\mathfrak{F}_0$	$\mathcal{G}$	$\sigma(\mathfrak{F}_0)$	$\mathcal{H}$	2^Ω
field	limit collection	smallest σ -field	completion of $\sigma(\mathfrak{F}_0)$	power set
μ	μ'	$\mu^*_{ \sigma(\mathfrak{F}_0)}$	$\mu^*_{ \mathcal{H}}$	μ^*
measure on $\mathfrak{F}_0$	set function	measure	measure	not countably additive

Table 2.1: Summary of the inclusions and the properties of the extensions of μ .

Theorem 2.1 (Existence of an extension). *A finite measure μ on a field $\mathfrak{F}_0$ can be extended to a measure on $\sigma(\mathfrak{F}_0)$.*

Proof. We have shown before that $\mathfrak{F}_0 \subseteq \mathcal{G} \subseteq \sigma(\mathfrak{F}_0) \subseteq \mathcal{H} \subseteq 2^\Omega$. Further, μ^* is an extension of μ to 2^Ω . Hence, the domain of μ^* covers $\sigma(\mathfrak{F}_0)$. Moreover μ^* is a finite measure on $\mathcal{H}$ by Lemma 2.5 and $\sigma(\mathfrak{F}_0) \subseteq \mathcal{H}$. Hence, the restriction of μ^* to $\sigma(\mathfrak{F}_0)$ is the desired finite measure on $\sigma(\mathfrak{F}_0)$. $\square$

With this result, we are able to extend μ from $\mathfrak{F}_0$ to $\sigma(\mathfrak{F}_0)$ and even more, to $\mathcal{H}$. Recall that it can be proved (see Sec. 2.3 for the details of the construction) that we cannot extend μ to a measure on the σ -field 2^Ω . However, the question whether there exist σ -fields that are larger than $\sigma(\mathfrak{F}_0)$ and $\mathcal{H}$ (but smaller than 2^Ω), which allow for an extension, is not answered by the preceding constructions. Within this thesis, we only refer to [Ben76, p. 40] which provides links to the related literature.

Although Thm. 2.1 allows us to extend any finite measure μ to the σ -field $\sigma(\mathfrak{F}_0)$, we do not know whether this extension is unique: More precisely, the question to be answered is: Does there exist another measure λ on $\sigma(\mathfrak{F}_0)$ such that $\mu = \lambda$ on $\mathfrak{F}_0$ but $\mu(A) \neq \lambda(A)$ for some set $A \in \sigma(\mathfrak{F}_0)$? The answer to this question will be the topic of the next section:

2.1.2 Uniqueness of the extension

Starting from a finite measure μ on some field $\mathfrak{F}_0$ of subsets of a set Ω , we have extended μ to a set function μ' on the class $\mathcal{G}$ that contains all limits of increasing sequences of sets in $\mathfrak{F}$; then, we have shown that the outer measure μ^* which is induced by μ' , is a finite measure on the class $\mathcal{H}$ of subsets of Ω . As $\sigma(\mathfrak{F}_0)$ is a subset of $\mathcal{H}$, we can consider μ^* as an extension of μ to the smallest σ -field generated by $\mathfrak{F}_0$. What remains to discuss is the uniqueness of our extension: Stated differently, does there exist another measure λ defined on $\sigma(\mathfrak{F}_0)$ such that μ and λ agree on sets in $\mathfrak{F}_0$ (i.e. $\mu^*(A) = \lambda(A)$ for all $A \in \mathfrak{F}_0$) while their extensions to $\sigma(\mathfrak{F}_0)$ differ (i.e. $\exists A \in \sigma(\mathfrak{F}_0). \mu^*(A) \neq \lambda^*(A)$)?

At the end of this section, we will answer this question in the negative, that is, the extension of μ is unique. The following theorem, the so-called monotone class theorem, is essential in proving this result. In fact, it provides the basis for a proof technique, where

it suffices to show a property on a *monotone class* to prove it for the entire σ -field. The only restriction is that the monotone class must be “large enough”, that is, it must contain at least all elements of the underlying field:

Definition 2.5 (Monotone class). Let $\mathcal{X}$ be a class of subsets of Ω . $\mathcal{X}$ is a monotone class iff for all collections $A_1, A_2, \dots \in \mathcal{X}$:

(a) $A_n \uparrow A \Rightarrow A \in \mathcal{X}$ and

(b) $A_n \downarrow A \Rightarrow A \in \mathcal{X}$.

Thus, any class of subsets of some set Ω which is closed under increasing and decreasing sequences is a monotone class.

Theorem 2.2 (Monotone class theorem). Let $\mathcal{X}$ be a monotone class over subsets of some set Ω and let $\mathfrak{F}_0$ be a field of subsets of Ω . If $\mathfrak{F}_0 \subseteq \mathcal{X}$, then $\sigma(\mathfrak{F}_0) \subseteq \mathcal{X}$.

Proof. A proof can be found in [ADD00, Thm. 1.6.2]. □

The monotone class theorem is extremely useful: We use it in the proof of Lemma 4.7 in Sec. 4.2.2 as well as in the next theorem to show that properties which hold for all elements in a field $\mathfrak{F}_0$ also hold for all elements in $\sigma(\mathfrak{F}_0)$.

The Carathéodory extension theorem is the main result of this section. It states that the extension of a finite measure μ from a field $\mathfrak{F}_0$ to the measure μ^* on $\sigma(\mathfrak{F})$ is unique. Moreover, it relaxes the restriction to finite measures that we have imposed so far:

Theorem 2.3 (Carathéodory extension theorem). Let μ be a σ -finite measure on a field $\mathfrak{F}_0$ of subsets of some set Ω . Then μ has a unique extension to a measure on $\sigma(\mathfrak{F}_0)$.

Proof. As the Carathéodory extension theorem is essential to measure theory and demonstrates a basic proof technique, we give a detailed proof here. It is split in two parts:

- We relax the restriction of μ of being a finite measure and allow μ to be σ -finite. Thus, there exist sets $A'_1, A'_2, \dots \in \mathfrak{F}_0$ such that $\bigcup_{i=1}^{\infty} A'_i = \Omega$ and $\mu(A'_i) < +\infty$ for all $i \in \mathbb{N}$. Now, define $A_n = A'_n \setminus \bigcup_{i=1}^{n-1} A'_i$. Then the sets A_n are pairwise disjoint and $\Omega = \bigcup_{n=1}^{\infty} A_n$ and $\mu(A_n) \leq \mu(A'_n) < +\infty$ for all $n \in \mathbb{N}$.

Now, define a family of measures μ_n on $\mathfrak{F}_0$ (for $n = 1, 2, \dots$) such that $\mu_n(A) = \mu(A \cap A_n)$. Each μ_n is a finite measure (because $\mu(A_n) < +\infty$) and has an extension μ_n^* to $\sigma(\mathfrak{F}_0)$. As the A_n are pairwise disjoint, it holds that $\mu(A) = \mu(A \cap \Omega) =$

$\mu(\bigcup_{n=1}^{\infty} (A \cap A_n)) = \sum_{n=1}^{\infty} \mu(A \cap A_n) = \sum_{n=1}^{\infty} \mu_n(A)$. Hence, the set function that is obtained by defining $\mu^*(A) = \sum_{n=1}^{\infty} \mu_n^*(A)$ for all $A \in \sigma(\mathfrak{F}_0)$ is an extension of μ . To prove that it is a measure, we check the condition of Def. 2.2: Let $B_1, B_2, \dots \in \sigma(\mathfrak{F}_0)$ be a sequence of pairwise disjoint sets in $\mathfrak{F}$. Then

$$\mu^*\left(\bigcup_{i=1}^{\infty} B_i\right) = \sum_{n=1}^{\infty} \mu_n^*\left(\bigcup_{i=1}^{\infty} B_i\right) = \sum_{n=1}^{\infty} \sum_{i=1}^{\infty} \mu_n^*(B_i) = \sum_{i=1}^{\infty} \sum_{n=1}^{\infty} \mu_n^*(B_i) = \sum_{i=1}^{\infty} \mu^*(B_i).$$

Therefore, μ^* is a measure on $\sigma(\mathfrak{F}_0)$.

- It remains to prove that the extension is unique: Therefore, suppose there exists another measure λ on $\sigma(\mathfrak{F}_0)$ such that $\mu(A) = \lambda(A)$ for all $A \in \mathfrak{F}_0$. Let $\lambda_n(A) = \lambda(A \cap A_n)$ for all $A \in \sigma(\mathfrak{F}_0)$. Note that we can define each λ_n directly on $\sigma(\mathfrak{F}_0)$ and not only on $\mathfrak{F}_0$ as it was the case for the measures μ_n ! Moreover, each λ_n is a finite measure on $\sigma(\mathfrak{F}_0)$, as it is bounded by $\lambda(A_n) = \mu(A_n)$, which is finite.

Our aim is to prove that λ and μ^* agree on $\sigma(\mathfrak{F}_0)$: For each A_n , consider the class $\mathcal{C}_n = \{A \in \sigma(\mathfrak{F}_0) \mid \lambda_n(A) = \mu_n^*(A)\}$, i.e. the class of all sets $A \in \sigma(\mathfrak{F}_0)$ for which λ_n and the extension of μ_n agree: First, we prove that each class $\mathcal{C}_n$ is a *monotone class*: Therefore, let $C_1, C_2, \dots \in \mathcal{C}_n$ such that $C_i \uparrow C$. Each C_i is an element of $\sigma(\mathfrak{F}_0)$ and as a σ -field, $\sigma(\mathfrak{F}_0)$ is closed under increasing sequences; hence $C \in \sigma(\mathfrak{F}_0)$. Thus, in order to show that $C \in \mathcal{C}_n$, it remains to prove that $\lambda_n(C) = \mu_n^*(C)$. Now $C_i \uparrow C$ implies that

$$\lim_{i \rightarrow \infty} \mu_n^*(C_i) = \mu_n^*(C) \quad \text{and} \quad \lim_{i \rightarrow \infty} \lambda_n(C_i) = \lambda_n(C).$$

But $\mu_n^*(C_i) = \lambda_n(C_i)$ for all $i \in \mathbb{N}$, as $C_i \in \mathcal{C}$. Thus $\lim_{i \rightarrow \infty} \mu_n^*(C_i) = \lim_{i \rightarrow \infty} \lambda_n(C_i)$. As the limits are equal, i.e. $\mu_n^*(C) = \lambda_n(C)$, we conclude that $C \in \mathcal{C}_n$.

Having established that each $\mathcal{C}_n$ is a monotone class, it is easy to see that $\mathfrak{F}_0 \subseteq \mathcal{C}_n$: From the extension, we know that $\mu_n = \mu_n^*$ on $\mathfrak{F}_0$; hence $\mu_n(A) = \mu_n^*(A) = \lambda_n(A)$ for all $A \in \mathfrak{F}_0$ and $\mathfrak{F}_0 \subseteq \mathcal{C}_n$. By Thm. 2.2, we conclude that $\sigma(\mathfrak{F}_0) \subseteq \mathcal{C}_n$ and thus, $\lambda_n(A) = \mu_n^*(A)$ for all $A \in \sigma(\mathfrak{F}_0)$. But then $\lambda(A) = \sum_{n=1}^{\infty} \lambda_n(A) = \sum_{n=1}^{\infty} \mu_n^*(A) = \mu^*(A)$. Hence $\lambda = \mu$ on $\sigma(\mathfrak{F}_0)$, proving uniqueness. $\square$

2.1.3 Approximate representations of elements in $\mathfrak{F}$

The difference between a field $\mathfrak{F}_0$ of subsets of Ω and the smallest σ -field $\sigma(\mathfrak{F}_0)$ generated by $\mathfrak{F}_0$ is that elements of $\sigma(\mathfrak{F}_0)$ may be obtained by taking countably infinite combinations of unions and intersections of elements in $\mathfrak{F}_0$. In contrast to $\sigma(\mathfrak{F}_0)$, the elements in $\mathfrak{F}_0$ are structurally simple, as they are constructed using only finitely many unions and intersections. Nevertheless, there is no bound on the number of such unions and intersections.

Intuitively, this leads to the following observation: If $\mathfrak{F}$ is the σ -field generated by a field $\mathfrak{F}_0$, and $A \in \mathfrak{F}$, we can construct a set $B \in \mathfrak{F}_0$ which approximates the set A arbitrarily closely by just taking enough unions and intersections of elements in $\mathfrak{F}_0$ when building the set B .

To make this precise, let $X, Y \subseteq \Omega$ and define the *set difference* $X \triangle Y$ of X and Y by $X \triangle Y = (X \setminus Y) \cup (Y \setminus X)$. Given a set $A \in \mathfrak{F}$, we can construct a set $B \in \mathfrak{F}_0$ by taking finitely many unions and intersections of elements in $\mathfrak{F}_0$ such that $\mu(A \triangle B) < \varepsilon$ for any predefined $\varepsilon > 0$.

Note however, that in general, the smaller ε is chosen, the more complex the unions and intersections needed for the construction of B become. The possibility of approximating elements in $\mathfrak{F}$ by those in $\mathfrak{F}_0$ is made precise in the following theorem:

Theorem 2.4 (Approximation theorem). *Let $(\Omega, \mathfrak{F}, \mu)$ be a measure space and $\mathfrak{F}_0$ be a field of subsets of Ω with $\sigma(\mathfrak{F}_0) = \mathfrak{F}$. Further, let μ be σ -finite on $\mathfrak{F}_0$. For all $\varepsilon > 0$ and $A \in \mathfrak{F}$ with $\mu(A) < +\infty$, there exists $B \in \mathfrak{F}_0$ such that $\mu(A \triangle B) < \varepsilon$.*

Proof. A proof can be found in [ADD00, Thm. 1.3.11]. □

The approximation theorem is used in Chapter 5 to construct finite representations of Borel-measurable functions.

2.2 The Borel σ -field and the Lebesgue measure

In this thesis, we consider systems that evolve in continuous-time, where time points are modeled by the set of nonnegative real numbers. The aim of this section is to construct a measure that allows us to quantify the “size” or “length” of any set of time-points, i.e. of any subset $A \subseteq \mathbb{R}_{\geq 0}$.

In the following, we apply the extension technique from Sec. 2.1 to derive a σ -field $\mathfrak{B}(\mathbb{R})$ over the set of real numbers $\mathbb{R}$. Further, we define the Lebesgue measure, which corresponds to the natural notion of “size” or “length” of subsets of $\mathbb{R}$.

2.2.1 The size of intervals

We strive to define a measure on (measurable) subsets of $\mathbb{R}_{\geq 0}$. A natural requirement is that the measure of any interval $(a, b]$ with $a, b \in \mathbb{R}_{\geq 0}$ and $a < b$ is its length, that is, we expect the measure of $(a, b]$ to be $b - a$.

Note that in the following, we use *right-semiclosed* intervals of the form $(a, b]$ to derive the Borel σ -field $\mathfrak{B}(\mathbb{R})$; however, as will become clear in the next paragraph, we also could have used any other type of interval (closed or open, or intervals of the form $(-\infty, a]$).

Definition 2.6 (Right-semiclosed interval). For $a, b \in \mathbb{R}^\infty$, the set $(a, b] = \{x \in \mathbb{R} \mid a < x \leq b\}$ is a right-semiclosed interval in $\mathbb{R}$.

Now, let μ be a set function defined on right-semiclosed intervals such that if $I = (a, b]$, then $\mu(I) = b - a$. In this way, μ formalizes the length of right-semiclosed intervals.

There is one subtle point in Def. 2.6: It states that any right-semiclosed interval on $\mathbb{R}$ is a subset of $\mathbb{R}$; as $+\infty, -\infty \notin \mathbb{R}$, we identify the set $(a, +\infty]$ with the set $\{x \in \mathbb{R} \mid a < x\}$ and define this set to be right-semiclosed. Similarly, we define $(-\infty, a] = \{x \in \mathbb{R} \mid x \leq a\}$ to be right-semiclosed. This convention is necessary, as it makes the class of right-semiclosed intervals closed under complement, which is required in Lemma 2.6.

Right-semiclosed intervals are a very restricted class of subsets of $\mathbb{R}$; for example, given a right-semiclosed interval $(a, b]$, we are not able to measure its complement $(a, b]^c = (-\infty, a] \cup (b, +\infty]$ or any other disjoint union of right-semiclosed intervals. To address this, we strive to extend the set function μ to a larger class of subsets of $\mathbb{R}$. In a first step, we therefore consider the class $\mathfrak{F}_0$ that consists of all finite disjoint unions of right-semiclosed intervals:

By definition, all elements A of $\mathfrak{F}_0$ have the form $A = (a_1, b_1] \cup (a_2, b_2] \cup \dots \cup (a_n, b_n]$ for some $n \in \mathbb{N}$ and $a_i, b_i \in \mathbb{R}^\infty$. Thus, it suffices to define $\mu(A) = \sum_{i=1}^n \mu((a_i, b_i])$ for all $A \in \mathfrak{F}_0$. Then the class $\mathfrak{F}_0$ of finite disjoint unions of right-semiclosed intervals forms a field:

Lemma 2.6. Let $\mathfrak{F}_0$ be the class of finite disjoint unions of right-semiclosed intervals in $\mathbb{R}$. Then $\mathfrak{F}_0$ is a field.

Proof. Let $\Omega = \mathbb{R}$. To show that $\mathfrak{F}_0$ is a field, we verify the conditions of Def. 2.1:

- (a) $\Omega \in \mathfrak{F}_0$ is satisfied as $\mathbb{R} = (-\infty, +\infty] \in \mathfrak{F}_0$. Note that by Def. 2.6, intervals of the form $\{x \in \mathbb{R} \mid a < x \leq +\infty\} = (a, +\infty]$ are right-semiclosed.
- (b) Let $A = \bigcup_{i=1}^n A_i$ with $A_i = (a_i, b_i]$ be a finite disjoint union of right-semiclosed intervals. Without loss of generality, we may assume that the A_i are ordered according to their lower interval bounds, i.e. let $a_i \leq a_{i+1}$ for $i = 1, 2, \dots, n-1$. First, we prove that $A \cup (a, b] \in \mathfrak{F}_0$ for any right-semiclosed interval $(a, b]$:

If $A \cap (a, b] = \emptyset$, then $A \cup (a, b] \in \mathfrak{F}_0$ and we are done. Otherwise, there exist $j, k \in \{1, \dots, n\}$, $j \leq k$ with $(a_i, b_i] \cap (a, b] \neq \emptyset$ for all $i \in \{j, j+1, \dots, k\}$ and $(a_i, b_i] \cap (a, b] = \emptyset$ for all other i . (see Fig. 2.1, where $j = 2$ and $k = 4$). To obtain a disjoint decomposition of the set $(\bigcup_{i=1}^n (a_i, b_i]) \cup (a, b]$, set $a_{\min} = \min\{a, a_j\}$ and $b_{\max} = \max\{b_k, b\}$ and replace $(\bigcup_{i=j}^k A_i) \cup (a, b] \subseteq A$ with the interval $(a_{\min}, b_{\max}]$: Therefore, define $C_i = A_i$ for $i < j$, $C_j = (a_{\min}, b_{\max}]$ and for $i > j$, define $C_i = A_{i+(k-j)}$.

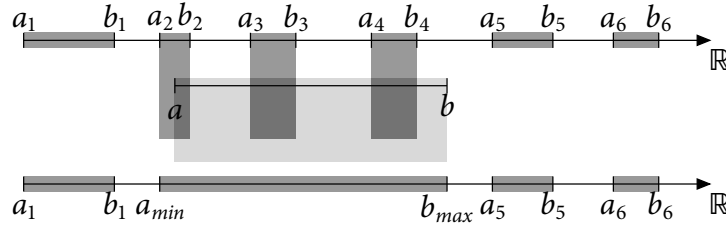


Figure 2.1: The union of an interval and a disjoint union of right-semiclosed intervals.

By construction it then follows that $C_i \cap C_j = \emptyset$ for $i \neq j$ and $(\bigcup_{i=1}^n A_i) \cup (a, b] = \bigcup_{i=1}^{n-(k-j)} C_i \in \mathfrak{F}_0$.

Now, let $A, B \in \mathfrak{F}_0$, i.e. $A = \bigcup_{i=1}^n A_i$ and $B = \bigcup_{i=1}^m B_i$ for some $n, m \in \mathbb{N}$. To complete the proof, we show that $A \cup B \in \mathfrak{F}_0$: Therefore, let $C_1 = A$ and $C_{i+1} = C_i \cup B_i$ for $i = 1, 2, \dots, m$. We prove that $C_i \in \mathfrak{F}_0$ by induction on i : By definition, $C_1 = A \in \mathfrak{F}_0$. For the induction step, let $C_i \in \mathfrak{F}_0$. By the above argument, $C_{i+1} = C_i \cup B_i \in \mathfrak{F}_0$. Hence, $C_{m+1} \in \mathfrak{F}_0$; now the claim follows, as $C_{m+1} = A \cup B$.

- (c) Let $A = \bigcup_{i=1}^n A_i \in \mathfrak{F}_0$ be defined as before and set $B_i = (b_{i-1}, a_i]$ for $1 \leq i \leq n+1$ with $b_0 = -\infty$ and $a_{n+1} = +\infty$. Then $A^c = \bigcup_{i=1}^{n+1} B_i$ and hence, $A^c \in \mathfrak{F}_0$. $\square$

With this result, we know that by extending μ from single intervals to the elements in $\mathfrak{F}_0$, we can already measure the complement and union of any finite combination of right-semiclosed intervals.

It can even be proved (cf. [MP90, p. 23] and [ADD00, Lemma 1.4.3]) that μ is countably additive on $\mathfrak{F}_0$, that is, if $A_1, A_2, \dots \in \mathfrak{F}_0$ is a countably infinite sequence of disjoint sets in $\mathfrak{F}_0$ with the property that their union $\bigcup_{i=1}^{\infty} A_i$ is again in $\mathfrak{F}_0$, then $\mu(\bigcup_{i=1}^{\infty} A_i) = \sum_{i=1}^{\infty} \mu(A_i)$. Hence, countable additivity on $\mathfrak{F}_0$ allows us to reason even about countably infinite unions of intervals, provided they do belong to $\mathfrak{F}_0$. However, such countable unions obviously are an exception, as $\mathfrak{F}_0$ is not a σ -field but just a field.

Example 2.4. As an example of a countably infinite union which is in $\mathfrak{F}_0$ and can be measured by μ without further extensions, let $A_i = (\frac{1}{2^i}, \frac{1}{2^{i-1}}]$ for $i = 1, 2, \dots$ be a countably infinite sequence of disjoint right-semiclosed intervals. Then $(\bigcup_{i=1}^{\infty} A_i) = (0, 1] \in \mathfrak{F}_0$ and therefore, $\mu(\bigcup_{i=1}^{\infty} A_i) = \mu((0, 1]) = 1$. However, this obviously does not hold in general: If $B_i = (1 - \frac{1}{2^{i-1}}, 1 - \frac{1}{2^i}]$, then $B_i \in \mathfrak{F}_0$ for all $i = 1, 2, \dots$ and $\bigcup_{i=1}^{\infty} B_i = (0, 1)$. But $(0, 1)$ is not right-semiclosed; hence, it is not in $\mathfrak{F}_0$ and therefore, not in the domain of μ . $\diamond$

As can be seen from the example, the structure of the elements in $\mathfrak{F}_0$ is too restricted. In the general case (cf. Sec. 2.1.1), the next step is to define the set function μ' (see Lemma 2.3), which extends μ to the class $\mathcal{G} = \{\bigcup_{i=1}^{\infty} A_i \mid A_i \in \mathfrak{F}_0\}$ of countable unions of elements in $\mathfrak{F}_0$. Although we do not go into the details here, note that the class $\mathcal{G}$ is still restricted; more specifically, it is not closed under complement:

Example 2.5. Reconsider the sequence of sets $B_i \in \mathfrak{F}_0$ as defined in Ex. 2.4 and let $G = (0, 1)$. If we define $G_n = \bigcup_{i=1}^n B_i$ for $n = 1, 2, \dots$, then $G_n \uparrow G$ and $G \in \mathcal{G}$. Therefore, with the extension of μ to μ' , we can measure the set $G = (0, 1)$. However, its complement $B^c = (-\infty, 0] \cup [1, +\infty]$ is still not in $\mathcal{G}$: To see this, note that by definition, the left-semiclosed interval $[1, +\infty]$ is not in $\mathfrak{F}_0$. Further, no increasing sequence $\{C_n\}_{n \in \mathbb{N}} \in \mathfrak{F}_0$ converges to a left-semiclosed interval. Hence $[1, +\infty] \notin \mathcal{G}$. $\diamond$

In order to extend μ to a larger class of subsets of $\mathbb{R}$, we now develop an extension to the smallest σ -field $\sigma(\mathfrak{F}_0)$ that is generated by $\mathfrak{F}_0$. To motivate this extension, observe that in contrast to $\mathfrak{F}_0$, the σ -field $\sigma(\mathfrak{F}_0)$ is closed under all countable unions and under complements.

2.2.2 Distribution functions and Lebesgue-Stieltjes measures

So far, μ is a measure on the field $\mathfrak{F}_0$ of finite disjoint unions of right-semiclosed intervals. Now, we apply the extension described in Sec. 2.1.1 to derive a measure on $\sigma(\mathfrak{F}_0)$:

Definition 2.7 (Borel σ -field). The Borel σ -field $\mathfrak{B}(\mathbb{R})$ is the smallest σ -field generated by the field $\mathfrak{F}_0$ of finite disjoint unions of right-semiclosed intervals, that is, $\mathfrak{B}(\mathbb{R}) = \sigma(\mathfrak{F}_0)$.

Any σ -field is closed under countable union and complement (cf. Def. 2.1). Therefore, we can imagine $\mathfrak{B}(\mathbb{R})$ also as the smallest σ -field that contains all right-semiclosed intervals. Moreover, the choice of right-semiclosed intervals for the construction of $\mathfrak{B}(\mathbb{R})$ is arbitrary. For example, $\mathfrak{B}(\mathbb{R})$ contains all closed intervals iff it contains all right-semiclosed intervals. To see this, note that

$$[a, b] = \bigcap_{n=1}^{\infty} \left(a - \frac{1}{n}, b \right] \quad \text{and} \quad (a, b] = \bigcup_{n=1}^{\infty} \left[a + \frac{1}{n}, b \right].$$

Similarly, it can be proved that $\mathfrak{B}(\mathbb{R})$ is the smallest σ -field that contains all left-semiclosed as well as all open intervals.

The extension of the measure μ from the field $\mathfrak{F}_0$ to the Borel σ -field $\mathfrak{B}(\mathbb{R}) = \sigma(\mathfrak{F}_0)$ is based on Carathéodory's extension theorem (Thm. 2.3). In the following, we generalize the idea of extending μ to $\mathfrak{B}(\mathbb{R})$ such that it also applies to cases, where the measure of an interval $(a, b]$ is not defined as the difference $b - a$:

Example 2.6 (Measure of the exponential distribution). Let $\lambda \in \mathbb{R}_{\geq 0}$ and define the function $\mu_\lambda((a, b]) = e^{-\lambda a} - e^{-\lambda b}$ for all right-semiclosed intervals $(a, b]$. As we will see later, μ_λ turns out to be the measure induced by the negative exponential distribution with rate λ . $\diamond$

To achieve greater flexibility, we do no longer define $\mu((a, b]) = b - a$ directly, but use a distribution function $F : \mathbb{R} \rightarrow \mathbb{R}$ instead, where we set $\mu((a, b]) = F(b) - F(a)$:

Definition 2.8 (Distribution function). A distribution function on $\mathbb{R}$ is a mapping $F : \mathbb{R} \rightarrow \mathbb{R}$ such that

- (a) F is increasing, i.e. $F(a) \leq F(b)$ for all $a \leq b$ and
- (b) F is right-continuous, i.e. $\lim_{x \rightarrow a^+} F(x) = F(a)$.

By the formula $\mu((a, b]) = F(b) - F(a)$, a distribution function F defines a measure μ on the Borel σ -field: For example, the distribution function $F(x) = x$ defines the measure μ that we have investigated so far, i.e. $\mu((a, b]) = b - a = F(b) - F(a)$. Further, the negative exponential distribution with rate λ is $F_\lambda(x) = 1 - e^{-\lambda x}$. Hence, the set function μ_λ in Ex. 2.6 is obtained directly by $F_\lambda(x)$.

In general, there is a one-to-one correspondence between distribution functions and the so-called class of Lebesgue-Stieltjes measures:

Definition 2.9 (Lebesgue-Stieltjes measure). A measure $\mu : \mathfrak{B}(\mathbb{R}) \rightarrow \mathbb{R}_{\geq 0}$ on $(\mathbb{R}, \mathfrak{B}(\mathbb{R}))$ is a Lebesgue-Stieltjes measure iff $\mu(I) < +\infty$ for all bounded intervals $I \subseteq \mathbb{R}$.

The class of Lebesgue-Stieltjes measures is the most prominent class of measures on the Borel σ -field. It is related to the definition of distribution functions in the following sense: Any measure that is defined by a distribution function is a Lebesgue-Stieltjes measure, and reversely, for any Lebesgue-Stieltjes measure, we can construct a corresponding distribution function:

Theorem 2.5 (Lebesgue-Stieltjes measures induce distribution functions). Let $\mu : \mathfrak{B}(\mathbb{R}) \rightarrow \mathbb{R}_{\geq 0}$ be a Lebesgue-Stieltjes measure and let $F : \mathbb{R} \rightarrow \mathbb{R}$ be such that $F(b) - F(a) = \mu((a, b])$. Then F is a distribution function.

Proof. Let $a, b \in \mathbb{R}$ and $a < b$. Then $F(b) - F(a) = \mu((a, b]) \geq 0$. This implies that $F(b) \geq F(a)$ and therefore, F is increasing. For right-continuity, let $x \in \mathbb{R}$ and let $x_1 > x_2 > x_3 > \dots$ be a decreasing sequence such that $\lim_{n \rightarrow \infty} x_n = x$. Then $F(x_n) - F(x) = \mu((x, x_n])$; further, as μ is a measure, it holds that $\lim_{n \rightarrow \infty} \mu((x, x_n]) = 0$. To see this, note that $\lim_{n \rightarrow \infty} (x, x_n] = \emptyset$, which has measure 0. This implies that $\lim_{n \rightarrow \infty} F(x_n) - F(x) = 0$ and $\lim_{n \rightarrow \infty} F(x_n) = F(x)$. Therefore, F is right-continuous. $\square$

For the proof of the reverse direction, we apply the extension results from Sec. 2.1.1:

Theorem 2.6 (Distribution functions induce Lebesgue-Stieltjes measures). *Let $F : \mathbb{R} \rightarrow \mathbb{R}$ be a distribution function and let μ be a function on right-semiclosed intervals such that $\mu((a, b]) = F(b) - F(a)$. Then μ extends uniquely to a measure on $\mathfrak{B}(\mathbb{R})$.*

Proof. As before, set $\mu((a, b]) = F(b) - F(a)$ to obtain a measure for right-semiclosed intervals. The first step in the extension is to define μ on $\mathfrak{F}_0$; therefore, let $A_1, A_2, \dots, A_n$ be disjoint right-semiclosed intervals in $\mathbb{R}$ and define $\mu(\bigcup_{i=1}^n A_i) = \sum_{i=1}^n \mu(A_i)$. This extends μ to a measure on the field $\mathfrak{F}_0$. To be able to apply Carathéodory's extension theorem that extends μ to $\sigma(\mathfrak{F}_0)$, we need to prove that μ is a σ -finite measure on the field $\mathfrak{F}_0$. First, note that μ is finitely additive on $\mathfrak{F}_0$; moreover, it can be proved that μ is also countably additive on $\mathfrak{F}_0$ (cf. [ADD00, Lemma 1.4.3]). To see that μ is σ -finite, note that $\mathbb{R} = \bigcup_{n=1}^{\infty} (-n, n]$ and that $\mu((-n, n]) = F(n) - F(-n) < +\infty$. Hence, by Carathéodory's extension theorem (Thm. 2.3), there exists a unique extension of μ to a measure on $\sigma(\mathfrak{F}_0) = \mathfrak{B}(\mathbb{R})$. $\square$

With Thm. 2.5 and Thm. 2.6, we have established a one-to-one correspondence between Lebesgue-Stieltjes measures and distribution functions. Thus, the measure μ on right-semiclosed intervals, that we defined by $\mu((a, b]) = b - a$ has a unique extension to the Borel σ -field. In fact, it is important enough to get its own name:

Definition 2.10 (Lebesgue measure). *The Lebesgue measure on $\mathfrak{B}(\mathbb{R})$ is the Lebesgue-Stieltjes measure induced by the distribution function $F(x) = x$.*

We slightly extend the definition of the Lebesgue measure: Let $\overline{\mathfrak{B}}(\mathbb{R})$ be the completion of $\mathfrak{B}(\mathbb{R})$ i.e. any element $A \in \overline{\mathfrak{B}}(\mathbb{R})$ can be expressed as a union $A = B \cup M$, where $B \in \mathfrak{B}(\mathbb{R})$ and $M \subseteq N$ is a subset of a set $N \in \mathfrak{B}(\mathbb{R})$ that has Lebesgue measure 0.

Definition 2.11 (Borel and Lebesgue measurable sets). *Let $\mathfrak{B}(\mathbb{R})$ the Borel σ -field, μ the Lebesgue measure and $\overline{\mathfrak{B}}(\mathbb{R})$ the completion of $\mathfrak{B}(\mathbb{R})$ w.r.t. μ . The elements in $\mathfrak{B}(\mathbb{R})$ are the Borel measurable sets. If $A \in \overline{\mathfrak{B}}(\mathbb{R})$, then A is a Lebesgue measurable set.*

To extend the Lebesgue measure μ to $\overline{\mathfrak{B}}(\mathbb{R})$, let $A \in \overline{\mathfrak{B}}(\mathbb{R})$. Then $A = B \cup M$, where $B, N \in \mathfrak{B}(\mathbb{R})$, $\mu(N) = 0$ and $M \subseteq N$. Therefore, we extend the Lebesgue measure μ from $\mathfrak{B}(\mathbb{R})$ to a measure on $\overline{\mathfrak{B}}(\mathbb{R})$ by setting $\mu(A) = \mu(B)$. As the difference between μ on $\mathfrak{B}(\mathbb{R})$ and $\overline{\mathfrak{B}}(\mathbb{R})$ is only w.r.t. sets of measure zero, we do not distinguish between μ and its extension to $\overline{\mathfrak{B}}(\mathbb{R})$; instead, we refer to both as the Lebesgue measure.

Another important property of the Lebesgue measure is translation invariance. It will be essential to prove the existence of sets that are not measurable.

Lemma 2.7 (The Lebesgue measure is translation invariant). *Let μ be the Lebesgue measure, $A \in \mathfrak{B}(\mathbb{R})$ and $b \in \mathbb{R}$. Then $A \oplus b \in \mathfrak{B}(\mathbb{R})$ and $\mu(A \oplus b) = \mu(A)$.*

Proof. First, let $A = \bigcup_{i=1}^n A_i \in \mathfrak{F}_0$ with pairwise disjoint right-semiclosed intervals A_i . Then $A \oplus b = \bigcup_{i=1}^n A_i \oplus b$ with each $A_i \oplus b$ being a right-semiclosed interval. Hence, $A \oplus b \in \mathfrak{F}_0$. Further, for each $A_i = (a_i, b_i]$ it holds that $\mu(A_i) = F(b_i) - F(a_i) = b_i - a_i = (b_i + b) - (a_i + b) = F(b_i + b) - F(a_i + b) = \mu(A_i \oplus b)$. Therefore $\mu(A) = \mu(\bigcup_{i=1}^n A_i) = \sum_{i=1}^n \mu(A_i) = \sum_{i=1}^n \mu(A_i \oplus b) = \mu(\bigcup_{i=1}^n (A_i \oplus b)) = \mu(A \oplus b)$, proving that the Lebesgue measure μ is translation invariant on $\mathfrak{F}_0$.

To extend this result to the Borel σ -field, we use the monotone class theorem (Thm. 2.2) and a proof technique which is also used in Thm. 4.7; in [ADD00], Ash calls it the “good sets principle”. The idea is as follows: Let

$$\mathfrak{C} = \{A \in \mathfrak{B}(\mathbb{R}) \mid A \oplus b \in \mathfrak{B}(\mathbb{R}) \wedge \mu(A \oplus b) = \mu(A)\}$$

be the class of good sets. First, we have to prove that $\mathfrak{C}$ is a monotone class:

- Let $A_1 \subseteq A_2 \subseteq \dots \in \mathfrak{C}$ be such that $A_n \uparrow A$. By definition of $\mathfrak{C}$, it follows that $A_n \oplus b \in \mathfrak{B}(\mathbb{R})$ for all $n \in \mathbb{N}$. Further, $A_1 \oplus b \subseteq A_2 \oplus b \subseteq \dots$. Hence, $A_n \oplus b \uparrow A \oplus b$. But as σ -fields are closed under increasing sequences (to see this, note that $A \oplus b = \bigcup_{n=1}^{\infty} A_n \oplus b$ and that $\mathfrak{B}(\mathbb{R})$ is closed under countable union), it follows that $A \oplus b \in \mathfrak{B}(\mathbb{R})$. Further, μ is a measure, hence $\mu(A \oplus b) = \lim_{n \rightarrow \infty} \mu(A_n \oplus b)$. By definition of $\mathfrak{C}$, $\mu(A_n \oplus b) = \mu(A_n)$. Therefore $\mu(A \oplus b) = \lim_{n \rightarrow \infty} \mu(A_n \oplus b) = \lim_{n \rightarrow \infty} \mu(A_n) = \mu(A)$. Thus $\mu(A \oplus b) = \mu(A)$ and $A \in \mathfrak{C}$.
- Let $A_1 \supseteq A_2 \supseteq \dots \in \mathfrak{C}$ such that $A_n \downarrow A$. Again, $A_n \oplus b \in \mathfrak{B}(\mathbb{R})$ and $A_n \oplus b \downarrow A \oplus b$. Further, σ -fields are closed under decreasing sequences as $A \oplus b = \bigcap_{n=1}^{\infty} (A_n \oplus b) = \left(\bigcup_{n=1}^{\infty} (A_n \oplus b)^c \right)^c$. Hence $A \oplus b \in \mathfrak{B}(\mathbb{R})$. Further, $\mu(A \oplus b) = \lim_{n \rightarrow \infty} \mu(A_n \oplus b) = \lim_{n \rightarrow \infty} \mu(A_n) = \mu(A)$. Hence, $A \in \mathfrak{C}$.

Thus, $\mathfrak{C}$ is a monotone class. Further, $\mathfrak{F}_0 \subseteq \mathfrak{C}$, as for each $A \in \mathfrak{F}_0$, it holds that $A \oplus b \in \mathfrak{F}_0$ and $\mu(A) = \mu(A \oplus b)$. By the monotone class theorem (Thm. 2.2), we conclude that $\sigma(\mathfrak{F}_0) \subseteq \mathfrak{C}$. Hence, $A \oplus b \in \mathfrak{B}(\mathbb{R})$ and $\mu(A) = \mu(A \oplus b)$ for all $A \in \mathfrak{B}(\mathbb{R})$ and $b \in \mathbb{R}$. $\square$

2.3 A set that is not Lebesgue measurable

Now that we have discussed the technical details that allow to derive Lebesgue-Stieltjes measures from distribution functions and right-semiclosed intervals, we now construct an example of a set that is not Lebesgue measurable.

Therefore, this section partly answers the question that we posed in the discussion following Thm. 2.1 in a more general setting. It turns out that $2^{\mathbb{R}} \setminus \overline{\mathfrak{B}(\mathbb{R})} \neq \emptyset$; hence,

although the extensions that we have discussed in Sec. 2.1.1 cover a very large class of subsets (namely $\overline{\mathfrak{B}}(\mathbb{R})$) of the real line, there exist sets that are not Lebesgue measurable. Even worse, there are uncountably many of them. However, the construction of these *Vitali sets* is nonconstructive and relies on the axiom of choice.

Let us start slowly with the definition of an equivalence relation:

Lemma 2.8. *Let $\mathbb{Q}$ denote the rationals and define a relation $\sim \subseteq \mathbb{R} \times \mathbb{R}$ such that*

$$\forall x, y \in \mathbb{R}. x \sim y \iff x - y \in \mathbb{Q}.$$

Then $\sim$ is an equivalence relation.

Proof. Reflexivity follows directly as $x - x = 0$ and $0 \in \mathbb{Q}$ for all $x \in \mathbb{R}$. For symmetry, let $x, y \in \mathbb{R}$ such that $x \sim y$. Then $x - y = z$ for some $z \in \mathbb{Q}$. Equivalently, $y - x = -z$. But $-z \in \mathbb{Q}$ and therefore $y \sim x$. For transitivity, let $x, y, z \in \mathbb{R}_{\geq 0}$ with $x \sim y$ and $y \sim z$. Further, let $x - y = u$ and $y - z = v$. Then $x - z = (u + y) - (y + v) = u - v$. Now $u, v \in \mathbb{Q}$; hence, $x - z = u - v \in \mathbb{Q}$ and therefore $x \sim z$. $\square$

As usual, let $[x]_{\sim} = \{y \in \mathbb{R} \mid x \sim y\}$ denote the equivalence class of $x \in \mathbb{R}$. Further, let $\mathfrak{R} = \{[x]_{\sim} \mid x \in \mathbb{R}\}$ be the set of all equivalence classes of $\sim$. Then $\mathfrak{R}$ partitions the set of real numbers, i.e. $\bigcup \mathfrak{R} = \mathbb{R}$.

Example 2.7. *Let $x \in \mathbb{Q}$. Its equivalence class $[x]_{\sim}$ is the set of all rational numbers, i.e. $[x]_{\sim} = \mathbb{Q}$ as $x - y \in \mathbb{Q}$ for all $y \in \mathbb{Q}$. As an example for an irrational number, consider the constant $\pi \in \mathbb{R}$. It holds $[\pi]_{\sim} = \{y \in \mathbb{R} \mid \exists u \in \mathbb{Q}. y = \pi + u\}$ and $[x]_{\sim} \neq [\pi]_{\sim}$.* $\diamond$

As it can be seen from the examples above, the definition of $\sim$ is not trivial; in fact, the set $\mathfrak{R}$ contains uncountably many equivalence classes, each of which consists of infinitely many elements.

For the construction of Vitali sets, we restrict to the subset of real numbers in $(0, 1]$. The idea is to pick from each equivalence class $[x]_{\sim} \in \mathfrak{R}$ exactly one representative; any set that contains a representative from each equivalence class is a Vitali set. Formally:

Definition 2.12 (Vitali set). *A Vitali set is a set $V \subseteq (0, 1]$ such that $|V \cap [x]_{\sim}| = 1$ for all $x \in \mathbb{R}$.*

Some remarks are in order: First, it turns out that there are uncountably many equivalence classes in $\mathfrak{R}$ (for a discussion, see [Kan91]). Second, each equivalence class is countably infinite: To see this, note that all elements y of any equivalence class $[x]_{\sim}$ differ in a rational number. Hence, the cardinality of $[x]_{\sim}$ is that of the rationals.

Hence, there are uncountably many possibilities to select a combination of representatives for each equivalence class so that we can construct uncountably many Vitali sets.

However, in this intuitive reasoning, we implicitly assume that it is possible to choose exactly one representative from each of the *uncountably many* equivalence classes in $\mathfrak{R}$. However, this assumption is not so clear: In fact, the existence of Vitali sets depends on the axiom of choice:

Axiom 2.1 (Axiom of choice). *Let $\mathcal{X}$ be a set. For any set $\mathfrak{X} \subseteq 2^{\mathcal{X}}$ with $\mathfrak{X} \neq \emptyset$, there exists a choice function $f : \mathfrak{X} \rightarrow \mathcal{X}$ such that $f(X) \in X$ for all $X \in \mathfrak{X}$.*

Therefore, if we set $\mathcal{X} = (0, 1]$ and $\mathfrak{X} = \{([x]_{\sim} \cap (0, 1]) \mid [x]_{\sim} \in \mathfrak{R}\}$, the axiom of choice states that we may select a representative in $([x]_{\sim} \cap (0, 1])$ for each equivalence class $[x]_{\sim} \in \mathfrak{R}$.

To prove that $V \notin \overline{\mathfrak{B}}(\mathbb{R})$, we have to investigate the Vitali sets a bit closer: Therefore, let V be a Vitali set, $v \in V$ an element of the Vitali set V and $q \in \mathbb{Q}$. Then $[v]_{\sim} = [v + q]_{\sim}$ as $(v + q) \sim v$. Moreover, if $q_1, q_2 \in \mathbb{Q}$ with $q_1 \neq q_2$ and $V \oplus q_i = \{v + q_i \in \mathbb{R} \mid v \in V\}$ for $i = 1, 2$, then $V \oplus q_1$ and $V \oplus q_2$ are both Vitali sets.

Furthermore it holds that $(V \oplus q_1) \cap (V \oplus q_2) = \emptyset$: To prove this, let $x \in (V \oplus q_1)$. Then there exists $v \in V$ such that $x = v + q_1$. Now assume that $x \in V \oplus q_2$. This implies that $x = v' + q_2$ for some $v' \in V$; further, $v \neq v'$ as $q_1 \neq q_2$. But $v' = x - q_2 = v + q_1 - q_2$ and $q_1 - q_2 \in \mathbb{Q}$; thus $v \sim v'$. Therefore $V \cap [v]_{\sim} \supseteq \{v, v'\}$, contradicting the definition of V . Hence, $x \notin V \oplus q_2$. The same argument applies for the reverse direction, i.e. $y \in V \oplus q_2$ implies $y \notin V \oplus q_1$. Hence, the two Vitali sets $V \oplus q_1$ and $V \oplus q_2$ are disjoint.

Another property used in the proof of the next theorem is that $(0, 1] \subseteq \bigcup_{q \in \mathbb{Q}} (V \oplus q)$. To establish this, fix some $x \in (0, 1]$ and consider its equivalence class $[x]_{\sim}$. By definition, there exists $v \in V$ such that $v \in [x]_{\sim}$. But then $x \sim v$ and $x = v + q$ for some $q \in \mathbb{Q}$. Hence, $x \in (V \oplus q)$ for some $q \in \mathbb{Q}$. Therefore it holds that $(0, 1] \subseteq \bigcup_{q \in \mathbb{Q}} (V \oplus q)$. We are now ready for the proof that Vitali sets are not Lebesgue measurable:

Theorem 2.7 (Vitali sets are not Lebesgue measurable). *Let $\overline{\mathfrak{B}}(\mathbb{R})$ be the Borel σ -field, completed w.r.t. the Lebesgue measure μ and let V be a Vitali set. Then $V \notin \overline{\mathfrak{B}}(\mathbb{R})$.*

Proof. Let μ be the Lebesgue measure on $\overline{\mathfrak{B}}(\mathbb{R})$ and assume that $V \in \overline{\mathfrak{B}}(\mathbb{R})$. Consider the sets $V \oplus \frac{1}{n}$ for $n \in \mathbb{N}_{>0}$. By definition, it holds that $(V \oplus \frac{1}{n}) \subseteq (0, 2]$ for all $n \in \mathbb{N}_{>0}$. Moreover, we have proved above, that the sets $V \oplus \frac{1}{n}$ and $V \oplus \frac{1}{m}$ are disjoint for $n, m \in \mathbb{N}_{>0}$ and $n \neq m$. Therefore $\bigcup_{n=1}^{\infty} (V \oplus \frac{1}{n}) \subseteq (0, 2]$. Thus

$$0 \leq \sum_{n=1}^{\infty} \mu \left(V \oplus \frac{1}{n} \right) = \mu \left(\bigcup_{n=1}^{\infty} \left(V \oplus \frac{1}{n} \right) \right) \leq \mu((0, 2]) = 2. \quad (2.3)$$

By Lemma 2.7, the Lebesgue measure μ is translation invariant. Hence $\mu(V \oplus \frac{1}{n}) = \mu(V)$ for all $n \in \mathbb{N}_{>0}$. Thus $\sum_{n=1}^{\infty} \mu(V \oplus \frac{1}{n}) = \sum_{n=1}^{\infty} \mu(V)$ and (2.3) implies $0 \leq \sum_{n=1}^{\infty} \mu(V) \leq 2$. The only solution to this inequality is $\mu(V) = 0$.

Applying Lemma 2.7 (translation invariance of μ) again, we obtain that $\mu(V \oplus c) = 0$ for all $c \in \mathbb{R}$. But as shown before, $(0, 1] \subseteq \bigcup_{q \in \mathbb{Q}} (V \oplus q)$. This implies

$$1 = \mu((0, 1]) \leq \mu\left(\bigcup_{q \in \mathbb{Q}} (V \oplus q)\right) = \sum_{q \in \mathbb{Q}} \mu(V \oplus q) = 0,$$

which is a contradiction. Hence $V \notin \overline{\mathfrak{B}}(\mathbb{R})$. $\square$

As a consequence of Thm. 2.7, we may conclude that although the extension techniques that we have developed in Sec. 2.1.1 extend a measure μ from a field $\mathfrak{F}_0$ to its generated σ -field $\sigma(\mathfrak{F})$ and moreover, to the completion of $\sigma(\mathfrak{F})$ w.r.t. μ , there generally remain uncountably many sets (like the Vitali sets in the case of the Borel σ -field), that are not measurable.

2.4 The Lebesgue integral

In order to define a path-based semantics of randomly timed systems like CTMDPs and IMCs, we need to integrate over uncountable sets of paths. Further, CTMDPs and IMCs are systems that evolve in continuous-time; hence, we need measures on the Borel σ -field to measure their behavior in the continuous-time domain.

To achieve this generality, we mostly do not use the Riemann integral, which only permits to integrate functions that map from the reals to the real numbers. Instead, we consider the more general Lebesgue integral, which accounts for Borel measurable functions that map from an arbitrary measurable space to the extended real numbers.

Although the set of Lebesgue integrable functions is a proper superset of Riemann integrable functions, we have to impose certain measurability conditions.

2.4.1 Measurable functions

To motivate the concept of measurable functions, let $(\Omega, \mathfrak{F}, \mu)$ be a measure space and let $h : \Omega \rightarrow \mathbb{R}^\infty$. Thus, the function h assigns to each element in Ω an extended real number. Now, assume that we are interested in the measure of the set of all $\omega \in \Omega$ for which $h(\omega) \in B$ for some interval $B \subseteq \mathbb{R}^\infty$. That is, we aim to compute the measure $\mu(h^{-1}(B))$ of the set $h^{-1}(B) = \{\omega \in \Omega \mid h(\omega) \in B\}$. As μ is a measure on $(\Omega, \mathfrak{F})$, it is a function $\mu : \mathfrak{F} \rightarrow \mathbb{R}_{\geq 0}^\infty$; hence, in order for $\mu(h^{-1}(B))$ to be well-defined, the set $h^{-1}(B)$ must be measurable, that is, it must hold that $h^{-1}(B) \in \mathfrak{F}$.

If we generalize this idea, we arrive at the definition of *measurable functions*:

Definition 2.13 (Measurable function). Let $(\Omega_1, \mathfrak{F}_1)$ and $(\Omega_2, \mathfrak{F}_2)$ be measurable spaces. Any function $f : \Omega_1 \rightarrow \Omega_2$ that satisfies $f^{-1}(B) \in \mathfrak{F}_1$ for all $B \in \mathfrak{F}_2$ is measurable with respect to the σ -fields $\mathfrak{F}_1$ and $\mathfrak{F}_2$.

We use the notation $f : (\Omega_1, \mathfrak{F}_1) \rightarrow (\Omega_2, \mathfrak{F}_2)$ to denote the fact that f is a measurable function with respect to the measurable spaces $(\Omega_1, \mathfrak{F}_1)$ and $(\Omega_2, \mathfrak{F}_2)$.

Measurable functions share many nice properties: For example, the composition of two measurable functions is again measurable:

Theorem 2.8 (Composition of measurable functions). *Let $f : (\Omega_1, \mathfrak{F}_1) \rightarrow (\Omega_2, \mathfrak{F}_2)$ and $g : (\Omega_2, \mathfrak{F}_2) \rightarrow (\Omega_3, \mathfrak{F}_3)$. Their composition $g \circ f$ is defined such that $(g \circ f)(\omega_1) = g(f(\omega_1))$ for all $\omega_1 \in \Omega_1$. Then, the function $g \circ f : \Omega_1 \rightarrow \Omega_3$ is measurable with respect to $\mathfrak{F}_1$ and $\mathfrak{F}_3$.*

Proof. The proof can be found in [ADD00, Lemma 1.5.7]. □

In the general setting above, we let h be defined between two measurable spaces; to link the definition to the Lebesgue integral, let $(\Omega_1, \mathfrak{F}_1)$ be some measurable space and set $(\Omega_2, \mathfrak{F}_2) = (\mathbb{R}^\infty, \mathfrak{B}(\mathbb{R}^\infty))$. Then $h : \Omega \rightarrow \mathbb{R}^\infty$ is measurable with respect to $(\Omega, \mathfrak{F})$ and $(\mathbb{R}^\infty, \mathfrak{B}(\mathbb{R}^\infty))$ iff $h^{-1}(B) \in \mathfrak{F}$ for all sets $B \in \mathfrak{B}(\mathbb{R}^\infty)$.

Definition 2.14 (Borel measurable function). *Let $(\Omega, \mathfrak{F})$ be a measurable space. A function $f : (\Omega, \mathfrak{F}) \rightarrow (\mathbb{R}^\infty, \mathfrak{B}(\mathbb{R}^\infty))$ is Borel measurable.*

In probability theory, Borel measurable functions are called random variables, i.e. a Borel measurable function $X : (\Omega, \mathfrak{F}) \rightarrow (\mathbb{R}, \mathfrak{B}(\mathbb{R}))$ is a *random variable*. Note that the Lebesgue integral also permits to integrate functions that map to $\{+\infty, -\infty\}$; however, within probability theory and also throughout this thesis, it suffices to consider the Borel σ -field $\mathfrak{B}(\mathbb{R})$ instead of the Borel σ -field $\mathfrak{B}(\mathbb{R}^\infty)$ over the extended reals.

Example 2.8 (A function that is not Borel measurable). *With the Vitali set construction from Sec. 2.3, it is straightforward to derive a function that is not Borel measurable: Let V be a Vitali set (hence, $V \notin \mathfrak{B}(\mathbb{R})$) and define $h : (\mathbb{R}, \mathfrak{B}(\mathbb{R})) \rightarrow (\mathbb{R}, \mathfrak{B}(\mathbb{R}))$ such that $h(x) = 1$ if $x \in V$ and $h(x) = 0$, otherwise. Then $h^{-1}(1) = V \notin \mathfrak{B}(\mathbb{R})$; hence, h is not Borel measurable.* ◇

Before we define the Lebesgue integral of Borel measurable functions, let us consider some properties of Borel measurable functions. As we have already seen, they are closed under composition. Moreover:

Theorem 2.9 (Pointwise limit of Borel measurable functions). *Let $(\Omega, \mathfrak{F})$ be a measurable space. If $h_1, h_2, \dots$ are Borel measurable functions such that $h_n(\omega) \rightarrow h(\omega)$ for all $\omega \in \Omega$ and $n \in \mathbb{N}$, then the function h (i.e. the pointwise limit of the h_n) is also Borel measurable.*

Proof. For a proof, see [ADD00, Thm. 1.5.4]. $\square$

Further, the class of Borel measurable functions is closed under algebraic operations:

Theorem 2.10. *Let h and h' be Borel measurable functions from $(\Omega, \mathfrak{F})$ to $(\mathbb{R}, \mathfrak{B}(\mathbb{R}))$. Provided they are well-defined, the functions $h + h'$, $h - h'$, $h \cdot h'$ and h/h' are Borel measurable.*

Proof. For a proof, see [ADD00, Thm. 1.5.6]. $\square$

2.4.2 The Lebesgue integral

With the introduction of Borel measurable functions, we are now ready to define the Lebesgue integral. Ultimately, we will define the Lebesgue integral of any Borel measurable function $h : (\Omega, \mathfrak{F}) \rightarrow (\mathbb{R}^\infty, \mathfrak{B}(\mathbb{R}^\infty))$ over some measure space $(\Omega, \mathfrak{F}, \mu)$. Therefore, we proceed stepwise; for the beginning, let us consider *simple functions*:

Definition 2.15 (Simple function). *Any Borel measurable function $h : (\Omega, \mathfrak{F}) \rightarrow (\mathbb{R}^\infty, \mathfrak{B}(\mathbb{R}^\infty))$ with a finite image is simple iff $|\{h(\omega) \mid \omega \in \Omega\}| < +\infty$.*

As a consequence, a simple function h takes on only finitely many values $x_1, x_2, \dots, x_n$, say. Hence, we can partition the domain Ω of h into finitely many disjoint sets, denoted $A_1, A_2, \dots, A_n \in \mathfrak{F}$, such that the elements in each set A_i map to the fixed value x_i . Formally, let $\{x_1, x_2, \dots, x_n\} = \{h(\omega) \mid \omega \in \Omega\}$ be the image of a simple function h and let $A_i = \{\omega \in \Omega \mid h(\omega) = x_i\}$. Then h can be written as the finite sum

$$h(\omega) = \sum_{i=1}^n x_i \cdot \mathbf{I}_{A_i}(\omega), \quad (2.4)$$

where we use the *indicator function* $\mathbf{I}$, which is defined for any subset X of a set $\mathcal{X}$ such that

$$\mathbf{I}_X : \mathcal{X} \rightarrow \{0, 1\} : x \mapsto \begin{cases} 1 & \text{if } x \in X \\ 0 & \text{otherwise.} \end{cases}$$

Hence, in Eq. (2.4), all summands with $\omega \notin A_i$ are 0, whereas for the (uniquely determined) set A_i with $\omega \in A_i$, we return the value x_i .

The idea to define the abstract Lebesgue integral of a simple function $h : (\Omega, \mathfrak{F}) \rightarrow (\mathbb{R}^\infty, \mathfrak{B}(\mathbb{R}^\infty))$ with respect to a measure space $(\Omega, \mathfrak{F}, \mu)$ is as follows: Let μ be a measure on $(\Omega, \mathfrak{F})$ and assume that as before, the sets $A_1, A_2, \dots, A_n$ partition the set Ω according

to the finitely many values $x_1, x_2, \dots, x_n$ that h takes on. Then we define the abstract Lebesgue integral of h as follows:

$$\int_{\Omega} h(\omega) \mu(d\omega) = \sum_{i=1}^n x_i \cdot \mu(A_i). \quad (2.5)$$

First, let us fix some notation: If ω is clear from the context (and μ is unary), we also use $\int_{\Omega} h d\mu$ to denote the Lebesgue integral as defined in Eq. (2.5).

According to Eq. (2.5), in order to compute $\int_{\Omega} h d\mu$, we multiply each value x_i that the simple function h can take on with the measure of its preimage under h .

Example 2.9 (Interpretation of the Lebesgue integral). *Informally, Fig. 2.2 depicts the construction of the abstract Lebesgue integral: In contrast to the Riemann integral, the Lebesgue integral computes the area under a curve by measuring each subset A_i of Ω , where the step function h takes on value x_i ; Fig. 2.2(a) depicts this partitioning of Ω according to the values that h takes on. Informally, the area that is under those segments of the graph of h , where h takes on, say value x_i , is given by the product of the measure of the segment and the height of x_i , that is, by $\mu(A_i) \cdot x_i$. Consequently we obtain the area under the curve of h (cf. Fig. 2.2(b)) by adding up the corresponding products for all values $x_1, x_2, \dots, x_n$. $\diamond$*

One further remark is in order here: The Lebesgue integral is defined w.r.t. an arbitrary measurable space $(\Omega, \mathfrak{F}, \mu)$. More concretely, notwithstanding its name, it is not limited to the Lebesgue measure or to the class of Lebesgue-Stieltjes measures!

Up to now, we have defined the abstract Lebesgue integral for simple functions only. To lift this restriction, we now strive for an extension of the defining Equation (2.5) to a larger class of functions. As a first step, consider the class of nonnegative Borel measurable functions: The idea is to approximate any nonnegative Borel measurable function h by a sequence of simple functions s that converges pointwise from below to h . Accordingly, we set

$$\int_{\Omega} h d\mu = \sup \left\{ \int_{\Omega} s d\mu \mid s \text{ is a simple function and } 0 \leq s \leq h \right\}.$$

This definition is justified by the following theorem:

Theorem 2.11 (Limit of simple functions). *Any nonnegative Borel measurable function is the limit of an increasing sequence of simple functions.*

Proof. A proof can be found in, e.g. [ADD00, Thm. 1.5.5]. □

Although within this thesis, we only need to consider the Lebesgue integral of nonnegative Borel measurable functions, the extension to arbitrary (also negative) Borel measurable functions is straightforward:

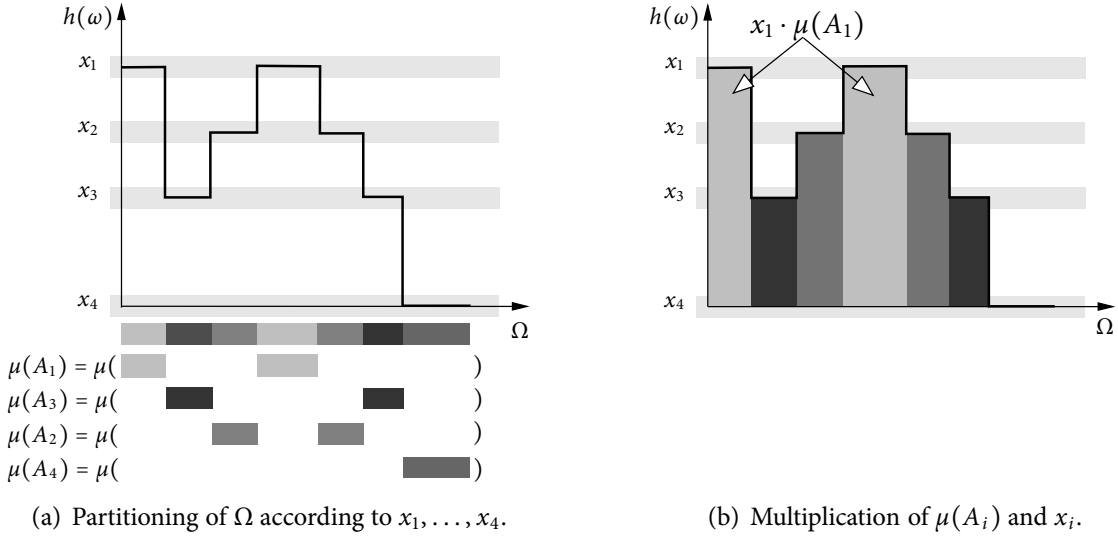


Figure 2.2: Deriving the Lebesgue integral of a simple function.

Let $h : (\Omega, \mathfrak{F}) \rightarrow (\mathbb{R}^\infty, \mathfrak{B}(\mathbb{R}^\infty))$ be an arbitrary Borel measurable function. Define the functions h^+ and h^- such that

$$h^+(\omega) = \begin{cases} h(\omega) & \text{if } h(\omega) \geq 0 \\ 0 & \text{otherwise} \end{cases} \quad h^-(\omega) = \begin{cases} -h(\omega) & \text{if } h(\omega) < 0 \\ 0 & \text{otherwise} \end{cases}$$

Obviously, this yields a decomposition of h into two nonnegative functions, i.e. $h = h^+(\omega) - h^-(\omega)$. Further, the functions h^+ and h^- are Borel measurable: To see this, we first show a more general result:

Lemma 2.9 (Maximum and minimum of Borel measurable functions). *Let $(\Omega, \mathfrak{F})$ be a measurable space and $h_1 : \Omega \rightarrow \mathbb{R}$ and $h_2 : \Omega \rightarrow \mathbb{R}$ be Borel measurable functions. Then their pointwise maximum and minimum are Borel measurable.*

Proof. We only prove the claim for the pointwise maximum, as the proof for the pointwise minimum is completely analogous. Formally, the pointwise maximum of h_1 and h_2 is the function $\max(h_1, h_2) : \Omega \rightarrow \mathbb{R} : \omega \mapsto \max\{h_1(\omega), h_2(\omega)\}$. To prove that $\max(h_1, h_2)$ is Borel measurable, it suffices to prove that $M = \{\omega \in \Omega \mid \max\{h_1(\omega), h_2(\omega)\} \leq c\} \in \mathfrak{F}$. To see this, note that the class $\{(-\infty, c] \mid c \in \mathbb{R}\}$ is a generator of $\mathfrak{B}(\mathbb{R})$. But $M = \{\omega \mid h_1(\omega) \leq c\} \cap \{\omega \mid h_2(\omega) \leq c\}$; from the fact that h_1 and h_2 are Borel measurable, we directly conclude that $\{\omega \mid h_1(\omega) \leq c\} \in \mathfrak{F}$ and $\{\omega \mid h_2(\omega) \leq c\} \in \mathfrak{F}$. As $\mathfrak{F}$ is closed under intersection, we derive that $M \in \mathfrak{F}$. $\square$

To extend the Lebesgue integral to a Borel measurable function $h = h^+(\omega) - h^-(\omega)$ given as above, note that $h^+ = \max(h, \underline{0})$ and $h^- = -\min(h, \underline{0})$, where $\underline{0}$ denotes the constant

(hence Borel measurable) function $\underline{0} : \Omega \rightarrow \mathbb{R}^\infty : \omega \mapsto 0$. With the result of Lemma 2.9, h^+ and h^- are Borel measurable. Thus, we can define the Lebesgue integral of h as the difference

$$\int_{\Omega} h \, d\mu = \int_{\Omega} h^+ \, d\mu - \int_{\Omega} h^- \, d\mu,$$

as long as the term does not have the form $(+\infty) - (+\infty)$, in which case the Lebesgue integral of h does not exist.

2.4.3 Properties of the Lebesgue integral

Even though it is much more general than the Riemann integral (cf. Sec. 2.4.4), the Lebesgue integral shares most of the properties that are commonly known from classical integration theory:

Theorem 2.12. *Let $(\Omega, \mathfrak{F}, \mu)$ be a measure space and $h : (\Omega, \mathfrak{F}) \rightarrow (\mathbb{R}^\infty, \mathfrak{B}(\mathbb{R}^\infty))$ be a Borel measurable function. The Lebesgue integral w.r.t. μ satisfies the following properties:*

- (a) *If $c \in \mathbb{R}$ is a constant and h a Borel measurable function such that $\int_{\Omega} h \, d\mu$ exists, then $\int_{\Omega} c \cdot h \, d\mu$ exists and $\int_{\Omega} c \cdot h \, d\mu = c \cdot \int_{\Omega} h \, d\mu$.*
- (b) *If h is nonnegative and $A \in \mathfrak{F}$, then*

$$\int_A h \, d\mu = \sup \left\{ \int_A s \, d\mu \mid s \text{ is a simple function and } 0 \leq s \leq h \right\}.$$

- (c) *If $\int_{\Omega} h \, d\mu$ exists, then $\int_A h \, d\mu$ exists for all $A \in \mathfrak{F}$.*

Proof. The proof can be found in [ADD00, Thm. 1.5.9]. □

Note, that with the property stated in Thm. 2.12(b) and Thm. 2.12(c), we obtain a means to compute the integral of a Borel measurable function over any set $A \in \mathfrak{F}$. Thus, we are no longer restricted to the abstract Lebesgue integral over the entire set Ω .

Theorem 2.13 (Monotone convergence). *Let $h_1 \leq h_2 \leq \dots$ be an increasing sequence of nonnegative Borel measurable functions from $(\Omega, \mathfrak{F})$ to $(\mathbb{R}^\infty, \mathfrak{B}(\mathbb{R}^\infty))$. Further, define $h(\omega) = \lim_{n \rightarrow \infty} h_n(\omega)$ for all $\omega \in \Omega$. Then $\int_{\Omega} h_n \, d\mu \rightarrow \int_{\Omega} h \, d\mu$ for $n \rightarrow \infty$.*

Proof. The proof can be found in [ADD00, Thm. 1.6.2]. □

In order to prove properties of the Lebesgue integral of a Borel measurable function, it is often useful to start with nonnegative simple functions; if we manage to prove the property for all simple functions, we know by Thm. 2.11, that any nonnegative Borel measurable function is the limit of an increasing sequence of nonnegative simple functions. Now, the monotone convergence theorem (Thm. 2.13) states that the Lebesgue integral of an increasing sequence of nonnegative simple functions converges to the integral of their limit. Thus, we have established the property on all nonnegative Borel measurable functions. What remains is the extension to arbitrary Borel measurable functions. This can often be done as in Sec. 2.4.2 by decomposing the function h in question into a positive and a negative part, i.e. $h = h^+ - h^-$. Within the thesis, we make use of this proof strategy in, for example, Lemma 4.2 on page 95.

Finally, to support the intuitive reasoning with Lebesgue integrals, we remark that they satisfy the usual additivity property:

Theorem 2.14 (Additivity). *Let h and g be Borel measurable functions on $(\Omega, \mathfrak{F})$. If $g + h$ is well defined (i.e. not of the form $+\infty - \infty$), then*

$$\int_{\Omega} (g + h) \, d\mu = \int_{\Omega} g \, d\mu + \int_{\Omega} h \, d\mu.$$

Proof. The proof can be found in [ADD00, Thm. 1.6.3]. □

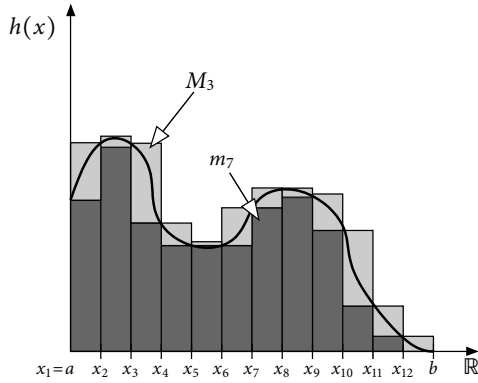
2.4.4 Comparison between the Lebesgue and Riemann integral

As we will see in this section, the Lebesgue integral is more versatile than the classical Riemann integral. More precisely, it will turn out that any Riemann integrable function is Lebesgue integrable w.r.t. the Lebesgue measure; hence, the Lebesgue integral extends the Riemann notion of integrability. Moreover, it is more versatile in the sense that it permits to integrate any Borel measurable function; moreover, the domain of integration and the corresponding measure may be given as an arbitrary measure space $(\Omega, \mathfrak{F}, \mu)$.

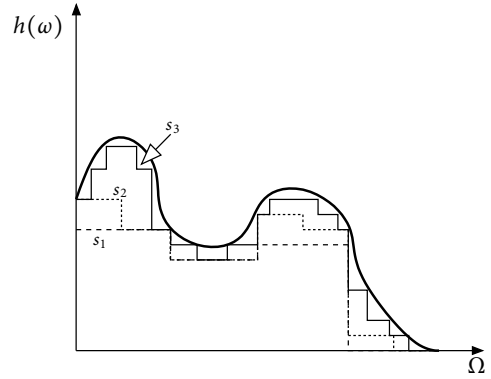
These arguments justify the use of Lebesgue integration within this thesis: We need to allow the domain of integration to be, e.g. the set of paths that describe the evolution of a system, and not just the real numbers. Further, we make heavy use of measurable functions which appear in the integral, but which are not Riemann integrable in general.

Figure 2.3(a) depicts the idea for the derivation of the Riemann integral: Let $[a, b] \subseteq \mathbb{R}$ be the domain of integration and let $x_1 < x_2 < x_3 < \dots < x_n$ with $x_1 = a$ and $x_n = b$ induce the partitioning $P = \bigcup_{i=1}^{n-1} \{[x_i, x_{i+1})\}$ of $[a, b]$. For $i = 1, 2, \dots, n-1$, the *upper* and *lower Riemann sums* are defined as

$$M_i = \sup \{h(x) \mid x \in [x_i, x_{i+1})\} \text{ and} \\ m_i = \inf \{h(x) \mid x \in [x_i, x_{i+1})\}, \text{ respectively.}$$



(a) Upper and lower Riemann sums for a partitioning $a < x_1 < \dots < b$ of the integration domain.



(b) The supremum of the limit of simple functions defines the Lebesgue integral.

Figure 2.3: The derivations of the Riemann and the Lebesgue integral.

Now, let $\alpha(x) = M_i$ and $\beta(x) = m_i$ for $x \in [x_i, x_{i+1})$. Thus, α and β are simple functions so that $U(P) = \int_a^b \alpha d\mu$ and $L(P) = \int_a^b \beta d\mu$ form the upper and lower sums given h and a partitioning P (see Fig. 2.3(a)).

Let $P_1, P_2, \dots$ be a sequence of partitions of $[a, b)$ such that P_{k+1} refines P_k and let $\alpha_1, \alpha_2, \dots$ and $\beta_1, \beta_2, \dots$ be the corresponding simple functions. Moreover, let $\|P_k\| = \max_{0 < i < |P_k|} (x_{i+1} - x_i)$ be the maximum length of the intervals in P_k . If we assume that the refinement from P_k to P_{k+1} is such that $\lim_{k \rightarrow \infty} \|P_k\| = 0$, i.e. if the length of all blocks of the refined partitions become infinitesimally small, then $\alpha_1 \geq \alpha_2 \geq \dots \geq h \geq \dots \geq \beta_2 \geq \beta_1$. Hence, $\lim_{k \rightarrow \infty} \alpha_k = \alpha$ and $\lim_{k \rightarrow \infty} \beta_k = \beta$, for the pointwise limits α and β . Further, from the definition of α and β one can derive that h is continuous (that is, $\lim_{x \rightarrow c} h(x) = h(c)$) iff $\alpha(x) = \beta(x)$ (to see this, consider the limit of a sequence of refined partitions).

With these preliminaries, the function h is *Riemann integrable* on $[a, b)$ iff

$$\int_a^b \alpha d\mu = \int_a^b \beta d\mu$$

holds independent of the partitions chosen to construct α and β . Accordingly, the Riemann integral of h on $[a, b)$ is then defined as the value $\int_a^b \alpha d\mu$ (or equivalently, $\int_a^b \beta d\mu$).

The next theorem provides another characterization of Riemann integrability. Moreover, it states that for bounded intervals, every Riemann integrable function is also Lebesgue integrable w.r.t. the Lebesgue measure:

Theorem 2.15. Let $I = [a, b)$ be an interval with $a, b \in \mathbb{R}$, $a < b$ and $h : I \rightarrow \mathbb{R}$.

(a) h is Riemann integrable on I iff h is continuous almost everywhere on I with respect to the Lebesgue measure.

(b) If h is Riemann integrable on I , then h is Lebesgue integrable on I . In this case

$$\int_I h(x) dx = \int_I h d\mu,$$

where μ is the Lebesgue measure.

Proof. The proof can be found in [ADD00, Thm. 1.7.1]. $\square$

In Thm. 2.15, the term “almost everywhere” needs some explanation: If $(\Omega, \mathfrak{F}, \mu)$ is a measurable space (in the case of Thm. 2.15, set $\Omega = \mathbb{R}$, $\mathfrak{F} = \mathfrak{B}(\mathbb{R})$ and μ the Lebesgue measure) a property holds almost everywhere on a set $A \in \mathfrak{F}$ w.r.t. measure μ iff the set $B \subseteq A$ of elements where it fails has measure 0. We denote this by stating that the property holds a.e. $[\mu]$. Note that when it comes to Lebesgue integration over a domain $A \in \mathfrak{F}$, it does not matter whether a property holds on all or almost all elements of A : In both cases, for the exceptional set B where the property is violated, it holds $\mu(B) = 0$. Hence, albeit the difference, the integrals are equal.

Now consider the converse direction: The *Dirichlet function* is an example of a function that is Lebesgue integrable w.r.t. the Lebesgue measure, but not Riemann integrable: Let $h : \mathbb{R} \rightarrow \{0, 1\}$ be such that

$$h(x) = \begin{cases} 1 & \text{if } x \in \mathbb{Q} \\ 0 & \text{otherwise.} \end{cases}$$

Then h is a nonnegative simple function and hence, it is Lebesgue integrable. Moreover, if $B \in \mathfrak{B}(\mathbb{R})$, we have that $\int_B h d\mu = \mu(B \cap \mathbb{Q})$, where μ is the Lebesgue measure. Further, the rationals are a measurable set in $\mathfrak{B}(\mathbb{R})$ and their Lebesgue measure is 0, i.e. $\mu(\mathbb{Q}) = 0$. Thus $\int_B h d\mu = 0$ for all $B \in \mathfrak{B}(\mathbb{R})$.

Further, h is not Riemann integrable, as it is discontinuous on any interval $[a, b]$ with $a < b$. To see this, note that for each block $I = [x_i, x_{i+1})$ of a partition P of $[a, b]$ and all $x \in I$ it holds that $\alpha(x) = 1$ as $I \cap \mathbb{Q} \neq \emptyset$; further, $\beta(x) = 0$ as $I \setminus \mathbb{Q} \neq \emptyset$.

Hence $\alpha_k \neq \beta_k$, no matter how fine the partition P is chosen. Thus, the upper and lower sums of the Riemann integral do not converge to the same limit.

2.5 Product σ -fields

The scope of this thesis is on finite-state systems. Their behavior is fully described by the path (or trajectory) along which they evolve. Among other things, on such a path we record which states have been visited and how long the system sojourned in each of those states. For example, the latter information (i.e. the sojourn time) is obtained as the outcome of a random experiment with a continuous probability distribution. As we

might expect, it turns out that any single such path has probability zero; therefore, we need measure theoretic arguments to measure *sets of paths*.

Now, a path may lead through arbitrarily many states, each of which involves random experiments. Hence, from a measure theoretic perspective, a path describes one outcome of a compound experiment which is composed of multiple stages.

This section is divided into four parts: First, we discuss in Sec. 2.5.1 how to construct multi-dimensional measurable spaces that capture events in compound random experiments.

With these results, the next natural question is: How can we measure these higher-dimensional events? We approach the technicalities slowly in Sec. 2.5.2, where we only describe the construction of 2-dimensional measures. After that, we extend these results in Sec. 2.5.3 and define measures on higher-dimensional product spaces.

Finally, in our systems, we also need to consider infinite paths. To formalize them, Sec. 2.5.4 introduces the cylinder set construction and the necessary tools to extend finite-dimensional product measures to the infinite case.

2.5.1 The construction of finite-dimensional product spaces

To obtain finite-dimensional product spaces, the starting point are measurable rectangles: Assume that two measurable spaces $(\Omega_1, \mathfrak{F}_1)$ and $(\Omega_2, \mathfrak{F}_2)$ are given. A natural first step to describe their product is to consider Cartesian products of the form $A_1 \times A_2$, where A_1 and A_2 are elements of the σ -fields $\mathfrak{F}_1$ and $\mathfrak{F}_2$, respectively.

As long as only finitely many measurable spaces are involved in this construction, these Cartesian products are (finite) measurable rectangles:

Definition 2.16 (Measurable rectangle). For $i = 1, 2, \dots, n$, let $(\Omega_i, \mathfrak{F}_i)$ be measurable spaces. A Cartesian product $\times_{i=1}^n A_i = A_1 \times A_2 \times \dots \times A_n$ with $A_i \in \mathfrak{F}_i$ for $i = 1, 2, \dots, n$ is a measurable rectangle. We use

$$\bigotimes_{i=1}^n \mathfrak{F}_i = \mathfrak{F}_1 \otimes \mathfrak{F}_2 \otimes \dots \otimes \mathfrak{F}_n = \{A_1 \times A_2 \times \dots \times A_n \mid A_i \in \mathfrak{F}_i\}$$

to denote the set of all measurable rectangles over the measurable spaces $(\Omega_i, \mathfrak{F}_i)$.

So far, the class of measurable rectangles is severely restricted: For example, it is neither closed under complement nor under any (finite or countably infinite) union.

Hence, we strive for an extension of the measurable rectangles to obtain a class of subsets of the entire space $\Omega = \Omega_1 \times \Omega_2 \times \dots \times \Omega_n$ that is closed under complement and countable union. Hence, we consider the smallest σ -field generated by the measurable rectangles:

Definition 2.17 (Product σ -field). For $i = 1, 2, \dots, n$, let $(\Omega_i, \mathfrak{F}_i)$ be measurable spaces. The product σ -field is the smallest σ -field generated by the measurable rectangles. It is denoted $\sigma(\otimes_{i=1}^n \mathfrak{F}_i)$.

Similar to the construction of the Borel σ -field $\mathfrak{B}(\mathbb{R})$, which is obtained as the smallest σ -field generated by the field of finite disjoint unions of right-semiclosed intervals (cf. Sec. 2.2), we can identify the field of finite disjoint unions of measurable rectangles as a generator of $\sigma(\otimes_{i=1}^n \mathfrak{F}_i)$. Hence, all results obtained in Sec. 2.1, and most notably, Carathéodory's extension theorem and the monotone class theorem carry over to the finite-dimensional case. In the next lemma, we prove that the class of finite disjoint unions of measurable rectangles is indeed a field; moreover, the lemma states that it generates the smallest σ -field over the measurable rectangles:

Lemma 2.10. Let $(\Omega_i, \mathfrak{F}_i)$ for $i = 1, 2, \dots, n$ be measurable spaces and define $\mathfrak{U}$ as the collection of finite disjoint unions of measurable rectangles in $\otimes_{i=1}^n \mathfrak{F}_i$. Then $\mathfrak{U}$ is a field and $\sigma(\otimes_{i=1}^n \mathfrak{F}_i) = \sigma(\mathfrak{U})$.

Proof. To prove that $\mathfrak{U}$ is a field, it is useful to first establish that $\mathfrak{U}$ is closed under finite intersection: Therefore, note that the set $\otimes_{i=1}^n \mathfrak{F}_i$ is already closed under finite intersection, for if $A, B \in \otimes_{i=1}^n \mathfrak{F}_i$, then $A = \times_{i=1}^n A_i$ and $B = \times_{i=1}^n B_i$ and $A \cap B = \times_{i=1}^n (A_i \cap B_i)$. As each $\mathfrak{F}_i$ is a σ -field, it follows that $(A_i \cap B_i) \in \mathfrak{F}_i$. Hence, $A \cap B \in \otimes_{i=1}^n \mathfrak{F}_i$.

With this observation we are ready to show that also $\mathfrak{U}$ is closed under finite intersection: Let $A, B \in \mathfrak{U}$. Then $A = \bigcup_{k=1}^m A_k$ and $B = \bigcup_{k'=1}^{m'} B_{k'}$, where $A_k = \times_{i=1}^n A_{k,i}$ and $B_{k'} = \times_{i=1}^n B_{k',i}$ with $A_{k,i}, B_{k',i} \in \mathfrak{F}_i$. Then

$$A \cap B = \left(\bigcup_{k=1}^m \bigtimes_{i=1}^n A_{k,i} \right) \cap \left(\bigcup_{k'=1}^{m'} \bigtimes_{i=1}^n B_{k',i} \right) = \bigcup_{k=1}^m \bigcup_{k'=1}^{m'} \bigtimes_{i=1}^n (A_{k,i} \cap B_{k',i}) = \bigcup_{k=1}^m \bigcup_{k'=1}^{m'} A_k \cap B_{k'}. \quad (2.6)$$

As shown before, $A_k, B_{k'} \in \otimes_{i=1}^n \mathfrak{F}_i$ implies that $A_k \cap B_{k'} \in \otimes_{i=1}^n \mathfrak{F}_i$. Hence, the disjoint union of the sets $(A_k \cap B_{k'})$ is also in $\mathfrak{U}$; but then Eq. (2.6) implies that $A \cap B \in \mathfrak{U}$.

Now we come to the proof that $\mathfrak{U}$ is a field. Therefore we verify properties (a), (b) and (c) of Def. 2.1:

- (a) For all $i \in \{1, 2, \dots, n\}$ it holds that $\Omega_i \in \mathfrak{F}_i$. Thus, the set $\Omega = \times_{i=1}^n \Omega_i$ is in $\mathfrak{U}$.
- (b) To show that $\mathfrak{U}$ is closed under complement, let $A \in \mathfrak{U}$. Then there exists $m \in \mathbb{N}$ such that $A = \bigcup_{k=1}^m A_k$ and $A_k \in \otimes_{i=1}^n \mathfrak{F}_i$ for all $k \in \{1, 2, \dots, m\}$. Hence, each set A_k is a Cartesian product of the form $A_k = A_{k,1} \times A_{k,2} \times \dots \times A_{k,n}$, where $A_{k,i} \in \mathfrak{F}_i$ for all $i = 1, 2, \dots, n$. Moreover, the complement A_k^c of each A_k has the form

$$A_k^c = (A_{k,1}^c \times A_{k,2} \times \dots \times A_{k,n}) \cup (A_{k,1} \times A_{k,2}^c \times A_{k,3} \times \dots \times A_{k,n})$$

$$\begin{aligned}
& \cup \dots \cup (A_{k,1} \times A_{k,2} \times A_{k,3} \times \dots \times A_{k,n}^c) \\
& \cup (A_{k,1}^c \times A_{k,2}^c \times A_{k,3} \times \dots \times A_{k,n}) \cup (A_{k,1}^c \times A_{k,2} \times A_{k,3}^c \times A_{k,4} \times \dots \times A_{k,n}) \\
& \cup \dots \cup (A_{k,1}^c \times A_{k,2}^c \times A_{k,3}^c \times \dots \times A_{k,n}^c).
\end{aligned}$$

Now, $A_{k,i} \in \mathfrak{F}_i$ implies that $A_{k,i}^c \in \mathfrak{F}_i$. Hence, each A_k^c is a finite disjoint union of measurable rectangles. Thus $A_k^c \in \mathfrak{U}$ for all $k = 1, 2, \dots, m$.

With these preliminaries, the proof that $A^c \in \mathfrak{U}$ is easy: By de Morgan's law, we have $A^c = (\cup_{k=1}^m A_k)^c = \cap_{k=1}^m A_k^c$. But as shown in the beginning, $\mathfrak{U}$ is closed under finite intersection. As each A_k^c is an element of $\mathfrak{U}$, we conclude that $A^c \in \mathfrak{U}$.

- (c) To prove that $\mathfrak{U}$ is closed under finite union, let $A, B \in \mathfrak{U}$. Then there exist constants $m, m' \in \mathbb{N}$ such that $A = \cup_{k=1}^m A_k$ and $B = \cup_{k'=1}^{m'} B_{k'}$ for sets $A_k, B_{k'} \in \otimes_{i=1}^n \mathfrak{F}_i$.

Now, let $C = A \cup B$. Then $C = \cup_{k=1}^{m+m'} C_k$, where for $k = 1, 2, \dots, m+m'$, we let $C_k = A_k$ for $k \leq m$ and $C_k = B_{k-m}$, otherwise. Then each C_k is a measurable rectangle.

By de Morgan's law, $C = (\cap_{k=1}^{m+m'} C_k^c)^c$ and by part (b), it holds that $C_k^c \in \mathfrak{U}$. As shown in the beginning, $\mathfrak{U}$ is closed under finite intersection, hence $\cap_{k=1}^{m+m'} C_k^c \in \mathfrak{U}$. Now $C \in \mathfrak{U}$ follows again by part (b).

Thus $\mathfrak{U}$ is a field. The fact that $\otimes_{i=1}^n \mathfrak{F}_i \subseteq \mathfrak{U}$ directly implies that $\sigma(\otimes_{i=1}^n \mathfrak{F}_i) \subseteq \sigma(\mathfrak{U})$. For the reverse direction, note that $\mathfrak{U} \subseteq \sigma(\otimes_{i=1}^n \mathfrak{F}_i)$. Hence, $\sigma(\mathfrak{U}) \subseteq \sigma(\otimes_{i=1}^n \mathfrak{F}_i)$. $\square$

Another property of product σ -fields that is used in this thesis is the following: If $(\Omega_i, \mathfrak{F}_i)$ is a measurable space for $i = 1, 2, \dots, n$, then $\sigma(\sigma(\otimes_{i=1}^{n-1} \mathfrak{F}_i) \otimes \mathfrak{F}_n) = \sigma(\otimes_{i=1}^n \mathfrak{F}_i)$. Hence, we may “append” a σ -field $\mathfrak{F}_n$ to a product σ -field $\sigma(\otimes_{i=1}^{n-1} \mathfrak{F}_i)$ by constructing the set of 2-dimensional measurable rectangles, where the first component is an element in $\otimes_{i=1}^{n-1} \mathfrak{F}_i$ and the second component is in $\mathfrak{F}_n$. Then, the smallest σ -field they generate coincides with the n -dimensional σ -field which we would have obtained if we began the construction right from the beginning with the class of n -dimensional measurable rectangles.

2.5.2 Measures on two-dimensional product spaces

We start the definition of measures on product spaces with the simple case of a two-dimensional σ -field. Therefore, let $(\Omega_1, \mathfrak{F}_1, \mu_1)$ and $(\Omega_2, \mathfrak{F}_2, \mu_2)$ be measure spaces and let $A \in \sigma(\mathfrak{F}_1 \otimes \mathfrak{F}_2)$ be an element of the product σ -field.

In probability theory, the event A corresponds to a set of outcomes $(\omega_1, \omega_2) \in A$ of a two-stage experiment, where ω_1 is the outcome of the first and ω_2 the outcome of the second experiment. If the two experiments are independent (that is, the outcome ω_1 of the first experiment does not alter the probability distribution of the second experiment) and $A = A_1 \times A_2$ is a measurable rectangle, we expect the measure of A to be the product of the measures of A_1 and A_2 , that is, $\mu(A) = \mu_1(A_1) \cdot \mu_2(A_2)$.

This idea transfers to arbitrary elements $A \in \sigma(\mathfrak{F}_1 \otimes \mathfrak{F}_2)$ (i.e. elements in $\sigma(\mathfrak{F}_1 \otimes \mathfrak{F}_2)$ that are not Cartesian products):

Example 2.10. Consider the two-stage random experiment where two fair dice are thrown. Then $\Omega_1 \times \Omega_2 = \{1, \dots, 6\}^2$ serves as the sample space and $\mathfrak{F}_i = 2^{\Omega_i}$ (for $i = 1, 2$) form the corresponding σ -fields.

An event such as $A = \{(\omega_1, \omega_2) \mid \omega_1 + \omega_2 \geq 6\}$ is not a measurable rectangle; for example, $(1, 5), (5, 1) \in A$ but $(1, 1) \notin A$. Thus, depending on the outcome ω_1 of the first experiment, we are interested in different events in $\mathfrak{F}_2$: If $\omega_1 = 1$, we are only interested in outcomes ω_2 of the second experiment that are in $A_2 = \{5, 6\}$; if $\omega_1 = 2$, we only measure those ω_2 which are in $A_2 = \{4, 5, 6\}$, etc. $\diamond$

If $A \in \sigma(\mathfrak{F}_1 \otimes \mathfrak{F}_2)$, we call the function $A(\omega_1) = \{\omega_2 \mid (\omega_1, \omega_2) \in A\}$ the *section* of A at ω_1 . Intuitively, for any given outcome ω_1 of the first experiment, the section $A(\omega_1)$ is the set of those outcomes ω_2 of the second stage, that make the “product outcome” (ω_1, ω_2) admissible with respect to the event A .

Hence, the general idea to obtain a measure μ on $\sigma(\mathfrak{F}_1 \otimes \mathfrak{F}_2)$ from *independent* measures μ_1 and μ_2 is to multiply the measure $\mu(d\omega_1)$ of any possible outcome of the first experiment with the measure of the section $A(\omega_1)$ of admissible outcomes $\omega_2 \in A(\omega_1)$. Formally, we obtain

$$\mu(A) = \int_{\Omega_1} \mu_2(A(\omega_1)) \mu_1(d\omega_1). \quad (2.7)$$

To motivate the next step, we come back to the special case of probability measures: Up to now, we have assumed that the outcome ω_1 of the first experiment does not alter the way we measure the events that occur in the second experiment. Now we drop this assumption and consider compound experiments where the probability measure for the second stage depends on the outcome of the first stage’s random experiment.

Formally, instead of a single measure μ_2 on $(\Omega_2, \mathfrak{F}_2)$ (as in (2.7)), we now assume that for each $\omega_1 \in \Omega_1$, we are given a separate measure $\mu_2(\omega_1, \cdot) : \mathfrak{F}_2 \rightarrow \mathbb{R}_{\geq 0}^\infty$ on $(\Omega_2, \mathfrak{F}_2)$. In this setting, we obtain a measure of the event $A \in \sigma(\mathfrak{F}_1 \otimes \mathfrak{F}_2)$ by multiplying the measure $\mu(d\omega_1)$ with the measure of the intersection $A(\omega_1)$. Formally, we obtain

$$\mu(A) = \int_{\Omega_1} \mu_2(\omega_1, A(\omega_1)) \mu_1(d\omega_1). \quad (2.8)$$

Note that for each $A_2 \in \mathfrak{F}_2$, the measures in $\{\mu_2(\omega_1, \cdot) \mid \omega_1 \in \Omega_1\}$ induce a function on Ω_1 , namely $\mu_2(\cdot, A_2) : \Omega_1 \rightarrow \mathbb{R}_{\geq 0}^\infty : \omega_1 \mapsto \mu_2(\omega_1, A_2)$. Further, note that for the integral in Eq. (2.8) to be well-defined, the function $\omega_1 \mapsto \mu_2(\omega_1, A(\omega_1))$ must be Borel measurable. It can be proved, that this is the case if we require the functions $\mu_2(\cdot, A_2) : \Omega_1 \rightarrow \mathbb{R}_{\geq 0}^\infty$ to be Borel measurable for all $A_2 \in \mathfrak{F}_2$, i.e. if $\mu_2(\cdot, A_2)$ is Borel measurable for all $A_2 \in \mathfrak{F}_2$, then $\mu_2(\omega_1, A(\omega_1))$ is Borel measurable w.r.t. $(\Omega_1, \mathfrak{F}_1)$ for all $A \in \sigma(\mathfrak{F}_1 \otimes \mathfrak{F}_2)$.

The construction of a 2-dimensional measure on $\sigma(\mathfrak{F}_1 \otimes \mathfrak{F}_2)$ from the measure μ_1 and the measures $\mu_2(\omega_1, \cdot)$ is described formally in the next theorem:

Theorem 2.16 (Two-dimensional product measure theorem). *Let $(\Omega_1, \mathfrak{F}_1, \mu_1)$ be a measure space with σ -finite measure μ_1 and let $(\Omega_2, \mathfrak{F}_2)$ be a measurable space. Moreover, for each $\omega_1 \in \Omega_1$, let $\mu_2(\omega_1, \cdot) : \mathfrak{F}_2 \rightarrow \mathbb{R}_{\geq 0}^\infty$ be a measure on $(\Omega_2, \mathfrak{F}_2)$ such that for each $A_2 \in \mathfrak{F}_2$, the induced function $\mu_2(\cdot, A_2) : \Omega_1 \rightarrow \mathbb{R}_{\geq 0}^\infty : \omega_1 \mapsto \mu_2(\omega_1, A_2)$ is Borel measurable.*

If the $\mu_2(\omega_1, \cdot)$ are uniformly σ -finite¹, then

$$\mu : \sigma(\mathfrak{F}_1 \otimes \mathfrak{F}_2) \rightarrow \mathbb{R}_{\geq 0}^\infty : A \mapsto \int_{\Omega_1} \mu_2(\omega_1, A(\omega_1)) \mu_1(d\omega_1)$$

is the unique measure such that for measurable rectangles $(A_1 \times A_2) \in \mathfrak{F}_1 \otimes \mathfrak{F}_2$ it holds

$$\mu(A_1 \times A_2) = \int_{A_1} \mu_2(\omega_1, A_2) \mu_1(d\omega_1).$$

The measure μ is σ -finite; it is a probability measure if μ_1 and each of the $\mu_2(\omega_1, \cdot)$ are probability measures.

Proof. The proof can be found in [ADD00, Thm. 2.6.2]. □

As we mostly consider probability measures (which are uniformly σ -finite), we can conclude from the above theorem, that given a probability space $(\Omega_1, \mathfrak{F}_1, \mu_1)$ and a family of probability measures $\{\mu_2(\omega_1, \cdot)\}_{\omega_1 \in \Omega_1}$ on the measurable space $(\Omega_2, \mathfrak{F}_2)$ such that the induced functions $\mu(\cdot, A_2)$ are Borel measurable for all $A_2 \in \mathfrak{F}_2$, the integral

$$\mu(A) = \int_{\Omega_1} \mu_2(\omega_1, A(\omega_1)) \mu_1(d\omega_1)$$

is the uniquely determined probability measure on the product σ -field $\sigma(\mathfrak{F}_1 \otimes \mathfrak{F}_2)$.

With Thm. 2.16, we have a means to construct measures on two-dimensional product σ -fields. Moreover, it enables us to define the Lebesgue integral in two dimensions. Therefore, let $\Omega = \Omega_1 \times \Omega_2$ and $\mathfrak{F} = \sigma(\mathfrak{F}_1 \otimes \mathfrak{F}_2)$ as before. Analogous to the derivation of the (one-dimensional) Lebesgue integral, if we are to derive the (abstract) Lebesgue integral of a two-dimensional function $h : \Omega \rightarrow \mathbb{R}^\infty$, we have to assume that h is Borel measurable, i.e.

$$h : (\Omega, \mathfrak{F}) \rightarrow (\mathbb{R}^\infty, \mathfrak{B}(\mathbb{R}^\infty)).$$

Then

$$\int_{\Omega} h(\omega_1, \omega_2) \mu(d\omega_1, d\omega_2) \tag{2.9}$$

¹The class of measures $\mu_2(\omega_1, \cdot)$ is uniformly σ -finite iff there exist $B_1, B_2, \dots \in \mathfrak{F}_2$ and corresponding $k_1, k_2, \dots$ such that $\Omega_2 = \bigcup_{i=1}^\infty B_i$ and $\mu(\omega_1, B_i) \leq k_i$ for all ω_1 .

is the abstract Lebesgue integral of h . However, Eq. (2.9) does not give a clue of how to compute with such integrals. More precisely, from an algebraic point of view, it is useful to express the integral $\int_{\Omega} h(\omega_1, \omega_2) \mu(d\omega_1, d\omega_2)$ as an iterated integral (in this case, as two integrals w.r.t. Ω_1 and Ω_2 , respectively). Furthermore, for algebraic reasoning it is often useful if we can exchange the order of integration. In the remainder of this thesis, we make heavy use of both techniques; this is backed up by Fubini's theorem, which is stated next:

Theorem 2.17 (Fubini's theorem). *Let $(\Omega_1, \mathfrak{F}_1)$ and $(\Omega_2, \mathfrak{F}_2)$ be measurable spaces, $\Omega = \Omega_1 \times \Omega_2$, $\mathfrak{F} = \sigma(\mathfrak{F}_1 \otimes \mathfrak{F}_2)$. Further, let $\mu_1 : \mathfrak{F}_1 \rightarrow \mathbb{R}_{\geq 0}^{\infty}$ be a σ -finite measure and let $\mu_2(\omega_1, \cdot) : \mathfrak{F}_2 \rightarrow \mathbb{R}_{\geq 0}^{\infty}$ be uniformly σ -finite measures for each $\omega_1 \in \Omega_1$ such that $\mu_2(\cdot, A) : (\Omega_1, \mathfrak{F}_1) \rightarrow (\mathbb{R}_{\geq 0}^{\infty}, \mathfrak{B}(\mathbb{R}_{\geq 0}^{\infty}))$ for all $A \in \mathfrak{F}_2$.*

If $h : (\Omega, \mathfrak{F}) \rightarrow (\mathbb{R}^{\infty}, \mathfrak{B}(\mathbb{R}^{\infty}))$ is a Borel measurable function such that the integral $\int_{\Omega} h(\omega_1, \omega_2) \mu(d\omega_1, d\omega_2)$ exists, then the integrals $\int_{\Omega_2} h(\omega_1, \omega_2) \mu_2(\omega_1, d\omega_2)$ also exist for almost all $\omega_1 \in \Omega_1$ $[\mu_1]$. Further, the function $\bar{h}(\omega_1) = \int_{\Omega_2} h(\omega_1, \omega_2) \mu_2(\omega_1, d\omega_2)$ is Borel measurable and

$$\int_{\Omega} h(\omega_1, \omega_2) \mu(d\omega_1, d\omega_2) = \int_{\Omega_1} \underbrace{\int_{\Omega_2} h(\omega_1, \omega_2) \mu_2(\omega_1, d\omega_2)}_{\bar{h}(\omega_1)} \mu_1(d\omega_1).$$

Proof. For a proof, we refer the reader to [ADD00, Thm. 2.6.4]. $\square$

In the case that the measures $\mu_2(\omega_1, \cdot)$ do not depend on ω_1 , i.e. if the measures $\mu_2(\omega_1, \cdot)$ in Thm. 2.17 are all equal and independent of ω_1 , we use $\mu = \mu_1 \times \mu_2$ to denote their product measure on $\sigma(\mathfrak{F}_1 \otimes \mathfrak{F}_2)$. In the special case that $\mu = \mu_1 \times \mu_2$, Fubini's theorem permits to change the order of integration:

Corollary 2.1 (Changing the order in iterated integration). *Let $(\Omega_1, \mathfrak{F}_1, \mu_1)$ and $(\Omega_2, \mathfrak{F}_2, \mu_2)$ be measure spaces with σ -finite measures μ_1 and μ_2 . Further, let $\Omega = \Omega_1 \times \Omega_2$, $\mathfrak{F} = \sigma(\mathfrak{F}_1 \otimes \mathfrak{F}_2)$ and $\mu = \mu_1 \times \mu_2$. If $h : (\Omega, \mathfrak{F}) \rightarrow (\mathbb{R}^{\infty}, \mathfrak{B}(\mathbb{R}^{\infty}))$ is a Borel measurable function such that $\int_{\Omega} h(\omega_1, \omega_2) \mu(d\omega_1, d\omega_2)$ exists, then*

$$\int_{\Omega} h(\omega_1, \omega_2) \mu(d\omega_1, d\omega_2) = \int_{\Omega_1} \int_{\Omega_2} h(\omega_1, \omega_2) \mu_2(d\omega_2) \mu_1(d\omega_1) \quad (2.10)$$

$$= \int_{\Omega_2} \int_{\Omega_1} h(\omega_1, \omega_2) \mu_1(d\omega_1) \mu_2(d\omega_2). \quad (2.11)$$

Proof. To see that Eq. (2.10) holds, apply Thm. 2.17 with $\mu_2(\omega_1, \cdot) = \mu_2$ for all $\omega_1 \in \Omega_1$. Then Eq. (2.11) follows by symmetry. $\square$

2.5.3 Measures on finite-dimensional product spaces

In the previous section, we have introduced the construction of a measure for two-dimensional product spaces and established the theory necessary to extend Lebesgue integration to those product spaces. With these preliminaries, the next step is rather straightforward; in fact, we only generalize the ideas of the two-dimensional case to all finite-dimensional product spaces:

Theorem 2.18 (Finite-dimensional product measure and Fubini's theorem). *Let $(\Omega_i, \mathfrak{F}_i)$, $i = 1, 2, \dots, n$ be measurable spaces, $\Omega = \times_{i=1}^n \Omega_i$, $\mathfrak{F} = \sigma(\otimes_{i=1}^n \mathfrak{F}_i)$ and let $\mu_1 : \Omega_1 \rightarrow \mathbb{R}_{\geq 0}^\infty$ be a σ -finite measure. For each $j \in \{2, 3, \dots, n-1\}$ and tuple $(\omega_1, \omega_2, \dots, \omega_j) \in \times_{i=1}^j \Omega_i$, let $\mu_j(\omega_1, \omega_2, \dots, \omega_j, \cdot) : \mathfrak{F}_{j+1} \rightarrow \mathbb{R}_{\geq 0}^\infty$ be a measure. Moreover, assume that the induced function*

$$\mu_j(\cdot, \cdot, \dots, \cdot, C) : \Omega_1 \times \Omega_2 \times \dots \times \Omega_j \times \mathfrak{F}_{j+1} \rightarrow \mathbb{R}_{\geq 0}^\infty$$

is Borel measurable for each $C \in \mathfrak{F}_{j+1}$, i.e. $\{(\omega_1, \omega_2, \dots, \omega_j) \mid \mu(\omega_1, \omega_2, \dots, \omega_j, C) \in B\} \in \sigma(\otimes_{i=1}^j \mathfrak{F}_i)$ for all $B \in \mathfrak{B}(\mathbb{R}_{\geq 0}^\infty)$. Let $\Omega = \times_{i=1}^n \Omega_i$ and $\mathfrak{F} = \sigma(\otimes_{i=1}^n \mathfrak{F}_i)$.

(a) Product measure theorem:

There exists a unique measure μ on $\mathfrak{F}$ such that for each measurable rectangle $A_1 \times A_2 \times \dots \times A_n \in \otimes_{i=1}^n \mathfrak{F}_i$ it holds:

$$\begin{aligned} \mu(A_1 \times A_2 \times \dots \times A_n) &= \int_{A_1} \mu_1(d\omega_1) \int_{A_2} \mu_2(\omega_1, d\omega_2) \\ &\quad \dots \int_{A_{n-1}} \mu_{n-1}(\omega_1, \omega_2, \dots, \omega_{n-2}, d\omega_{n-1}) \int_{A_n} \mu_n(\omega_1, \omega_2, \dots, \omega_{n-1}, d\omega_n). \end{aligned}$$

(b) Fubini's theorem:

If $h : (\Omega, \mathfrak{F}) \rightarrow (\mathbb{R}^\infty, \mathfrak{B}(\mathbb{R}^\infty))$ is a Borel measurable function such that the integral $\int_\Omega h(\omega_1, \omega_2, \dots, \omega_n) \mu(d\omega_1, d\omega_2, \dots, d\omega_n)$ exists, then

$$\begin{aligned} \int_\Omega h \, d\mu &= \int_{\Omega_1} \mu_1(d\omega_1) \int_{\Omega_2} \mu_2(\omega_1, d\omega_2) \dots \int_{\Omega_{n-1}} \mu_{n-1}(\omega_1, \omega_2, \dots, \omega_{n-2}, d\omega_{n-1}) \\ &\quad \int_{\Omega_n} h(\omega_1, \omega_2, \dots, \omega_n) \mu_n(\omega_1, \omega_2, \dots, \omega_{n-1}, d\omega_n) \end{aligned}$$

and each of the integrals w.r.t. $\mu_j(\omega_1, \omega_2, \dots, \omega_j, \cdot)$, $j = 1, 2, \dots, n$ exists for almost all $(\omega_1, \omega_2, \dots, \omega_j)$ and is a Borel measurable function $(\times_{i=1}^j \Omega_i, \otimes_{i=1}^j \mathfrak{F}_i) \rightarrow (\mathbb{R}^\infty, \mathfrak{B}(\mathbb{R}^\infty))$.

Proof. The proof can be found in, e.g. [ADD00, Thm. 2.6.7]. □

So far, the product measure theorem (Thm. 2.18(a)) only states how to define the measure μ on measurable rectangles in $\otimes_{i=1}^n \mathfrak{F}_i$. To obtain the measure of an arbitrary element in their product σ -field $\mathfrak{F} = \sigma(\otimes_{i=1}^n \mathfrak{F}_i)$, we use Fubini's theorem: Obviously, if $A \in \mathfrak{F}$, then the indicator function $\mathbf{I}_A$ is Borel measurable with respect to $\mathfrak{F}$. If we set $h = \mathbf{I}_A$ in Thm. 2.18(b), we obtain an explicit representation of the measure of A in form of the iterated integral

$$\mu(A) = \int_{\Omega_1} \mu_1(d\omega_1) \int_{\Omega_2} \mu_2(\omega_1, d\omega_2) \cdots \int_{\Omega_{n-1}} \mu_{n-1}(\omega_1, \omega_2, \dots, \omega_{n-2}, d\omega_{n-1}) \\ \int_{\Omega_n} h(\omega_1, \omega_2, \dots, \omega_n) \mu_n(\omega_1, \omega_2, \dots, \omega_{n-1}, d\omega_n).$$

Similar to the two-dimensional case and Corollary 2.1, Fubini's theorem also permits to change the order of integration in the n -dimensional case provided that the measures $\mu_j(\omega_1, \omega_2, \dots, \omega_{j-1}, \cdot)$ do not depend on the values of $\omega_1, \omega_2, \dots, \omega_{j-1}$, i.e. if for each $j = 1, 2, \dots, n$ there exists a measure μ_j on $\mathfrak{F}_j$ such that $\mu_j(\omega_j) = \mu_j(\omega_1, \omega_2, \dots, \omega_{j-1}, \cdot)$ for all $(\omega_1, \dots, \omega_{j-1}) \in \times_{i=1}^{j-1} \Omega_i$. If this is the case, we denote by $\mu = \mu_1 \times \mu_2 \times \cdots \times \mu_n$ the product measure on $\mathfrak{F}$.

Now, let $\mu = \mu_1 \times \mu_2 \times \cdots \times \mu_n$ be such a measure and let $h : (\Omega, \mathfrak{F}) \rightarrow (\mathbb{R}^\infty, \mathfrak{B}(\mathbb{R}^\infty))$ be a Borel measurable function such that the integral $\int_\Omega h d\mu$ exists. Then

$$\int_\Omega h d\mu = \int_{\Omega_{i_1}} \mu_{i_1}(d\omega_{i_1}) \int_{\Omega_{i_2}} \mu_{i_2}(d\omega_{i_2}) \cdots \int_{\Omega_{i_n}} h(\omega_{i_1}, \omega_{i_2}, \dots, \omega_{i_n}) \mu_{i_n}(d\omega_{i_n})$$

for any permutation $(i_1, i_2, \dots, i_n)$ of $\{1, 2, \dots, n\}$.

2.5.4 Infinite product spaces

In this section, we extend the construction of the product σ -field and the corresponding product measure theorem to countably infinite products and the corresponding measures. A fundamental tool in this construction are the so-called *cylinder sets*:

Definition 2.18 (Cylinder set). For $i = 1, 2, \dots$, let $(\Omega_i, \mathfrak{F}_i)$ be measurable spaces. Then $\Omega = \times_{i=1}^\infty \Omega_i$ is the set of all tuples of the form $(\omega_1, \omega_2, \dots)$, where $\omega_i \in \Omega_i$ for all $i \in \mathbb{N}$. For a set $B^n \subseteq \times_{i=1}^n \Omega_i$, let

$$B_n = \left\{ (\omega_1, \omega_2, \dots) \in \times_{i=1}^\infty \Omega_i \mid (\omega_1, \omega_2, \dots, \omega_n) \in B^n \right\}$$

be the cylinder set induced by the base B^n . The cylinder set B_n is measurable iff $B^n \in \sigma(\otimes_{i=1}^n \mathfrak{F}_i)$. If $B^n \in \times_{i=1}^n \Omega_i$ is a Cartesian product, its induced cylinder B_n is an infinite rectangle; moreover, if $B^n \in \otimes_{i=1}^n \mathfrak{F}_i$, then B_n is an infinite measurable rectangle. We use $\otimes_{i=1}^\infty \mathfrak{F}_i$ and $\sigma(\otimes_{i=1}^\infty \mathfrak{F}_i)$ to denote the class of infinite measurable rectangles, respectively the smallest σ -field they generate.

Informally, a cylinder set B_n with base $B^n \subseteq \times_{i=1}^n \Omega_i$ can be constructed as follows: Take each finite sequence $(\omega_1, \omega_2, \dots, \omega_n)$ in the base B^n and extend it in all possible ways, that is, by all infinite extensions $(\omega_{n+1}, \omega_{n+2}, \dots) \in \times_{i=n+1}^\infty \Omega_i$. The result is a set of infinite sequences $\{(\omega_1, \omega_2, \dots, \omega_n, \omega_{n+1}, \omega_{n+2}, \dots) \in \times_{i=1}^\infty \Omega_i \mid (\omega_{n+1}, \omega_{n+2}, \dots) \in \times_{i=n+1}^\infty \Omega_i\}$ that have $(\omega_1, \omega_2, \dots, \omega_n)$ as a common prefix. Then the cylinder set B_n is the union of all those extensions of n -tuples in B^n . Figure 2.4 depicts the cylinder set construction. To ease notation, we sometimes also use $\text{Cyl}(B^n)$ to denote the cylinder B_n induced by the base B^n .

Cylinder sets are the building block of infinite product spaces. Therefore, let us consider their properties in more detail: For the moment, let $\mathfrak{W} = \{B_n \mid B^n \in \sigma(\otimes_{i=1}^n \mathfrak{F}_i)\}$ be the class of measurable cylinders.

Then $\mathfrak{W}$ is a field: Obviously, $\Omega_1 \in \mathfrak{F}_1$ is a measurable base of Ω ; hence $\Omega \in \mathfrak{W}$. For the closure under complement, note that the complement of a measurable cylinder is induced by the complement of its base, which is measurable. Moreover, any measurable cylinder $A_m \in \mathfrak{W}$ with a finite base A^m can also be represented by a longer base, i.e. by a base of the form $A^n = A^m \times \times_{i=m+1}^n \Omega_i$, where $n > m$. To see that $\mathfrak{W}$ is closed under finite union, let $A_m, B_n \in \mathfrak{W}$ be measurable cylinders with bases A^m and B^n and assume w.l.o.g. that $m < n$. Then $(A^m \times \times_{i=m+1}^n \Omega_i) \cup B^n \in \sigma(\otimes_{i=1}^n \mathfrak{F}_i)$ is a measurable base of $A_m \cup B_n$. Therefore $A_m \cup B_n \in \mathfrak{W}$. With this extension in mind, it can be proved along the same lines as for the proof of Lemma 2.10, that the class of finite disjoint unions of (infinite) measurable rectangles forms a field.

In order to relate measurable cylinders and measurable rectangles, we note that both classes generate the same σ -field. Hence, the infinite product σ -field $\sigma(\mathfrak{W})$ and the σ -field $\sigma(\otimes_{i=1}^\infty \mathfrak{F}_i)$ generated by the class of infinite measurable rectangles are the same:

Lemma 2.11. *For $i = 1, 2, \dots$, let $(\Omega_i, \mathfrak{F}_i)$ be measurable spaces. Then*

$$\sigma\left(\bigotimes_{i=1}^\infty \mathfrak{F}_i\right) = \sigma(\mathfrak{W}). \quad (2.12)$$

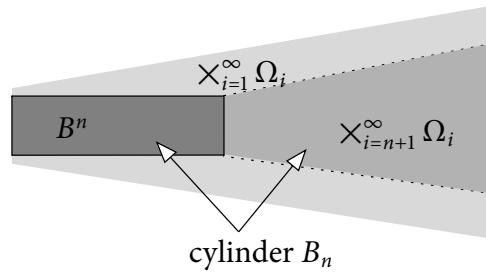


Figure 2.4: Construction from a cylinder set from a finite-dimensional base.

Proof. Let n range over $\mathbb{N}$. Further, recall that by definition it holds that $\sigma(\otimes_{i=1}^{\infty} \mathfrak{F}_i) = \sigma(\{B_n \mid B^n \in \otimes_{i=1}^n \mathfrak{F}_i\})$ and $\sigma(\mathfrak{W}) = \sigma(\{B_n \mid B^n \in \sigma(\otimes_{i=1}^n \mathfrak{F}_i)\})$. We split the proof of Eq. (2.12) in two parts:

- $\sigma(\otimes_{i=1}^{\infty} \mathfrak{F}_i) \subseteq \sigma(\mathfrak{W})$: This follows directly, as the fact that $\otimes_{i=1}^n \mathfrak{F}_i \subseteq \sigma(\otimes_{i=1}^n \mathfrak{F}_i)$ implies $\sigma(\{B_n \mid B^n \in \otimes_{i=1}^n \mathfrak{F}_i\}) \subseteq \sigma(\{B_n \mid B^n \in \sigma(\otimes_{i=1}^n \mathfrak{F}_i)\})$.
- $\sigma(\mathfrak{W}) \subseteq \sigma(\otimes_{i=1}^{\infty} \mathfrak{F}_i)$: To establish this direction, we have to prove the inclusion $\sigma(\{B_n \mid B^n \in \sigma(\otimes_{i=1}^n \mathfrak{F}_i)\}) \subseteq \sigma(\{B_n \mid B^n \in \otimes_{i=1}^n \mathfrak{F}_i\})$. Hence, it suffices to show that all measurable cylinders B_n are in $\sigma(\otimes_{i=1}^{\infty} \mathfrak{F}_i)$. Therefore, let $B_n \in \mathfrak{W}$ be a measurable cylinder with a measurable base $B^n \in \sigma(\otimes_{i=1}^n \mathfrak{F}_i)$ of length n .

With the good sets principle, let

$$\mathfrak{C} = \left\{ B^n \in \sigma\left(\bigotimes_{i=1}^n \mathfrak{F}_i\right) \mid B_n \in \sigma\left(\bigotimes_{i=1}^{\infty} \mathfrak{F}_i\right) \right\}$$

be the class of measurable bases which induce cylinders in $\sigma(\otimes_{i=1}^{\infty} \mathfrak{F}_i)$. Note that by definition, both $\sigma(\otimes_{i=1}^n \mathfrak{F}_i)$ and $\sigma(\otimes_{i=1}^{\infty} \mathfrak{F}_i)$ are σ -fields; hence, $\mathfrak{C}$ is a monotone class. Further, the field $\mathfrak{U}$ of finite disjoint unions of measurable rectangles of dimension n (cf. Def. 2.10 on page 43) is contained in $\mathfrak{C}$. Then $\sigma(\mathfrak{U}) \subseteq \mathfrak{C}$ by the monotone class theorem. By Lemma 2.10 it holds that $\sigma(\mathfrak{U}) = \sigma(\otimes_{i=1}^n \mathfrak{F}_i)$; hence $\sigma(\otimes_{i=1}^n \mathfrak{F}_i) \subseteq \mathfrak{C}$.

Therefore, let $B_n \in \mathfrak{W}$ be a measurable cylinder with base $B^n \in \sigma(\otimes_{i=1}^n \mathfrak{F}_i)$. Then $B^n \in \mathfrak{C}$. Thus, $B_n \in \sigma(\otimes_{i=1}^{\infty} \mathfrak{F}_i)$, proving the claim. $\square$

As a consequence of Lemma 2.11, from now on we use $\sigma(\otimes_{i=1}^{\infty} \mathfrak{F}_i)$ instead of $\sigma(\mathfrak{W})$ to denote the smallest σ -field generated by the class of measurable cylinders.

With these definitions, we are ready to derive the product measure theorem for the infinite case. As within this thesis we only need to consider probability measures on infinite product spaces, we restrict the exposition to the case of probability spaces² and do not consider the case of arbitrary measure spaces.

Theorem 2.19 (Ionescu-Tulcea extension theorem). *Let $(\Omega_n, \mathfrak{F}_n)$ (with $n = 1, 2, \dots$) be measurable spaces and $\mathfrak{F} = \sigma(\otimes_{i=1}^{\infty} \mathfrak{F}_i)$. Further, let P_1 be a probability measure on $\mathfrak{F}_1$ and for each $n = 2, 3, \dots$ and for all $(\omega_1, \omega_2, \dots, \omega_{n-1}) \in \times_{i=1}^{n-1} \Omega_i$, let $P_n(\omega_1, \omega_2, \dots, \omega_{n-1}, \cdot)$ be a probability measure on $\mathfrak{F}_n$. Assume that for each $n = 1, 2, \dots$ and $A \in \mathfrak{F}_n$, the induced function $P_n(\cdot, A) : \times_{i=1}^{n-1} \Omega_i \rightarrow [0, 1] : (\omega_1, \omega_2, \dots, \omega_{n-1}) \mapsto P_n(\omega_1, \omega_2, \dots, \omega_{n-1}, A)$ is measurable w.r.t. $\sigma(\otimes_{i=1}^{n-1} \mathfrak{F}_i)$.*

²In fact, the Ionescu-Tulcea construction that is used in the proof of Thm. 2.19 does not apply to arbitrary measures in a straightforward way.

For $n = 1, 2, \dots$, let $P'_n : \sigma(\otimes_{i=1}^n \mathfrak{F}_i) \rightarrow [0, 1]$ be the induced probability measure on the n -dimensional product σ -fields (cf. Thm. 2.18), that is

$$P'_n(B^n) = \int_{\Omega_1} P_1(d\omega_1) \int_{\Omega_1} P_2(\omega_1, d\omega_2) \cdots \int_{\Omega_{n-1}} P_{n-1}(\omega_1, \omega_2, \dots, \omega_{n-2}, d\omega_{n-1}) \\ \int_{\Omega_n} \mathbf{I}_{B^n}(\omega_1, \omega_2, \dots, \omega_n) P_n(\omega_1, \omega_2, \dots, \omega_{n-1}, d\omega_n)$$

for all $B^n \in \sigma(\otimes_{i=1}^n \mathfrak{F}_i)$. There exists a unique probability measure P on $\mathfrak{F}$ such that for all $n \in \mathbb{N}$, P agrees with P'_n on n -dimensional measurable cylinders, that is, for all measurable bases $B^n \in \sigma(\otimes_{i=1}^n \mathfrak{F}_i)$ it holds

$$P'_n(B^n) = P(B_n),$$

where $B_n = \{(\omega_1, \omega_2, \dots) \in \Omega \mid (\omega_1, \omega_2, \dots, \omega_n) \in B^n\}$ is the cylinder induced by B^n .

Proof. We do not go into the details here, but refer the reader to [ADD00, Thm. 2.7.2]. $\square$

The intuition for the infinite product measure theorem is simple: In fact, it states that if we have measures for all measurable bases, we obtain a unique measure on measurable cylinders if we define the probability $P(B_n)$ of a cylinder $B_n \in \sigma(\otimes_{i=1}^n \mathfrak{F}_i)$ as the probability $P'_n(B^n)$ of its base $B^n \in \sigma(\mathfrak{F}_{i=1}^n \mathfrak{F}_i)$.

The fact, that simply setting $P(B_n) = P'_n(B^n)$ yields a well-defined probability measure is not so clear. To see this, recall that any given cylinder B_m has several bases of different lengths. For example, if $n > m$, any extension of B^m of the form $B^n = B^m \times \times_{i=m+1}^n \Omega_i$ induces the cylinder B_m . However, a consequence of the uniqueness property in Thm. 2.19 is that the probabilities of those alternative bases all coincide with $P'_n(B^n)$. Therefore, it does not matter which base we choose to define the measure of B_n .

2.6 Concluding remarks

In this section, we gave a survey of the measure theoretic foundations that are needed throughout this thesis. As will turn out in the next section, we need the cylinder set construction and infinite product σ -fields to define infinite trajectories of continuous-time Markov decision processes and interactive Markov chains. Moreover, the Lebesgue integral and Fubini's theorem allow to define their semantics precisely.

Most of the material presented here is based on chapters one, two and four of the excellent book “Probability & measure theory” by Robert Ash and Catherine Doléans-Dade [ADD00]. Other references include the book “Real Analysis and Probability” by Richard Dudley [Dud02] and “Probability and Measure” by Patrick Billingsley [Bil95]. Some remarks about the difficulties in extending the notion of length to a class of subsets

of the reals that is larger than the Lebesgue measurable sets can be found in the German book “Stochastik für Informatiker” by Rudolf Mathar [MP90] and in [Ben76].

The proofs of Lemma 2.6 and Lemma 2.10, as well as the proof of Lemma 2.11 are mostly omitted in the literature. Hence, they have been proved anew and adapted to Ash’s notation, which is used throughout this thesis.

Finally, some enlightening details about the Vitali set construction, especially about the cardinality of the sets involved, can be found in the English translation of Kanovei’s paper [Kan91]. The remaining material presented in Sec. 2.3 is mostly based on a lecture note [vRS92] from Radboud University, Nijmegen.

3 An overview of stochastic models

Proof is the idol before whom
the pure mathematician
tortures himself.

(Sir Arthur Eddington)

In this thesis, we discuss a variety of probabilistic and stochastic models that describe the system behavior either in discrete or in continuous time. Therefore, this chapter introduces the basic models that we will use throughout the thesis. For each model, we try to convey its informal behavior before we formally define its semantics.

As our models evolve in time, their behaviors are described as the outcomes of compound random experiments which can be formalized in an infinite-dimensional product space, where each dimension corresponds to a fixed time-point. We refer to Sec. 2.5.4 for the probability theoretic construction of such spaces. The underlying mathematical tool that allows us to reason about these models formally, is called a *stochastic process*.

Accordingly, this overview chapter starts by shortly introducing the concepts of discrete and continuous stochastic processes. Then we discuss the special cases of discrete- and continuous-time Markov chains in more detail, as their properties are essential for the class of models that we are confronted with. Most of the material that we provide here is based on the standard textbook [Kul95], which provides an excellent introduction to Markov processes.

In the second part of this chapter, we introduce nondeterminism in Markov chains and thereby obtain discrete- and continuous-time Markov *decision processes*, where the latter are at the core of our studies in the forthcoming chapters. Discrete-time Markov decision processes are discussed in the textbook [Put94] in great detail. Moreover, [Put94] contains an introduction to continuous-time Markov decision processes in Chapter 11.

3.1 Stochastic processes

As we aim at an algorithmic verification mechanism for nondeterministic and stochastic systems, we are mostly interested in the subclass of stochastic processes that have a finite state space, as they can be stored in finite memory. Within the scope of this thesis, we therefore restrict to systems that have a finite state space. In this setting, a stochastic process is defined as follows:

Definition 3.1 (Stochastic process). A stochastic process on a finite state space $\mathcal{S}$ is a collection $\{X_t\}_{t \in T}$ of random variables X_t , where the parameter t ranges over a parameter set T . Each X_t takes on values that are in the finite state space $\mathcal{S}$.

Usually, the parameter t is interpreted as time; accordingly, for $t \in T$, the value of X_t is the state that is occupied by the stochastic process at time t . In case of a *discrete stochastic process*, the parameter set T is a subset of $\mathbb{N}$ (finite or countably infinite), whereas for *continuous stochastic processes*, the set T is a connected subset of $\mathbb{R}_{\geq 0}$. To ease notation, we use the natural numbers to refer to the discrete time parameters and the nonnegative reals for the continuous time domain.

To describe one possible evolution of a stochastic process, let $\pi : T \rightarrow \mathcal{S}$ be a function such that $\pi(t) \in \mathcal{S}$ describes the state that the stochastic process occupies at time t . Each such function describes a trajectory of the underlying stochastic process; in mathematics, each $\pi : T \rightarrow \mathcal{S}$ is called a *sample path* of the stochastic process.

Now, a stochastic process $\{X_t\}_{t \in T}$ evolves randomly along one of its sample paths. Therefore, a sample path can be seen as one possible outcome of the compound random experiment that is associated with the entire stochastic process. To link this view of a stochastic processes to the measure theoretic results of the previous chapter, we define the *sample space* of a stochastic process as the collection of all its sample paths, i.e. we set $\Omega = \{\pi : T \rightarrow \mathcal{S}\} = \mathcal{S}^T$. Accordingly, each random variable X_t is a measurable function $X_t : (\Omega, \mathfrak{F}) \rightarrow (\mathcal{S}, 2^{\mathcal{S}})$, where $\mathfrak{F}$ denotes the σ -field generated by the measurable cylinders¹. Hence, given a sample path $\pi : T \rightarrow \mathcal{S}$, the random variable X_t maps π to the state that is occupied on π at time t , i.e. $X_t(\pi) = \pi(t)$.

Now, let P be a probability measure on the measurable space $(\Omega, \mathfrak{F})$. If we are interested in the probability that at time $t \in T$, the stochastic process is in state $s \in \mathcal{S}$, we have to compute the probability measure of the set of all sample paths that are in state s at time t . Formally, this probability can be denoted as follows:

$$\begin{aligned} P(\{X_t = s\}) &= P(X_t^{-1}(s)) \\ &= P(\{\pi : T \rightarrow \mathcal{S} \mid X_t(\pi) = s\}) \\ &= P(\{\pi : T \rightarrow \mathcal{S} \mid \pi(t) = s\}). \end{aligned}$$

3.2 Markov chains

Markov chains are a prominent subclass of stochastic processes; they are particularly popular, as their analysis is relatively easy, while their expressiveness suffices to describe a broad variety of stochastic phenomena that change randomly over time. As instances of

¹In Chapter 2 we only defined countably infinite cylinders. This concept can be generalized to uncountable cylinders (see [ADD00]); however, it is not needed in the context of the thesis.

stochastic processes, Markov chains can either be discrete or continuous, depending on their underlying notion of time.

The *Markov property* distinguishes Markov chains from other stochastic processes: Informally, it states that the behavior of a Markov chain in a state $s \in \mathcal{S}$ at time $t \in T$ is independent of the states that have been visited before. Hence, it only depends on the current state s and the global time t . This locality makes Markov chains especially attractive for an analysis. The formal definition of a Markov chain with a finite state space can be stated as follows:

Definition 3.2 (Markov chain). A stochastic process $\{X_t\}_{t \in T}$ with a finite state space $\mathcal{S}$ and a parameter set T is a Markov chain iff for all $n \in \mathbb{N}$ and for all decreasing sequences of time instances $t_{n+1} > t_n > t_{n-1} > \dots > t_1 > t_0 \in T$ and states $s_i \in \mathcal{S}$, it holds that

$$\begin{aligned} P(\{X_{t_{n+1}} = s_{n+1}\} \mid X_{t_n} = s_n, X_{t_{n-1}} = s_{n-1}, X_{t_{n-2}} = s_{n-2}, \dots, X_{t_0} = s_0) \\ = P(\{X_{t_{n+1}} = s_{n+1}\} \mid X_{t_n} = s_n). \end{aligned}$$

Definition 3.2 formalizes the Markov property: Intuitively, it states that if the current time is t_n and $n+1$ steps of a Markov chain have been observed at time points $t_0 < t_1 < t_2 < \dots < t_n$, the probability to be in state s_{n+1} at time $t_{n+1} > t_n$ does only depend on the state s_n that is occupied at the current time t_n and *not* on the states $s_0, s_1, \dots, s_{n-1}$, that have been occupied before at times $t_0, t_1, \dots, t_{n-1}$.

Note that the Markov property does not state that being in state s at time t implies that the probability to be in state s_{n+1} at a later time $t' = t + \delta \in T$ is the same for all times $t \in T$.

Hence, in general, the probability to move from a state s_n within δ time units to state s_{n+1} may vary depending on the time t at which we are in state s . Stated differently, the future behavior of a Markov chain may depend on the current time t .

However, within this thesis, we assume the Markov chains to be invariant to time shifts. Such Markov chains are called *time-homogeneous*:

Definition 3.3 (Time-homogeneous Markov chain). A Markov chain $\{X_t\}_{t \in T}$ is time-homogeneous iff for all states $s, s' \in \mathcal{S}$ and for all times $t' > t \in T$ it holds that

$$P(\{X_{t'} = s'\} \mid X_t = s) = P(\{X_{t'-t} = s'\} \mid X_0 = s).$$

In the following, we restrict to time-homogeneous Markov chains and discuss their discrete- and continuous-time variants. By definition, both share the appealing property that for a given current state s , the future evolution of the Markov chain is completely determined by the state s alone. In particular, it does neither depend on the states that

have been visited in the past (Markov property), nor does it depend on the amount of time that has passed (time-homogeneity).

In the next section, we start the discussion with the conceptually simple model of discrete-time Markov chains.

3.2.1 Discrete-time Markov chains

The elements of the parameter set T of a *discrete-time Markov chain* (DTMC) are interpreted as discrete-time steps. Accordingly, the set T is usually identified with the natural numbers.

The values of the random variables X_n of a DTMC $\{X_n\}_{n \in \mathbb{N}}$ are understood as the state that the DTMC occupies after n time steps have passed. As before, the Markov property states that the probability to move from the state $X_n = s_n$ to a state $X_{n+1} = s_{n+1}$ is independent of the trajectory that led into state s_n . Moreover, we assume any DTMC to be time-homogeneous. In the discrete-time setting, this implies that

$$P(\{X_{n+1} = s'\} \mid X_n = s) = P(\{X_{m+1} = s'\} \mid X_m = s) \quad (3.1)$$

for all discrete time points $m, n \in \mathbb{N}$. Hence, the probability to move from state s to state s' does neither depend on the state sequence that has been traversed before, nor does it depend on the number of time steps that have passed.

Let $\{X_n\}_{n \in \mathbb{N}}$ be a DTMC and define

$$p_{s,s'} = P(\{X_1 = s'\} \mid X_0 = s). \quad (3.2)$$

Then $p_{s,s'}$ is the probability to move from state s to state s' , independent of the number of steps or the trajectory taken so far. Taking the $p_{s,s'}$ together, they form the *one-step transition probability matrix* $\mathbf{P}$, where $\mathbf{P} \in [0, 1]^{\mathcal{S} \times \mathcal{S}}$ is defined by $\mathbf{P}(s, s') = p_{s,s'}$. Note that there are no deadlock states in the definition of a Markov chain; therefore $\sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') = 1$ holds for all states $s \in \mathcal{S}$.

Definition 3.4 (Stochastic matrix). A matrix $\mathbf{P} \in [0, 1]^{\mathcal{S} \times \mathcal{S}}$ is stochastic iff for all $s \in \mathcal{S}$ it holds $\sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') = 1$.

From the definition, it comes as no surprise that the one-step transition probability matrix $\mathbf{P}$ of a DTMC is a stochastic matrix, i.e. the probabilities to move from a state s to some successor state $s' \in \mathcal{S}$ sum up to one:

Lemma 3.1. Let $\{X_n\}_{n \in \mathbb{N}}$ be a DTMC. Its one-step transition probability matrix $\mathbf{P}$ is a stochastic matrix.

Proof. Let $s \in \mathcal{S}$. Then it holds

$$\begin{aligned} \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') &= \sum_{s' \in \mathcal{S}} p_{s, s'} = \sum_{s' \in \mathcal{S}} P(\{X_1 = s'\} \mid X_0 = s) \\ &= P(\{X_1 \in \mathcal{S}\} \mid X_0 = s) = \frac{P(\{X_1 \in \mathcal{S} \wedge X_0 = s\})}{P(\{X_0 = s\})} = \frac{P(\{X_0 = s\})}{P(\{X_0 = s\})} = 1. \quad \square \end{aligned}$$

We are nearly done in completely describing a DTMC: The only missing item is an *initial distribution* which specifies the probability to start in a certain state s . We use $\nu \in \text{Distr}(\mathcal{S})$ to denote an initial distribution and interpret $\nu(s)$ as the probability to start in state $s \in \mathcal{S}$.

Now recall, that the random variable X_0 describes the state in which the DTMC starts. Hence, ν specifies the probability distribution associated with the random variable X_0 . As we will see, the initial distribution and the matrix $\mathbf{P}$ uniquely determine a DTMC, which is characterized by the probabilities

$$P(\{X_n = s\}) = P(X_n^{-1}(s)) = P(\{\pi : \mathbb{N} \rightarrow \mathcal{S} \mid \pi(n) = s\}).$$

According to our previous remark, an initial distribution ν serves as the probability distribution of the random variable X_0 , i.e. $P(\{X_0 = s\}) = \nu(s)$. Moreover, by the Markov property and time-homogeneity, each $p_{s, s'}$ is equal to the conditional probability $P(\{X_{n+1} = s'\} \mid X_n = s)$. As we have the probability distribution for X_0 fixed by ν , we can use the conditional probability $P(\{X_1 = s'\} \mid X_0 = s)$ to obtain the probability distribution for X_1 , that is

$$P(\{X_1 = s'\}) = \sum_{s \in \mathcal{S}} P(\{X_0 = s\}) \cdot P(\{X_1 = s'\} \mid X_0 = s).$$

In the same way, we obtain the probability $P(\{X_2 = s'\}) = \sum_{s \in \mathcal{S}} P(\{X_1 = s\}) \cdot P(\{X_2 = s'\} \mid X_1 = s)$ from the probability $P(\{X_1 = s\})$. Obviously, this inductive idea extends to all X_n . Formally, we obtain the probability distribution of X_n by the matrix vector multiplication

$$P \circ X_n^{-1} = \vec{\nu} \cdot \mathbf{P}^n, \quad (3.3)$$

where $\vec{\nu} = (\nu(s_0), \nu(s_1), \dots, \nu(s_n))$. Equation (3.3) formalizes the *transient behavior* of a DTMC. Having the one-step transition probability matrix $\mathbf{P}$ and the initial distribution ν , one can compute the probability distribution for each random variable X_n of the associated DTMC. We conclude that a DTMC is completely described by $\mathbf{P}$ and ν :

Theorem 3.1. *A DTMC is uniquely determined by a one-step transition probability matrix $\mathbf{P} \in [0, 1]^{\mathcal{S} \times \mathcal{S}}$ and an initial distribution $\nu \in \text{Distr}(\mathcal{S})$.*

Proof. The proof follows directly from the Markov property and the restriction to time-homogeneous DTMCs. Its details can be found in [Kul95, Thm. 2.2]. $\square$

Theorem 3.1 allows for another interpretation of DTMCs: From a modeling point of view, a DTMC can be imagined as a transition system model, where each transition from a state s to a successor state s' is labeled with the probability $p_{s,s'}$ and moreover, the state changes occur at discrete clock ticks that are global to the system.

Therefore, for the remainder of the thesis, we define a DTMC as follows:

Definition 3.5 (Discrete-time Markov chain). A discrete-time Markov chain is a tuple $\mathcal{D} = (\mathcal{S}, \mathbf{P}, \nu)$, where $\mathcal{S}$ is a finite, nonempty set of states, $\mathbf{P} : \mathcal{S} \times \mathcal{S} \rightarrow [0, 1]$ is a stochastic matrix and $\nu \in \text{Distr}(\mathcal{S})$ is an initial distribution.

This definition allows for a graphical representation of DTMCs, the so-called *state transition diagram*. We introduce this representation by means of an example:

Example 3.1. Consider the DTMC $\mathcal{D}$ depicted in Fig. 3.1: The state space is the set $\mathcal{S} = \{s_0, s_1, s_2, s_3\}$. Moreover, the initial distribution and the one-step transition probability matrix are given as follows:

$$\vec{\nu} = \left(\frac{1}{2}, \frac{1}{2}, 0, 0 \right) \quad \text{and} \quad \mathbf{P} = \begin{pmatrix} 0 & \frac{1}{2} & \frac{1}{3} & \frac{1}{6} \\ 0 & \frac{1}{2} & \frac{1}{2} & 0 \\ 0 & \frac{3}{4} & 0 & \frac{1}{4} \\ 0 & 0 & 1 & 0 \end{pmatrix}.$$

With these two ingredients, we can compute the probability distribution of any random variable X_n of the DTMC's stochastic process $\{X_n\}_{n \in \mathbb{N}}$. For example, we obtain the following distributions for the first two time steps in $\mathcal{D}$:

$$P(X_1 = \cdot) = \vec{\nu} \cdot \mathbf{P} = \left(0, \frac{1}{2}, \frac{5}{12}, \frac{1}{12} \right) \quad \text{and} \quad P(X_2 = \cdot) = \vec{\nu} \cdot \mathbf{P}^2 = \left(0, \frac{9}{16}, \frac{1}{3}, \frac{5}{48} \right). \quad \diamond$$

In Sec. 3.1, a *sample path* of a stochastic processes $\{X_t\}_{t \in T}$ is defined as a function $\pi : T \rightarrow \mathcal{S}$ with the intuition that if the outcome of the stochastic process is π , $\pi(t) = s$ means that the process is in state s at time point $t \in T$. Thus, in the special case of a DTMC $\{X_n\}_{n \in \mathbb{N}}$, a sample path is a function $\pi : \mathbb{N} \rightarrow \mathcal{S}$. However, in the remainder of the thesis we use an alternative (but equivalent) representation of sample paths, that is directly related to the transition diagram of a DTMC. Instead of using a function $\mathbb{N} \rightarrow \mathcal{S}$, we denote sample paths as countably infinite sequences of states.

Using this notation, a path in a DTMC has the form

$$\pi = s_0 \rightarrow s_1 \rightarrow s_2 \rightarrow s_3 \rightarrow \dots$$

and describes the sequence of states that have been traversed in the state transition diagram of $\mathcal{D}$. The link to the sample path definition in Sec. 3.1 is established by noting that each infinite sequence π is in a one-to-one correspondence with a sample path $\pi' : \mathbb{N} \rightarrow \mathcal{S} : n \mapsto \pi[n]$, where $\pi[n] = s_n$ denotes the $(n+1)$ -th state on π .

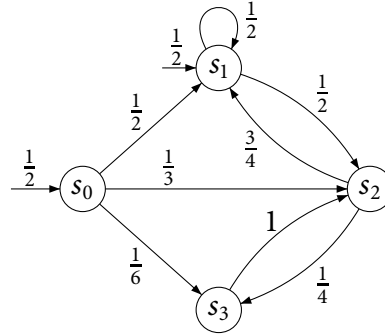


Figure 3.1: The state transition diagram of the DTMC.

Additionally, we sometimes also consider finite paths. Accordingly, we let $Paths_{\mathcal{D}}^n$ denote the sets of paths of length n in $\mathcal{D}$, where the *length* of a finite path π is denoted $|\pi|$ and determined by the number of states on π . Consequently, $Paths_{\mathcal{D}}^* = \bigcup_{n=0}^{\infty} Paths_{\mathcal{D}}^n$ is the set of all finite paths in $\mathcal{D}$ and $Paths_{\mathcal{D}}^{\omega}$ denotes the set of all infinite paths in $\mathcal{D}$. In the following, the reference to $\mathcal{D}$ is omitted whenever it is clear from the context.

The geometric distribution and the memoryless property

A DTMC is closely related to a geometric distribution: Imagine a sequence of random experiments, which either succeed with probability $p \in (0, 1]$ or which fail with probability $(1 - p)$. Now, let X be a random variable for the number of trials that we need to undertake until we succeed for the first time. Formally, we can describe the probability that the n -th experiment is the first that succeeds as follows [ADD00, p. 328]:

$$P(\{X = n\}) = (1 - p)^{n-1} p.$$

Hence, the probabilities $P(\{X = n\})$ for $n = 1, 2, 3, \dots$ form a geometric sequence. To see this, note that $P(\{X = n + 1\})$ is obtained by multiplying $P(\{X = n\})$ with the constant factor $(1 - p)$. With these preliminaries, the geometric distribution has the discrete probability distribution function $F(n)$ given by

$$F(n) = P(\{X \leq n\}) = \sum_{i=1}^n P(\{X = i\}). \quad (3.4)$$

The last term in Eq. (3.4) is a geometric series. Using the well-known formula $\sum_{k=0}^n ar^k = \frac{a(r^{n+1} - 1)}{r - 1}$, we can express $F(n)$ as follows (where $a = 1$ and $r = (1 - p)$):

$$F(n) = \sum_{k=1}^n (1 - p)^{k-1} p = p \sum_{k=0}^{n-1} (1 - p)^k = p \cdot \frac{(1 - p)^n - 1}{(1 - p) - 1} = 1 - (1 - p)^n.$$

Hence, we can also interpret $F(n)$ as the probability that we do not see n failures of the random experiment in a row.

An interesting property of the geometric distribution is that it is memoryless:

Theorem 3.2 (The geometric distribution is memoryless). *Let X be a random variable with a geometric distribution with parameter $p \in (0, 1)$. Then*

$$P(\{X > n + k\} \mid X > n) = P(\{X > k\}). \quad (3.5)$$

Hence, the geometric distribution is memoryless. Moreover, all discrete probability distributions that are memoryless are geometrically distributed.

Proof. We first prove Eq. (3.5):

$$P(\{X > n + k\} \mid X > n) = \frac{P(\{X > n + k \cap X > n\})}{P(\{X > n\})} = \frac{P(\{X > n + k\})}{P(\{X > n\})}.$$

From the derivation of the probability distribution function F , we know that $P(\{X > x\}) = 1 - P(\{X \leq x\}) = 1 - F(x)$. Hence

$$\frac{P(\{X > n + k\})}{P(\{X > n\})} = \frac{1 - F(n + k)}{1 - F(n)} = \frac{1 - (1 - (1 - p)^{n+k})}{1 - (1 - (1 - p)^n)} = \frac{(1 - p)^{n+k}}{(1 - p)^n} = (1 - p)^k.$$

Now $P(\{X > k\}) = 1 - F(k) = (1 - p)^k$, thereby proving Eq. (3.5).

We prove that the geometric distribution is the only discrete probability distribution which is memoryless: We proceed by contraposition and assume that Y is a discrete random variable which is memoryless, but not geometrically distributed. Further, let $F_Y^c(y) = P(\{Y > y\})$. As Y is memoryless, it must hold that $P(\{Y > n + k\} \mid Y > n) = P(\{Y > k\})$. By the law of total probability, we obtain

$$\begin{aligned} F_Y^c(n + k) &= P(\{Y > n + k\}) \\ &= P(\{Y > n + k\} \mid Y > n) \cdot P(\{Y > n\}) \\ &= P(\{Y > k\}) \cdot P(\{Y > n\}) \\ &= F_Y^c(k) \cdot F_Y^c(n) \end{aligned}$$

for all $n, k \in \mathbb{N}$. Therefore $F_Y^c(2) = F_Y^c(1)^2$ (choose $n=k=1$) and $F_Y^c(3) = F_Y^c(1) \cdot F_Y^c(2)$ (with $k=1$ and $n=2$). But then $F_Y^c(3) = F_Y^c(1)^3$.

According to this reasoning, we have that $F_Y^c(m) = F_Y^c(1)^m$ for all $m \in \mathbb{N}_{>0}$. Now, the only discrete function g that satisfies $g(m) = g(1)^m$ has the form $g(m) = q^m$ for some $q \in \mathbb{R}$. Hence, $F_Y^c(m) = q^m$ for some $q \in (0, 1)$. Moreover, $F_Y(m) = 1 - F_Y^c(m) = 1 - q^m$, which is the distribution function of the geometric distribution with parameter $p = 1 - q$. Hence we obtain a contradiction, as the random variable Y is geometrically distributed. $\square$

To see how the geometric distribution is related to our definition of a discrete-time Markov chain, recall that we require a DTMC to have the Markov property. Now, the time that a DTMC spends in a given state s is geometrically distributed. To see this, let $(\mathcal{S}, \mathbf{P}, \nu)$ be a

DTMC and fix an arbitrary state $s \in \mathcal{S}$. If the random variable N (defined on $\{1, 2, 3, \dots\}$) describes the number of time steps that the DTMC remains in state s , then

$$\begin{aligned} P(\{N = 1\}) &= 1 - p_{s,s} \\ P(\{N = 2\}) &= (1 - p_{s,s}) \cdot p_{s,s} \\ P(\{N = 3\}) &= (1 - p_{s,s}) \cdot p_{s,s}^2 \\ &\vdots = \vdots \end{aligned}$$

Hence, we have that N is geometrically distributed with $p = 1 - p_{s,s}$.

This is not a random coincidence: Intuitively, the Markov property states that the information that the DTMC has been in a state s for a certain amount of time already, must not influence the distribution of the remaining sojourn time.

At this point, we conclude the discussion of DTMCs, inevitably leaving many theoretical gaps open. However, we have covered the fundamental properties that we will need in the remainder of this thesis. An otherwise important topic that we have ignored completely, is the definition of a DTMC's *steady state*. It can be imagined as the probability to be in a given state of the DTMC after a (very) long time. However, in the controlled Markov processes that we investigate later, steady states generally do not exist. Hence, we do not go into the details here but refer to the broad selection of literature about the topic, for example [Kul95].

3.2.2 Continuous-time Markov chains

After having introduced discrete-time Markov chains, this section discusses their continuous-time analogue. A *continuous-time Markov chain* (CTMC) is a Markov chain $\{X_t\}_{t \in T}$ with parameter set $T = \mathbb{R}_{\geq 0}$, such that each random variable X_t describes the state of the CTMC at time point t .

Compared to DTMCs, the definition of CTMCs is slightly more involved. Similar to DTMCs, the Markov property also applies to CTMCs: If a CTMC is in state $s_n \in \mathcal{S}$ at time $t_n \in \mathbb{R}_{\geq 0}$, its future behavior does not depend on the states $s_{n-1}, s_{n-2}, \dots, s_1, s_0$, that have been observed at some time points $t_{n-1} > t_{n-2} > \dots > t_1 > t_0 \in \mathbb{R}_{\geq 0}$. Formally, the Markov property for CTMCs is stated as follows: Let $A \subseteq \mathcal{S}$ be a set of states and $n \in \mathbb{N}$. For all decreasing sequences of time points $t_{n+1} > t_n > \dots > t_1 > t_0 \in \mathbb{R}_{\geq 0}$ and states $s_n, s_{n-1}, \dots, s_1, s_0$, it holds that [Hav00, Sec. 4.1]

$$\begin{aligned} P(\{X_{t_{n+1}} \in A\} \mid X_{t_n} = s_n, X_{t_{n-1}} = s_{n-1}, \dots, X_{t_1} = s_1, X_{t_0} = s_0) \\ = P(\{X_{t_{n+1}} \in A\} \mid X_{t_n} = s_n). \end{aligned} \quad (3.6)$$

It is important to note a subtle difference to the discrete-time case: There, Eq. (3.1) (see page 58) summarizes the Markov property and time-homogeneity for DTMCs by considering a discrete time step. As this discrete notion of a time step does not exist in CTMCs, the probability $P(\{X_{t_{n+1}} \in A\} \mid X_{t_n} = s_n)$ in Eq. (3.6) depends on the amount of

time $\Delta_t = t_{n+1} - t_n$ that has passed since the last time (t_n in our notation), the state of the CTMC has been observed. We will come back to this point later, when we discuss the transition probabilities of a CTMC.

The second important property of a CTMC is time-homogeneity. Together with the Markov property in Eq. (3.6), time homogeneity is expressed as follows:

$$P(\{X_{t_{n+1}} \in A\} \mid X_{t_n} = s_n) = P(\{X_{\Delta_t} \in A\} \mid X_0 = s_n). \quad (3.7)$$

Therefore, Eq. (3.7) implies that the future behavior of a CTMC depends only on Δ_t and on the current state s_n . In particular, it does neither depend on the previous history (by Eq. (3.6)) nor on the amount of time t that has passed (cf. Eq. (3.7)) before the current state was entered at time t_n . In Eq. (3.7) and Eq. (3.6) we may interpret t_n as the current time and $t + \Delta_t$ as the time in the future, when we observe the state of the CTMC again.

For a time period $\Delta_t > 0$, the probability to move from the current state s_n to a state in the set $A \subseteq \mathcal{S}$ within Δ_t time units is determined by some parameter $\lambda \in \mathbb{R}_{>0}$ such that

$$P(\{X_{\Delta_t} \in A\} \mid X_0 = s_n) = \lambda \cdot \Delta_t + o(\Delta_t), \quad (3.8)$$

where the second summand $o(\Delta_t)$ denotes the probability that multiple transitions occur within time interval $[0, \Delta_t)$. The Landau notation $o(\Delta_t)$ that is used in Eq. (3.8) is defined such that for functions $f, g : \mathbb{R} \rightarrow \mathbb{R}$ it holds that $f \in o(g) \iff \lim_{x \rightarrow \infty} \frac{f(x)}{g(x)} = 0$. Therefore, for small enough Δ_t , the probability that we “miss” intermediate transitions can safely be ignored.

With these remarks, we can interpret Eq. (3.8) as follows: If the time Δ_t that has passed since the last observation of the CTMC’s state is short enough, the probability to move from state s_n to a state in the set A scales linearly with parameter $\lambda > 0$.

Hence, the knowledge about the current state of a CTMC and the parameters λ completely describe the future behavior of a CTMC. In the discrete-time case, the number of steps that a DTMC sojourns in a state is geometrically distributed (cf. Sec. 3.2.1). Similarly, the Markov property and time-homogeneity imply that the sojourn times in a CTMC obey the *exponential distribution*. Before we continue the discussion of the behavior of CTMCs, let us shortly recall the important properties of the exponential distribution:

The exponential distribution

The exponential distribution is a continuous probability distribution which is determined by a rate parameter $\lambda \in \mathbb{R}_{>0}$. Figure 3.2 plots its cumulative distribution function for different rate parameters.

Definition 3.6 (Exponential distribution). Let $\lambda \in \mathbb{R}_{>0}$ be a rate, $t, z \in \mathbb{R}_{\geq 0}$ and

$$f_\lambda(t) = \lambda e^{-\lambda t} \text{ and}$$

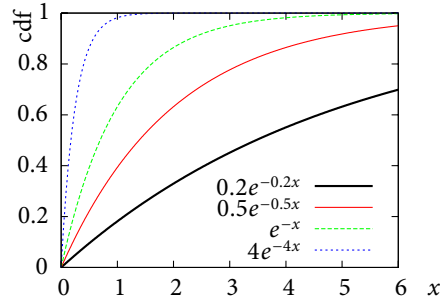


Figure 3.2: Plot of the exponential distribution (cdf) for rates $\lambda = 0.2, 0.5, 1$ and 4 .

$$F_\lambda(z) = \int_0^z f_\lambda(t) dt = \int_0^z \lambda e^{-\lambda t} dt = 1 - e^{-\lambda z}.$$

Then f_λ is the probability density function and F_λ the cumulative distribution function of the negative exponential distribution.

From Def. 3.6, we can directly conclude:

Corollary 3.1. *The rate $\lambda \in \mathbb{R}_{>0}$ uniquely determines an exponential distribution.*

In contrast to DTMC, where a transition between a pair $(s, s') \in \mathcal{S} \times \mathcal{S}$ of states are taken at discrete time steps with a certain probability $\mathbf{P}(s, s')$, CTMCs are continuous stochastic processes. Therefore, the transitions in a CTMC are characterized by a *transition rate* $\mathbf{R}(s, s')$. In a CTMC, the value $\mathbf{R}(s, s')$ is interpreted as the rate of an exponential distribution which governs the transition's delay. Similar to the DTMC case, a CTMC is completely characterized by its *transition rate matrix* $\mathbf{R}$ and an initial distribution. As we have seen in Sec. 3.2.1, the time that a DTMC stays in the same state (given by the number of discrete time ticks) obeys a geometric distribution. The exponential distribution is its counterpart in the continuous-time domain:

Theorem 3.3 (The exponential distribution is memoryless). *Let X be a random variable with an exponential distribution. Then*

$$P(\{X > x + k\} \mid X > k) = P(\{X > x\}) \quad (3.9)$$

for all $x, k \in \mathbb{R}_{\geq 0}$. Hence, the exponential distribution is memoryless. Moreover, any continuous distribution which is memoryless is an exponential distribution.

Proof. The proof is similar to that of Thm. 3.2 and can be found in, e.g. [Kul95, p. 189]. $\square$

If X and Y are two independent, exponentially distributed random variables with rates λ_1 and λ_2 , then the *minimum* of X and Y is again exponentially distributed:

Lemma 3.2. *Let $X \sim \text{Exp}(\lambda_1)$ and $Y \sim \text{Exp}(\lambda_2)$ be independent random variables with rates $\lambda_1, \lambda_2 \in \mathbb{R}_{>0}$. Then $P(\{\min(X, Y) \leq z\}) = (1 - e^{-(\lambda_1 + \lambda_2)z})$ for all $z \in \mathbb{R}_{\geq 0}$.*

Proof. For the proof, we consider the joint distribution of X and Y :

$$\begin{aligned}
 P(\{\min(X, Y) \leq z\}) &= P_{X,Y}(\{(x, y) \in \mathbb{R}_{\geq 0}^2 \mid \min(x, y) \leq z\}) \\
 &= \int_0^\infty \left(\int_0^\infty \mathbf{I}_{\min(x,y) \leq z}(x, y) \cdot \lambda_1 e^{-\lambda_1 x} \cdot \lambda_2 e^{-\lambda_2 y} dy \right) dx \\
 &= \int_0^z \int_x^\infty \lambda_1 e^{-\lambda_1 x} \cdot \lambda_2 e^{-\lambda_2 y} dy dx + \int_0^z \int_y^\infty \lambda_1 e^{-\lambda_1 x} \cdot \lambda_2 e^{-\lambda_2 y} dx dy \\
 &= \int_0^z \lambda_1 e^{-\lambda_1 x} \cdot e^{-\lambda_2 x} dx + \int_0^z \lambda_2 e^{-\lambda_2 y} \cdot e^{-\lambda_1 y} dy \\
 &= \int_0^z \lambda_1 e^{-(\lambda_1 + \lambda_2)x} dx + \int_0^z \lambda_2 e^{-(\lambda_1 + \lambda_2)y} dy \\
 &= \int_0^z (\lambda_1 + \lambda_2) \cdot e^{-(\lambda_1 + \lambda_2)t} dt = (1 - e^{-(\lambda_1 + \lambda_2)z}). \quad \square
 \end{aligned}$$

Hence the class of exponential distributions is closed under minimum.

In a similar way, we can prove that the probability that the outcome of the random experiment associated with X is less than that of Y is given by the fraction $\frac{\lambda_1}{\lambda_1 + \lambda_2}$:

Lemma 3.3. *For two independent random variables $X \sim \text{Exp}(\lambda_1)$ and $Y \sim \text{Exp}(\lambda_2)$ with rates $\lambda_1, \lambda_2 \in \mathbb{R}_{>0}$ it holds $P(\{X \leq Y\}) = \frac{\lambda_1}{\lambda_1 + \lambda_2}$.*

Proof. Again by the joint distribution function:

$$\begin{aligned}
 P(\{X \leq Y\}) &= P_{X,Y}(\{(x, y) \in \mathbb{R}_{\geq 0}^2 \mid x \leq y\}) = \int_0^\infty \lambda_2 e^{-\lambda_2 y} \left(\int_0^y \lambda_1 e^{-\lambda_1 x} dx \right) dy \\
 &= \int_0^\infty \lambda_2 e^{-\lambda_2 y} (1 - e^{-\lambda_1 y}) dy = 1 - \int_0^\infty \lambda_2 e^{-\lambda_2 y} e^{-\lambda_1 y} dy \\
 &= 1 - \int_0^\infty \lambda_2 e^{-(\lambda_1 + \lambda_2)y} dy = 1 - \frac{\lambda_2}{\lambda_1 + \lambda_2} \cdot \int_0^\infty (\lambda_1 + \lambda_2) e^{-(\lambda_1 + \lambda_2)y} dy \\
 &= 1 - \frac{\lambda_2}{\lambda_1 + \lambda_2} = \frac{\lambda_1}{\lambda_1 + \lambda_2}. \quad \square
 \end{aligned}$$

Obviously, we obtain $P(\{Y \leq X\}) = \frac{\lambda_2}{\lambda_1 + \lambda_2}$ in exactly the same way. Moreover, we can prove in the same way as in Lemma 3.3, that the probability that the value of the i -th random variable is the smallest of a sequence of independent random variables $X_k \sim \text{Exp}(\lambda_k)$ for $k = 1, 2, \dots, n$ is $\frac{\lambda_i}{\sum_{k=1}^n \lambda_k}$. Finally, as the exponential distribution is continuous, we have for any exponentially distributed random variable X that $P(\{X = c\}) = 0$ for all $c \in \mathbb{R}_{\geq 0}$.

With these preliminaries, we are ready to fully describe the behavior of a CTMC:

The definition of continuous-time Markov chains

A continuous-time Markov chain is defined by its transition rates $\mathbf{R}(s, s')$: For states s and s' , the value of $\mathbf{R}(s, s')$ specifies the *rate* of the transition that leads from state s to its successor state s' . If no such transition exists then $\mathbf{R}(s, s') = 0$. The values $\mathbf{R}(s, s') \in \mathbb{R}_{\geq 0}$ form the *transition rate matrix* of a CTMC. Roughly, it is the continuous-time counterpart of a DTMC's one-step transition probability matrix.

If $X_{s,s'} \sim \text{Exp}(\mathbf{R}(s, s'))$ denotes the random variable that is distributed with rate $\mathbf{R}(s, s')$, then $X_{s,s'}$ can be understood as the delay that is needed for the transition from state s to state s' to execute. For multiple successor states, consider the situation depicted in Fig. 3.3: Here, transitions lead from state s_0 to states s_1, s_2 and s_3 . Each of them has an exponentially distributed delay, described by the rates $\mathbf{R}(s_0, s_1)$, $\mathbf{R}(s_0, s_2)$ and $\mathbf{R}(s_0, s_3)$, respectively. Two obvious questions arise if we consider the behavior in state s_0 :

- (a) What is the probability to take the transition to, say, state s_2 ?
- (b) How long is the sojourn in state s_0 ?

The three transitions that leave state s_0 compete for execution, that is, the first transition whose delay expires, executes and determines the successor state. Therefore, we may reformulate question (a) and ask for the probability that the delay of the transition that leads to state s_2 expires before the delays of the other two transitions. Formally, this corresponds to the probability that the sample drawn for the random variable X_{s_0, s_2} is less than the samples drawn for X_{s_0, s_1} and X_{s_0, s_3} :

$$P(\{X_{s_0, s_2} \leq X_{s_0, s_1}\} \cap \{X_{s_0, s_2} \leq X_{s_0, s_3}\}).$$

As the random variables are independent, we obtain in the same way as in the proof of Lemma 3.3 that

$$P(\{X_{s_0, s_2} \leq X_{s_0, s_1}\} \cap \{X_{s_0, s_2} \leq X_{s_0, s_3}\}) = \frac{\mathbf{R}(s_0, s_2)}{\mathbf{R}(s_0, s_1) + \mathbf{R}(s_0, s_2) + \mathbf{R}(s_0, s_3)}.$$

The situation depicted in Fig. 3.3 is known as a *race condition*, as the outgoing transitions compete for execution according to their associated rates.

To answer question (b), note that the sojourn time in state s_0 is governed by the time that it takes for the first transition to execute. As this equals the minimum delay of the

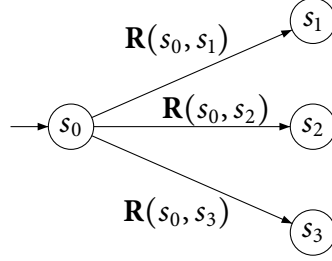


Figure 3.3: Race condition in a (fragment) CTMC.

outgoing transitions, the sojourn time in state s_0 is described by the random variable $Y_0 = \min \{X_{s_0, s_1}, X_{s_0, s_2}, X_{s_0, s_3}\}$. By Lemma 3.2, we conclude that the probability distribution of the sojourn time Y_0 in state s_0 is

$$\begin{aligned} P(\{Y_0 \leq z\}) &= P(\{\min(X_{s_0, s_1}, X_{s_0, s_2}, X_{s_0, s_3}) \leq z\}) = 1 - e^{-(R(s_0, s_1) + R(s_0, s_2) + R(s_0, s_3))z} \\ &= 1 - e^{-E(s_0)z}. \end{aligned}$$

Hence, the sojourn in state s_0 is exponentially distributed with the sum of the rates of all transitions that leave state s_0 . Formally, this sum is the *exit rate* of state s_0 and defined as $E(s_0) = \sum_{s' \in \mathcal{S}} R(s_0, s') = R(s_0, s_1) + R(s_0, s_2) + R(s_0, s_3)$. Thus, the sojourn time Y in a state s is obtained by the equation

$$P(\{Y \leq z\}) = \int_0^z E(s) e^{-E(s)t} dt = (1 - e^{-E(s)z}).$$

As in the case of DTMCs, we also use state transition diagrams to graphically represent CTMCs, where we augment the transitions with the corresponding entry in the CTMC's rate matrix (instead of the probabilities that are given by a DTMC's one-step transition probability matrix).

Definition 3.7 (Continuous-time Markov chain). A continuous-time Markov chain is a tuple $(\mathcal{S}, \mathbf{R}, \nu)$, where $\mathcal{S}$ is the finite set of states, $\mathbf{R} : \mathcal{S} \times \mathcal{S} \rightarrow \mathbb{R}_{\geq 0}$ is the two-dimensional rate matrix and $\nu \in \text{Distr}(\mathcal{S})$ is an initial distribution.

As done in [BHHK03], we assume that the CTMC does not contain deadlock states and require that $E(s) = \sum_{s' \in \mathcal{S}} R(s, s') > 0$ for all states $s \in \mathcal{S}$.

If we abstract from the sojourn times in a CTMC, we obtain its *embedded DTMC*: Let $(\mathcal{S}, \mathbf{R}, \nu)$ be a CTMC. Its embedded DTMC $(\mathcal{S}, \mathbf{P}, \nu)$ is given by the probability matrix $\mathbf{P}$ defined as $\mathbf{P}(s, s') = \frac{R(s, s')}{E(s)}$. Intuitively, for states $s, s' \in \mathcal{S}$ the value of $\mathbf{P}(s, s')$ is the probability that the transition that leads from state s to state s' in the underlying CTMC executes first. In this way, the embedded DTMC abstracts from the timing information in a CTMC and only considers its time-abstract behavior.

3.3 Nondeterminism in stochastic models

In the previous section, we discussed discrete and continuous time Markov chains. These models are complete in the sense, that their underlying stochastic process is uniquely determined. In this section, we extend the notion of a Markov chains and also allow that nondeterministic choices may occur in the model. Thereby, we arrive at the definition of discrete- and continuous-time Markov decision processes.

We follow the same route as in Sec. 3.2 and consider discrete-time Markov decision processes first. Afterwards, Sec. 3.3.2 discusses continuous-time Markov decision processes in detail.

3.3.1 Discrete time Markov decision processes

Discrete-time Markov decision processes [Bel57, How71, Ber95, Put94] (MDPs) have already been discovered in the late 1950's. They are applied widely in mathematics and operations research. Moreover, with value iteration [Bel57] and policy iteration [How60], two techniques exist which are well understood and permit to solve MDPs algorithmically.

In computer science, MDPs are of particular interest: As discovered by Vardi [Var85], they allow us to model the behavior of randomized distributed algorithms. An example of such an algorithm is a leader election protocol, where ties are broken by probabilistic choices [IR90].

Furthermore, the support of nondeterminism in MDPs allows us to use abstraction techniques such as simulation relations to reduce the state space of discrete-time Markov chains and MDPs [DJJL01]. In this application, states with different behavior are grouped together, yielding a set of different possible probabilistic behaviors. As the identity of the underlying states is hidden in the abstract system, their different behaviors give rise to nondeterministic choices. In this way, abstracting DTMCs yields discrete-time Markov decision processes.

Each state of an MDP is equipped with a finite set of one-step transition probability distributions, each of which is uniquely identified by an action. Hence, the actions indicate the nondeterministic choices available in a state.

Definition 3.8 (Discrete-time Markov decision process). A discrete-time Markov decision process (MDP) is a tuple $(\mathcal{S}, Act, \mathbf{P}, \nu)$, where $\mathcal{S}$ and Act are finite, nonempty sets of states and actions and $\nu \in \text{Distr}(\mathcal{S})$ is an initial distribution. Moreover, $\mathbf{P} : \mathcal{S} \times Act \times \mathcal{S} \rightarrow [0, 1]$ is a three-dimensional probability matrix which satisfies $\sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \in \{0, 1\}$.

Let $\mathcal{M} = (\mathcal{S}, Act, \mathbf{P}, \nu)$ be an MDP. An action $\alpha \in Act$ is *enabled* in a state $s \in \mathcal{S}$ iff $\sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') = 1$. Accordingly, the set

$$Act(s) = \left\{ \alpha \in Act \mid \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') = 1 \right\}$$

is the set of *enabled actions in state s* . We require that $|Act(s)| > 0$ for all states $s \in \mathcal{S}$. Note that this is no restriction, as any *deadlock state* with $Act(s) = \emptyset$ is never left. Therefore, it can safely be equipped with a self-loop transition $\mathbf{P}(s, \alpha, s) = 1$ for some action $\alpha \in Act$ without altering the MDP's semantics.

As for DTMCs and CTMCs, the initial distribution ν quantifies the probability that the MDP starts in a certain state. We say that a state $s \in \mathcal{S}$ is an *initial state* of the MDP $\mathcal{M}$ if its initial distribution is degenerate and of the form $\nu = \{s \mapsto 1\}$. In this case, the evolution of the MDP definitely starts in state s . In principle, Def. 3.8 could be extended to allow for sets of initial distributions. However, to simplify the technicalities, throughout this thesis, we assume that nondeterministic models are equipped with a fixed initial distribution.

The behavior of an MDP can be described as follows: The first state of the MDP is determined by the initial distribution ν . When entering a state $s \in \mathcal{S}$, each enabled action $\alpha \in Act(s)$ corresponds to one possibility to resolve the nondeterministic choices that are represented by the set of enabled actions $Act(s)$. More precisely, each action identifies a probability distribution $\mathbf{P}(s, \alpha, \cdot) \in Distr(\mathcal{S})$, where $\mathbf{P}(s, \alpha, s')$ is the probability that the MDP moves from state s to *successor state s'* . In general, several actions are enabled in state s , denoting different probability distributions. Therefore, to reason about probability measures in MDPs, it is necessary to resolve the nondeterminism by choosing one action from the set $Act(s)$.

Markov decision processes, whose set of enabled actions are singletons, i.e. if $|Act(s)| = 1$ holds for all $s \in \mathcal{S}$, are semantically equivalent to DTMCs. To see this, note that such an MDP does not contain any nondeterministic choices as only one selectable action remains in each state. Conversely, each DTMC can be construed as an MDP of the above form. Therefore, the class of DTMCs is a proper subclass of MDPs.

Note that the Markov property also holds for MDPs, that is, after an action $\alpha \in Act(s)$ has been chosen, its effect only depends on the current state s and not on the states that have been traversed before.

Example 3.2. Figure 3.4 depicts an MDP with initial state s_0 . A nondeterministic choice occurs between actions α and β upon entering state s_1 ; all other states are deterministic, that is, their sets of enabled actions are singletons. If action α is chosen in state s_1 , the probabilities to move to states s_2 or back to s_0 are $\mathbf{P}(s_1, \alpha, s_2) = \frac{1}{3}$ and $\mathbf{P}(s_1, \alpha, s_0) = \frac{2}{3}$, respectively. For action β , the probability to reach state s_0 or state s_2 is zero; instead, we stay in state s_1 with probability $\mathbf{P}(s_1, \beta, s_1) = \frac{3}{4}$ and move to state s_3 with the remaining probability $\mathbf{P}(s_1, \beta, s_3) = \frac{1}{4}$. $\diamond$

Formally, a *path* is a finite or infinite sequence of states and actions. Whereas paths in MDPs are *time abstract*, we need to consider *time-dependent* paths later. To distinguish between the two variants, we mark the sets of time-abstract paths with subscript *abs*. According to this notation, $Paths_{abs}^n = \mathcal{S} \times (Act \times \mathcal{S})^n$ denotes the set of all paths of length n ; similarly, $Paths_{abs}^* = \bigcup_{n=0}^{\infty} Paths_{abs}^n$ and $Paths_{abs}^\omega = \mathcal{S} \times (Act \times \mathcal{S})^\omega$ denote the sets of finite

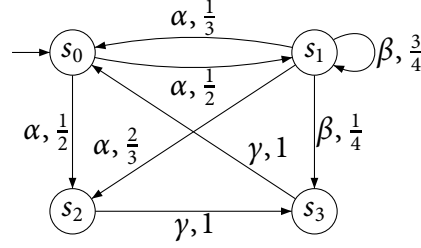


Figure 3.4: An example of a discrete-time Markov decision process.

and infinite paths, resp. For notational convenience, we describe paths in the form

$$\pi = s_0 \xrightarrow{\alpha_0} s_1 \xrightarrow{\alpha_1} s_2 \xrightarrow{\alpha_2} \dots$$

If π is a finite path that ends in state s_n , we use $\pi \downarrow = s_n$ and $|\pi| = n$ to denote the last state on π and the length of π , respectively. Informally, a time abstract path records the states and actions that are traversed by an MDP and thereby describes one instance of the random behaviors of an MDP together with the actions that have been chosen.

At this stage, we cannot assign probabilities to sets of paths in $2^{\text{Paths}_{abs}^\omega}$: To see why, reconsider state s_1 from the MDP in Ex. 3.2. Up to now, we cannot answer questions like “What is the probability that being in state s_1 , the next state is s_3 ?”, as the probability depends on whether action α or action β are chosen in state s_1 . Even if we assume that β is chosen, this does not imply that β is chosen again, if state s_1 is re-entered later.

Schedulers solve this problem by quantifying the nondeterministic choices in each state of an MDP. In the following definition, we slightly generalize the intuition of a scheduler and consider *randomized* schedulers, which can not only decide for a single action, but may also yield a probability distribution over the next actions:

Definition 3.9 (MDP scheduler). Let $\mathcal{M} = (\mathcal{S}, \text{Act}, \mathbf{P}, \nu)$ be an MDP. An MDP scheduler for $\mathcal{M}$ is a mapping $D : \text{Paths}_{abs}^* \rightarrow \text{Distr}(\text{Act})$ such that $D(\pi)(\alpha) > 0$ implies $\alpha \in \text{Act}(\pi \downarrow)$ for all $\pi \in \text{Paths}_{abs}^*$.

The condition in Def. 3.9 implies that if a scheduler assigns a positive probability to an action α , this action is indeed enabled in the current state $\pi \downarrow$.

The combination of an MDP $\mathcal{M}$ and an MDP scheduler D for $\mathcal{M}$ uniquely determines the probabilistic behavior of the MDP. Informally, when $\mathcal{M}$ enters a state after it has traversed path π , the scheduler D resolves the nondeterministic choice between the available actions in the current state $\pi \downarrow$. If action α is chosen, the resulting probability distribution $\mathbf{P}(\pi \downarrow, D(\pi \downarrow), \cdot)$ governs the next state that is occupied by the MDP.

To measure probabilities in an MDP, we use the smallest σ -field of subsets of $\text{Paths}_{abs}^\omega$, that is generated by the measurable cylinders (cf. Sec. 2.5); we denote it by $\mathfrak{F}_{\text{Paths}_{abs}^\omega}$. Hence,

the elements of $\mathfrak{F}_{Paths_{abs}^\omega}$ have the form

$$\{\pi \in Paths_{abs}^\omega \mid \pi \in \Pi^n\}$$

for some cylinder base $\Pi^n \subseteq Paths_{abs}^n$. Observe that in the discrete-time setting, no measurability issues arise as all sets are finite or countably infinite. Therefore, we do not need to restrict ourselves to *measurable* cylinder bases but can simply assume that $\Pi^n \subseteq Paths_{abs}^n$. Accordingly, we use $\mathfrak{F}_{Paths_{abs}^n} = 2^{Paths_{abs}^n}$ as the σ -field over subsets of paths of length n .

The definition of the probability measure of MDPs is standard and can be found in, for example [dA97]. However, to ease the understanding of the probability measure for continuous-time Markov decision processes which is introduced in Sec. 3.3.2, we restate the definition for MDPs here in the same notation:

Definition 3.10 (Probability measure). *Let $\mathcal{M} = (\mathcal{S}, Act, \mathbf{P}, v)$ be an MDP and D be an MDP scheduler for $\mathcal{M}$. The probability measure $Pr_{v,D}^n$ on $(Paths_{abs}^n, 2^{Paths_{abs}^n})$ is inductively defined as follows:*

$$\begin{aligned} Pr_{v,D}^0 : \mathfrak{F}_{Paths_{abs}^0} &\rightarrow [0,1] : \Pi \mapsto \sum_{s \in \Pi} v(\{s\}) \quad \text{and} \\ Pr_{v,D}^{n+1} : \mathfrak{F}_{Paths_{abs}^{n+1}} &\rightarrow [0,1] : \Pi \mapsto \sum_{\pi \in Paths_{abs}^n} Pr_{v,D}^n(\{\pi\}) \sum_{\alpha \in Act} D(\pi)(\alpha) \sum_{s' \in \mathcal{S}} \mathbf{I}_\Pi(\pi \xrightarrow{\alpha} s') \cdot \mathbf{P}(\pi \downarrow, \alpha, s'). \end{aligned}$$

Note that $\mathbf{I}_\Pi$ is an indicator function such that $\mathbf{I}_\Pi(\pi) = 1$ if $\pi \in \Pi$ and 0, otherwise. Definition 3.10 inductively derives a family of probability measures, each defined on sets of paths of some (finite) length n : Note that a set of paths of length 0 is just a set of states. Obviously, the probability to start in a state from the set $\Pi^0 \subseteq \mathcal{S}$ is given by the sum of the initial probabilities for all states $s \in \Pi^0$.

By the inductive definition, we may rely on the measure for sets of paths of length n in measuring paths of length $n+1$. More precisely, in Def. 3.10 we obtain the probability of a set of paths $\Pi \subseteq Paths_{abs}^{n+1}$ by multiplying the probability of all paths of length n with all one-step extensions; the indicator $\mathbf{I}_\Pi$ is then used to project on the set Π .

At this point, the definition of $Pr_{v,D}^n$ might appear overly complex. However, this generality allows us to define the probability measures in the continuous-time case (cf. Sec. 3.3.2) in a very similar way. Let us formally prove that Def. 3.10 indeed coincides with the semantics of MDPs that is found in the literature: As each $\mathfrak{F}_{Paths_{abs}^n}$ belongs to a discrete probability space, the measure of a set of paths $\Pi^n \subseteq Paths_{abs}^n$ is defined by the sum of the probabilities of all elements in Π^n . To map our definition to the standard notation, as given in [dA97, Sec. 3.1.2], note that the probability of a single path $\pi = s_0 \xrightarrow{\alpha_0} s_1 \xrightarrow{\alpha_1} \dots \xrightarrow{\alpha_{n-1}} s_n$ is given by the product

$$p(D, \pi) = v(\pi[0]) \cdot \prod_{i=0}^{n-1} \mathbf{P}(\pi[i], \alpha_i, \pi[i+1]) \cdot D\left(s_0 \xrightarrow{\alpha_0} s_1 \xrightarrow{\alpha_1} \dots \xrightarrow{\alpha_{i-1}} s_i\right)(\alpha_i).$$

Now a simple inductive proof shows that our definition of the probability $Pr_{v,D}^n(\Pi^n)$ of the set of paths $\Pi \subseteq Paths_{abs}^n$ coincides with that used in [dA97, BK08]:

Lemma 3.4 (Probability measure). *Let $\mathcal{M} = (\mathcal{S}, Act, \mathbf{P}, v)$ be an MDP, D an MDP scheduler for $\mathcal{M}$ and $\Pi^n \subseteq Paths_{abs}^n$ for some $n \in \mathbb{N}$. Then*

$$Pr_{v,D}^n(\Pi) = \sum_{\pi \in \Pi^n} p(D, \pi). \quad (3.10)$$

Proof. We prove Eq. (3.10) by induction on n :

1. The induction base follows trivially, as $Pr_{v,D}^0(\Pi^0) = \sum_{s \in \Pi^0} v(\{s\}) = \sum_{\pi \in \Pi^0} p(D, \pi)$.
2. In the induction step ($n \rightsquigarrow n+1$), we use as induction hypothesis that $Pr_{v,D}^n(\Pi^n) = \sum_{\pi \in \Pi^n} p(D, \pi)$ holds for all $\Pi^n \subseteq Paths_{abs}^n$. Then

$$\begin{aligned} Pr_{v,D}^{n+1}(\Pi^{n+1}) &= \sum_{\pi \in Paths_{abs}^n} Pr_{v,D}^n(\{\pi\}) \sum_{\alpha \in Act} D(\pi)(\alpha) \sum_{s' \in \mathcal{S}} \mathbf{I}_{\Pi^{n+1}}(\pi \xrightarrow{\alpha} s') \cdot \mathbf{P}(\pi \downarrow, \alpha, s') \\ &= \sum_{\pi \in Paths_{abs}^n} p(D, \pi) \sum_{\alpha \in Act} D(\pi)(\alpha) \sum_{s' \in \mathcal{S}} \mathbf{I}_{\Pi^{n+1}}(\pi \xrightarrow{\alpha} s') \cdot \mathbf{P}(\pi \downarrow, \alpha, s') \\ &= \sum_{\pi \in Paths_{abs}^n} \sum_{\alpha \in Act} \sum_{s' \in \mathcal{S}} \mathbf{I}_{\Pi^{n+1}}(\pi \xrightarrow{\alpha} s') \cdot p(D, \pi) \cdot D(\pi)(\alpha) \cdot \mathbf{P}(\pi \downarrow, \alpha, s') \\ &= \sum_{\pi \in Paths_{abs}^n} \sum_{\alpha \in Act} \sum_{s' \in \mathcal{S}} \mathbf{I}_{\Pi^{n+1}}(\pi \xrightarrow{\alpha} s') \cdot p(D, \pi \xrightarrow{\alpha} s') \\ &= \sum_{\pi \in \Pi^{n+1}} p(D, \pi \xrightarrow{\alpha} s'). \quad \square \end{aligned}$$

Ultimately, we are interested in the probability measure on sets of infinite paths. The probability measures $Pr_{v,D}^n$ for sets of paths with length n that are obtained in Def. 3.10 directly extend to a unique probability measure on the σ -field $\mathfrak{F}_{Paths_{abs}^\omega}$. Recall that $\mathfrak{F}_{Paths_{abs}^\omega}$ is the smallest σ -field generated by the measurable cylinders.

The measure theoretical arguments that justify the cylinder set construction have been discussed in detail in Sec. 2.5. Here, we only state the definition of the probability measure $Pr_{v,D}^\omega$ on cylinders. Given a cylinder $B_n \in \mathfrak{F}_{Paths_{abs}^\omega}$ with cylinder base $B^n \in \mathfrak{F}_{Paths_{abs}^n}$, we define

$$Pr_{v,D}^\omega(B_n) = Pr_{v,D}^n(B^n).$$

By the Ionescu-Tulcea extension theorem (Thm. 2.19 on page 51), this definition suffices to uniquely determine the probability of all events in $\mathfrak{F}_{Paths_{abs}^\omega}$. Therefore, we have completed the construction of the probability space $(Paths_{abs}^\omega, \mathfrak{F}_{Paths_{abs}^\omega}, Pr_{v,D}^\omega)$ that is associated with an MDP $\mathcal{M} = (\mathcal{S}, Act, \mathbf{P}, v)$ and an MDP scheduler D for $\mathcal{M}$.

The MDP schedulers that we have considered so far are *history dependent*: Upon entering a state s of an MDP, the decision taken by an MDP scheduler D depends not only on the current state s , but on the history π that led into s . In particular, D 's decision may be different each time state s is entered.

However, in many cases, simpler schedulers suffice. More precisely, if the measure of interest is the maximal (or minimal) probability to reach a set of goal states in an MDP, a deterministic and positional scheduler exists which induces the optimal probabilities [dA97],[BK08, Lemma 10.102]. An MDP scheduler D is positional iff $D(\pi) = D(\pi')$ for all $\pi, \pi' \in \text{Paths}_{abs}^*$ with $\pi \downarrow = \pi' \downarrow$; moreover, it is *deterministic* iff for all $s \in \mathcal{S}$ there exists $\alpha \in \text{Act}$ such that $D(\pi) = \{\alpha \mapsto 1\}$.

The situation becomes more complicated if we aim at finding a scheduler that optimizes (i.e. maximizes or minimizes) the reachability of a set of goal states within a certain number of steps. For such *step-bounded reachability probabilities*, the class of deterministic *hop-counting schedulers* suffices. A scheduler is *hop counting* iff $D(\pi) = D(\pi')$ for all $\pi, \pi' \in \text{Paths}_{abs}^*$ with $\pi \downarrow = \pi' \downarrow$ and $|\pi| = |\pi'|$.

Example 3.3. Reconsider the MDP $\mathcal{M}$ depicted in Fig. 3.4. The positional MDP schedulers D_α and D_β are uniquely determined by $D_\alpha(s_1) = \{\alpha \mapsto 1\}$ and $D_\beta(s_1) = \{\beta \mapsto 1\}$. The induced probability to reach state s_3 within 2 steps is derived as follows:

We consider the event $\diamond^{\leq 2} \{s_3\} = \{\pi \in \text{Paths}_{abs}^\omega \mid \exists k \leq 2. \pi[k] = s_3\}$ and compute the probabilities $Pr_{v,D_\alpha}^\omega(\diamond^{\leq 2} \{s_3\})$ and $Pr_{v,D_\beta}^\omega(\diamond^{\leq 2} \{s_3\})$, respectively:

$$\begin{aligned} Pr_{v,D_\alpha}^\omega(\diamond^{\leq 2} \{s_3\}) &= Pr_{v,D_\alpha}^\omega(\text{Cyl}(\{s_0 \xrightarrow{\alpha_0} s_2 \xrightarrow{\gamma} s_3\})) = \frac{1}{2} \quad \text{and} \\ Pr_{v,D_\beta}^\omega(\diamond^{\leq 2} \{s_3\}) &= Pr_{v,D_\beta}^\omega(\text{Cyl}(\{s_0 \xrightarrow{\alpha_0} s_2 \xrightarrow{\gamma} s_3, s_0 \xrightarrow{\alpha} s_1 \xrightarrow{\beta} s_3\})) = \frac{1}{2} + \frac{1}{2} \cdot \frac{1}{4} = \frac{5}{8}. \quad \diamond \end{aligned}$$

3.3.2 Continuous time Markov decision processes

The focus of this thesis is on the analysis of stochastic models which combine nondeterminism and exponentially distributed delays. More precisely, we strive to extend continuous-time Markov chains (cf. Sec. 3.2.2) with nondeterministic choices. Following the nomenclature in the discrete-time case, where nondeterministic extensions of DTMCs are referred to as MDPs (cf. Sec. 3.3.1), the corresponding continuous-time model is called a *continuous-time Markov decision process* (CTMDP) [Mil68b, Mil68a, Put94].

The behavior in a state of a CTMC is completely determined by the exponentially distributed delays of its outgoing transitions. This is not the case in a CTMDP, where transitions are labeled with both, a rate of an exponential distribution (as in CTMCs) and an action, which names a nondeterministic choice.

Intuitively, the behavior of a CTMDP is as follows: Upon entering a state, one of the actions that are available according to the state's outgoing transitions must be chosen nondeterministically. After that, the behavior in that state is governed by the exponentially

distributed delays of those transitions, that correspond to the chosen action. The definition of a CTMDP differs from that of an MDP in that the transition probability matrix is replaced by a *rate matrix* which specifies the transitions' delay time distribution:

Definition 3.11 (Continuous-time Markov decision process). A continuous-time Markov decision process (CTMDP) is a tuple $\mathcal{C} = (\mathcal{S}, Act, \mathbf{R}, \nu)$ where $\mathcal{S}$ and Act are finite, nonempty sets of states and actions, $\mathbf{R} : \mathcal{S} \times Act \times \mathcal{S} \rightarrow \mathbb{R}_{\geq 0}$ is a three-dimensional rate matrix and $\nu \in Distr(\mathcal{S})$ is an initial distribution.

If $\mathbf{R}(s, \alpha, s') = \lambda$ and $\lambda > 0$, an α -transition with *rate* λ leads from state s to state s' . λ is the rate of the negative exponential distribution which governs the transition's delay. Therefore, the α -transition executes in time interval $[a, b] \subseteq \mathbb{R}_{\geq 0}$ with probability $\eta_\lambda([a, b]) = \int_a^b \lambda e^{-\lambda t} dt = (e^{-\lambda a} - e^{-\lambda b})$. The function η_λ corresponds to the cumulative distribution function of the exponential distribution with rate λ . It extends to a probability measure on the Borel σ -field $\mathfrak{B}(\mathbb{R}_{\geq 0})$ in the standard way.

Similar to the semantics of MDPs, the actions of the transitions that leave a state $s \in \mathcal{S}$ of a CTMDP constitute the set of *enabled actions* in that state:

$$Act(s) = \{\alpha \in Act \mid \exists s' \in \mathcal{S}. \mathbf{R}(s, \alpha, s') > 0\}.$$

The *exit rate* of a state $s \in \mathcal{S}$ under action α is the sum of the rates of all α -transitions that leave that state; formally, $E(s, \alpha) = \sum_{s' \in \mathcal{S}} \mathbf{R}(s, \alpha, s')$. Note that in general, the exit rate of a state differs depending on the enabled action that is considered.

Upon entering state s , an action from the set $Act(s)$ is chosen nondeterministically, say α . The exit rate of state s under action α determines its sojourn time: By choosing α , all transitions that are labeled with actions $\beta \neq \alpha$ get blocked. The subsequent behavior in state s equals that of a CTMC (cf. Sec. 3.2.2): The remaining α -transitions compete in a race, which is won by the α -transition whose randomly drawn delay expires first. Hence, the sojourn time in state s is governed by the minimum of the exponentially distributed delays of all outgoing α -transitions. The random variable that describes the minimum of exponential distributions is again exponentially distributed, namely with the sum $E(s, \alpha)$ of the rates of the competing α -transitions.

At the same time, the probability to move to a given α -successor state s' of s is also determined by the outcome of the race: It corresponds to the event that an α -transition which leads to state s' executes first. When leaving state s with action α , the probability to jump to a successor state s' is denoted $\mathbf{P}(s, \alpha, s')$, where $\mathbf{P} : \mathcal{S} \times Act \times \mathcal{S} \rightarrow [0, 1]$ is the three-dimensional transition probability matrix defined by $\mathbf{P}(s, \alpha, s') = \frac{\mathbf{R}(s, \alpha, s')}{E(s, \alpha)}$ if $E(s, \alpha) > 0$ and $\mathbf{P}(s, \alpha, s') = 0$, otherwise. In this way, each CTMDP $(\mathcal{S}, Act, \mathbf{R}, \nu)$ induces the *embedded MDP* $(\mathcal{S}, Act, \mathbf{P}, \nu)$, which abstracts from the CTMDP's timed behaviors and only considers its branching probabilities.

Similar to MDPs, we assume that $Act(s) \neq \emptyset$ for all states $s \in \mathcal{S}$ of a CTMDP. In this way, we avoid deadlock states which complicate the definition of the underlying stochastic process. Note that for our purposes (i.e. for timed reachability analysis and CSL model checking), this is no restriction as all deadlock states $s \in \mathcal{S}$ can easily be equipped with a self-loop of the form $\mathbf{R}(s, \alpha, s) = 1$ for some arbitrary $\alpha \in Act$. As we assume that a deadlock state is never left, this yields an equivalent CTMDP that satisfies our requirement.

Example 3.4. When entering state s_1 of the CTMDP in Fig. 3.5, one action from the set of enabled actions $Act(s_1) = \{\alpha, \beta\}$ is chosen nondeterministically, say α . Next, the rate of the α -transition determines its exponentially distributed delay. Hence for a single α -transition, the probability to go from s_1 to s_3 within time t is $1 - e^{-\mathbf{R}(s_1, \alpha, s_3)t} = 1 - e^{-0.1t}$.

In Fig. 3.5 a race occurs in state s_1 if action β is chosen: Two β -transitions (to states s_2 and s_3) with rates $\mathbf{R}(s_1, \beta, s_2) = 15$ and $\mathbf{R}(s_1, \beta, s_3) = 5$ become available and state s_1 is left as soon as the first transition executes. The sojourn time in state s_1 is exponentially distributed with rate $E(s_1, \beta) = \mathbf{R}(s_1, \beta, s_2) + \mathbf{R}(s_1, \beta, s_3) = 20$. The probability $\mathbf{P}(s_1, \beta, s_2)$ to move to state s_2 is $\mathbf{R}(s_1, \beta, s_2)/E(s_1, \beta) = 0.75$. $\diamond$

We call a CTMDP *deterministic* iff $|Act(s)| = 1$ for all states $s \in \mathcal{S}$. In this case, no nondeterministic choices exist and the CTMDP corresponds to a CTMC. Reversely, any CTMC corresponds to a deterministic CTMDP. Therefore, CTMDPs are a conservative extension of CTMCs.

The measurable space

To measure the probability of events in a CTMDP, we use paths to represent a single outcome of the associated random experiment. Opposed to the paths for MDPs that were defined in Sec. 3.3.1, the *timed paths* of a CTMDP also capture the sojourn times in each state. In this way, a timed path describes the complete trajectory of the CTMDP:

Definition 3.12 (Timed paths). Let $\mathcal{C} = (\mathcal{S}, Act, \mathbf{R}, \nu)$ be a CTMDP. $Paths^n(\mathcal{C}) = \mathcal{S} \times (Act \times \mathbb{R}_{\geq 0} \times \mathcal{S})^n$ is the set of paths of length n in $\mathcal{C}$; the set of finite paths in $\mathcal{C}$ is defined as $Paths^*(\mathcal{C}) = \bigcup_{n \in \mathbb{N}} Paths^n$, and $Paths^\omega(\mathcal{C}) = (\mathcal{S} \times Act \times \mathbb{R}_{\geq 0})^\omega$ is the set of infinite paths in $\mathcal{C}$. Accordingly, $Paths(\mathcal{C}) = Paths^*(\mathcal{C}) \cup Paths^\omega(\mathcal{C})$ denotes the set of all paths in $\mathcal{C}$.

We write $Paths$ instead of $Paths(\mathcal{C})$ whenever $\mathcal{C}$ is clear from the context. Moreover, if no ambiguity arises, we refer to the time-abstract paths in MDPs and the timed paths in CTMDPs simply as *paths*.

A single timed path is denoted $\pi = s_0 \xrightarrow{\alpha_0, t_0} s_1 \xrightarrow{\alpha_1, t_1} \dots \xrightarrow{\alpha_{n-1}, t_{n-1}} s_n$ where $|\pi| = n$ is the length of π and $\pi \downarrow = s_n$ is the last state of π . We use $abs(\pi) = s_0 \xrightarrow{\alpha_0} s_1 \xrightarrow{\alpha_2} \dots \xrightarrow{\alpha_{n-1}} s_n$ to refer to the *time-abstract* path induced by π .

For $k \leq |\pi|$, $\pi[k] = s_k$ is the $(k+1)$ -th state on π ; if $k < |\pi|$, $\delta(\pi, k) = t_k$ is the time spent in state s_k . If $i < j \leq |\pi|$ then $\pi[i..j]$ denotes the path-infix $s_i \xrightarrow{\alpha_i, t_i} s_{i+1} \xrightarrow{\alpha_{i+1}, t_{i+1}} \dots \xrightarrow{\alpha_{j-1}, t_{j-1}} s_j$

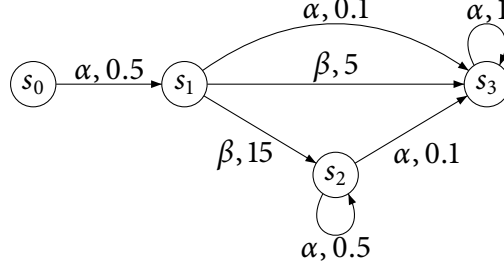


Figure 3.5: Example of a CTMDP.

of π . Finally, for infinite path π , we use $\pi@t$ to denote the state that is occupied on π at time point $t \in \mathbb{R}_{\geq 0}$. Formally, $\pi@t = \pi[k]$ where $k \in \mathbb{N}$ is the smallest index such that $\sum_{i=0}^k t_i > t$. If no such k exists, $\pi@t$ is undefined.

Note that Def. 3.12 does not impose any semantic restrictions on paths. In particular, the set *Paths* may contain paths which do not comply with the rate matrix of the underlying CTMDP. However, the definition of the probability measure (cf. Def. 3.15 on page 80) justifies this, as it assigns probability zero to such sets of paths.

To define the probability space that is induced by a CTMDP and a scheduler, we rely on the measure theoretic results from Chapter 2.

Our goal is to measure the probability of (measurable) sets of paths. Therefore, we first define a σ -field of sets of *combined transitions* which we later use to define σ -fields of sets of finite and infinite paths. The concept of a combined transition goes back to [WJ06, Joh07]. Informally, a combined transition is a tuple (α, t, s') which entangles the decision for action α with the time-point t at which the CTMDP moves to successor state s' . Formally, for a CTMDP $\mathcal{C} = (\mathcal{S}, Act, \mathbf{R}, \nu)$, let $\Omega = Act \times \mathbb{R}_{\geq 0} \times \mathcal{S}$ be the set of combined transitions in $\mathcal{C}$. To define a probability space on Ω , note that $\mathcal{S}$ and Act are finite; hence, the corresponding σ -fields are defined as $\mathfrak{F}_{Act} = 2^{Act}$ and $\mathfrak{F}_{\mathcal{S}} = 2^{\mathcal{S}}$. Any combined transition occurs at some time point $t \in \mathbb{R}_{\geq 0}$, so that we can use the Borel σ -field $\mathfrak{B}(\mathbb{R}_{\geq 0})$ to measure the corresponding subsets of $\mathbb{R}_{\geq 0}$.

Any path $\pi = s_0 \xrightarrow{\alpha_0, t_0} s_1 \xrightarrow{\alpha_1, t_1} \dots \xrightarrow{\alpha_{n-1}, t_{n-1}} s_n$ of length n can be extended by a combined transition $m = (\alpha_n, t_n, s_{n+1})$ to a path of length $n + 1$. This extension is denoted $\pi \circ m$. Hence, any path can be regarded as an initial state and a (finite or infinite) concatenation of combined transitions from the set Ω . Obviously, this is closely linked to the definition of product σ -fields which are discussed in detail in Sec. 2.5.

Recall that a Cartesian product is a *measurable rectangle* if its constituent sets are elements of their respective σ -fields. For example, in our case the set $A \times T \times S'$ is a measurable rectangle if $A \in \mathfrak{F}_{Act}$, $T \in \mathfrak{B}(\mathbb{R}_{\geq 0})$ and $S' \in \mathfrak{F}_{\mathcal{S}}$. We use $\mathfrak{F}_{Act} \otimes \mathfrak{B}(\mathbb{R}_{\geq 0}) \otimes \mathfrak{F}_{\mathcal{S}}$ to denote the set of all measurable rectangles². It generates the desired σ -field $\mathfrak{F}$ of sets of combined transitions, i.e. $\mathfrak{F} = \sigma(\mathfrak{F}_{Act} \otimes \mathfrak{B}(\mathbb{R}_{\geq 0}) \otimes \mathfrak{F}_{\mathcal{S}})$.

Now $\mathfrak{F}$ may be used to infer the σ -fields $\mathfrak{F}_{Paths^n}$ of sets of paths of length n : $\mathfrak{F}_{Paths^n}$ is

²Recall our notation: $\mathfrak{F}_{Act} \otimes \mathfrak{B}(\mathbb{R}_{\geq 0}) \otimes \mathfrak{F}_{\mathcal{S}}$ is not a Cartesian product itself; instead, it is the set of all Cartesian products. For details, see Def. 2.16 on page 42.

generated by the set of measurable (path) rectangles, that is

$$\mathfrak{F}_{Paths^n} = \sigma(\{S_0 \times M_1 \times \dots \times M_n \mid S_0 \in \mathfrak{F}_S, M_i \in \mathfrak{F}, 1 \leq i \leq n\}).$$

Intuitively, $\mathfrak{F}_{Paths^n}$ consists of all possible (even countably infinite) unions and intersections of measurable path rectangles of length n .

Example 3.5. For the CTMDP in Fig. 3.5, the event “from s_1 we directly reach state s_3 within 0.5 time units” and the event “action α is chosen in state s_1 and we remain in s_1 for less than 0.2 or more than 1 time units” are described by the Cartesian products $\Pi_1 = \{s_1\} \times Act \times [0, 0.5] \times \{s_3\}$ and $\Pi_2 = \{s_1\} \times \{\alpha\} \times ([0, 0.2) \cup (1, \infty)) \times S$. Π_1 and Π_2 are measurable rectangles whereas their union $\Pi_1 \cup \Pi_2$ is an element of the σ -field $\mathfrak{F}_{Paths^1}$. $\diamond$

The σ -field of sets of infinite paths is obtained by applying the cylinder set construction which is discussed in detail in Sec. 2.5.4: A set C^n of paths of length n is called a *cylinder base*; it induces the infinite *cylinder* $C_n = \{\pi \in Paths^\omega \mid \pi[0..n] \in C^n\}$. A cylinder C_n is *measurable* if $C^n \in \mathfrak{F}_{Paths^n}$; C_n is called an *infinite rectangle* if $C^n = S_0 \times A_0 \times T_0 \times \dots \times A_{n-1} \times T_{n-1} \times S_n$ and $S_i \subseteq S$, $A_i \subseteq Act$ and $T_i \subseteq \mathbb{R}_{\geq 0}$. It is a *measurable infinite rectangle*, if $S_i \in \mathfrak{F}_S$, $A_i \in \mathfrak{F}_{Act}$ and $T_i \in \mathfrak{B}(\mathbb{R}_{\geq 0})$. We obtain the desired σ -field of sets of infinite paths as the minimal σ -field generated by the set of measurable cylinders; formally, $\mathfrak{F}_{Paths^\omega} = \sigma(\bigcup_{n=0}^{\infty} \{C_n \mid C^n \in \mathfrak{F}_{Paths^n}\})$. Finally, the σ -field $\mathfrak{F}_{Paths^*}$ over finite and infinite paths is the smallest σ -field generated by the disjoint union $\bigcup_{n=0}^{\infty} \mathfrak{F}_{Paths^n} \sqcup \mathfrak{F}_{Paths^\omega}$.

The probability measure

As for MDPs, we use schedulers to define the semantics for CTMDPs. More precisely, a CTMDP and a scheduler induce a unique probability measure on the measurable spaces that we have defined above.

A scheduler quantifies the probability of the next action based on the history of the system: If state s is reached via finite path π , the scheduler yields a probability distribution over $Act(\pi \downarrow)$. The class of *measurable schedulers* that we use here has been defined in [WJ06, Joh07]. A measurable scheduler can incorporate the complete information from the history π that led into the current state when making its decision. In particular, it may yield different decisions depending on the time that has passed on π or in single states on π .

In fact, there exists a plethora of scheduler classes which differ both in the information they can base their decision on as well as on the time, their decision is due. A detailed discussion of this topic follows in Chapter 4. For now, we do not go into those subtle details and stick to the general definition of measurable schedulers:

Definition 3.13 (Measurable scheduler). Let $\mathcal{C} = (S, Act, \mathbf{R}, v)$ be a CTMDP. A mapping $D : Paths^* \times \mathfrak{F}_{Act} \rightarrow [0, 1]$ is a measurable scheduler iff $D(\pi, \cdot) \in Distr(Act(\pi \downarrow))$ for all $\pi \in Paths^*$ and the functions $D(\cdot, A) : Paths^* \rightarrow [0, 1]$ are measurable for all $A \in \mathfrak{F}_{Act}$. We use GM to denote the set of all measurable schedulers.

In Def. 3.13, the measurability condition states that for any measurable set of probabilities $B \in \mathfrak{B}([0, 1])$ and any set of actions $A \in \mathfrak{F}_{Act}$, the set $\{\pi \in Paths^* \mid D(\pi, A) \in B\}$ belongs to $\mathfrak{F}_{Paths^*}$ (for details, we refer to [WJ06]).

Similar to the MDP definition, the support restriction $D(\pi, \cdot) \in Distr(Act(\pi \downarrow))$ states that whenever $D(\pi)(\alpha) > 0$, the action α is enabled in state $\pi \downarrow$. This prevents a measurable scheduler to choose actions that are not available in the current state.

Note that we can equivalently specify any GM-scheduler $D : Paths^* \times \mathfrak{F}_{Act} \rightarrow [0, 1]$ as a mapping $D' : Paths^* \rightarrow Distr(Act)$ by setting $D'(\pi)(A) = D(\pi, A)$ for all $\pi \in Paths^*$ and $A \in \mathfrak{F}_{Act}$; to further simplify notation, we also use $D(\pi, \cdot)$ to refer to this distribution.

To derive a probability measure on $\mathfrak{F}_{Paths^\omega}$, we first define a probability measure on combined transitions, i.e. on the measurable space $(\Omega, \mathfrak{F})$:

Definition 3.14 (Probability on combined transitions). Let $\mathcal{C} = (\mathcal{S}, Act, \mathbf{R}, \nu)$ be a CT-MDP and D a GM-scheduler on $\mathcal{C}$. For all $\pi \in Paths^*(\mathcal{C})$, we define the probability measure $\mu_D(\pi, \cdot) : \mathfrak{F} \rightarrow [0, 1]$ such that

$$\mu_D(\pi, M) = \int_{Act} D(\pi, d\alpha) \int_{\mathbb{R}_{\geq 0}} \eta_{E(\pi \downarrow, \alpha)}(dt) \int_{\mathcal{S}} \mathbf{I}_M(\alpha, t, s') \mathbf{P}(\pi \downarrow, \alpha, ds'). \quad (3.11)$$

Here, we use $\mathbf{I}_M(\alpha, t, s)$ to denote the indicator for the set $M \subseteq \Omega$, that is, $\mathbf{I}_M(\alpha, t, s) = 1$ if the combined transition $(\alpha, t, s) \in M$ and $\mathbf{I}_M(\alpha, t, s) = 0$, otherwise. Intuitively, for a given finite path π and a set M of combined transitions, $\mu_D(\pi, M)$ is the probability to continue from $\pi \downarrow$ by one of the combined transitions in M . For a measurable rectangle $A \times T \times S' \in \mathfrak{F}$ and time interval T , we obtain

$$\mu_D(\pi, A \times T \times S') = \sum_{\alpha \in A} D(\pi, \{\alpha\}) \cdot \mathbf{P}(\pi \downarrow, \alpha, S') \cdot \int_T E(\pi \downarrow, \alpha) \cdot e^{-E(\pi \downarrow, \alpha)t} dt, \quad (3.12)$$

which is the probability to leave $\pi \downarrow$ via some action from the set A within time interval T to a state in S' .

Lemma 3.5. For any $\pi \in Paths^*$, the function $\mu_D(\pi, \cdot) : \mathfrak{F} \rightarrow [0, 1]$ is a probability measure on $(\Omega, \mathfrak{F})$.

Proof. This follows from [ADD00, Theorem 2.6.7], for $D(\pi, \cdot)$ is a probability measure and all $\eta_{E(\pi \downarrow, \alpha)}$ as well as $\mathbf{P}(\pi \downarrow, \alpha, \cdot)$ are probability measures for $\alpha \in Act(\pi \downarrow)$. $\square$

To extend this to a probability measure on $\mathfrak{F}_{Paths^\omega}$, we assume an *initial distribution* $\nu \in Distr(\mathcal{S})$ for the probability to start in a certain state s and inductively append sets of combined transitions.

As the probability measures in Def. 3.15 (see below) depend on the Lebesgue integral of a function involving the measure μ_D , we have to show that $\mu_D : \text{Paths}^* \times \mathfrak{F} \rightarrow [0, 1]$ is measurable in its first argument, i.e. that for all $M \in \mathfrak{F}$ and $B \in \mathfrak{B}([0, 1])$ it is the case that $\mu_D(\cdot, M)^{-1}(B) \in \mathfrak{F}_{\text{Paths}^*}$. The following theorem stems from [WJ06] and is restated here only for the sake of completeness:

Theorem 3.4 (Combined transition measurability [WJ06, Theorem 1]). *Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ be a CTMDP and D a GM-scheduler. For all $A \in \mathfrak{F}_{\text{Act}}$ it holds: $D(\cdot, A) : \text{Paths}^* \rightarrow [0, 1]$ is measurable iff $\forall M \in \mathfrak{F}, \mu_D(\cdot, M) : \text{Paths}^* \rightarrow [0, 1]$ is measurable.*

Hence $\mu_D : \text{Paths}^* \times \mathfrak{F} \rightarrow [0, 1]$ is measurable in its first argument whenever D is a GM-scheduler. Note also, that the restriction $\mu_D : \text{Paths}^n \times \mathfrak{F} \rightarrow [0, 1]$ is measurable with respect to $\mathfrak{F}_{\text{Paths}^n}$. With these preconditions, we can define the probability measure on sets of finite paths as follows:

Definition 3.15 (Probability measure). *Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ be a CTMDP. The probability measure on $(\text{Paths}^n, \mathfrak{F}_{\text{Paths}^n})$ is defined inductively as follows:*

$$\begin{aligned} Pr_{\nu, D}^0 : \mathfrak{F}_{\text{Paths}^0} &\rightarrow [0, 1] : \Pi \mapsto \sum_{s \in \Pi} \nu(\{s\}) \quad \text{and} \\ Pr_{\nu, D}^{n+1} : \mathfrak{F}_{\text{Paths}^{n+1}} &\rightarrow [0, 1] : \Pi \mapsto \int_{\text{Paths}^n} Pr_{\nu, D}^n(d\pi) \int_{\Omega} \mathbf{I}_{\Pi}(\pi \circ m) \mu_D(\pi, dm). \end{aligned}$$

Informally, Def. 3.15 derives the probability measure $Pr_{\nu, D}^{n+1}$ on sets of paths Π of length $n+1$ by multiplying the probability $Pr_{\nu, D}^n(d\pi)$ of a path π of length n with the probability $\mu_D(\pi, dm)$ of a combined transition m such that the concatenation $\pi \circ m$ is a path from the set Π .

One further remark is in order here: Formally, we have not yet proved that the nested integral in the definition of $Pr_{\nu, D}^{n+1}$ yields a measurable function with respect to $\mathfrak{F}_{\text{Paths}^n}$. To bridge this gap, we first show that the functions

$$f_{\Pi} : \text{Paths}^{n-1} \rightarrow [0, 1] : \pi \mapsto \int_{\Omega} \mathbf{I}_{\Pi}(\pi \circ m) \mu_D(\pi, dm)$$

are measurable for all $\Pi \in \mathfrak{F}_{\text{Paths}^n}$. To see this, first note that $\{m \in \Omega \mid \pi \circ m \in \Pi\} \in \mathfrak{F}$ for all $\pi \in \text{Paths}^{n-1}$: If $\Pi = S_0 \times M_0 \times \cdots \times M_{n-1}$ is a measurable rectangle such that $M_i \in \mathfrak{F}$ for $0 \leq i < n$, we obtain

$$\{m \in \Omega \mid \pi \circ m \in \Pi\} = \begin{cases} M_{n-1} & \text{if } \pi \in S_0 \times M_0 \times \cdots \times M_{n-2} \\ \emptyset & \text{otherwise.} \end{cases}$$

Hence, for measurable rectangle Π , the set $\{m \in \Omega \mid \pi \circ m \in \Pi\}$ is measurable.

Now, let $\Pi = \Pi_1 \cup \Pi_2$ and $M_i = \{m \in \Omega \mid \pi \circ m \in \Pi_i\}$ for $i = 1, 2$. By the induction hypothesis, $M_i \in \mathfrak{F}$; further, $\{m \in \Omega \mid \pi \circ m \in \Pi\} = M_1 \cup M_2$. As $\mathfrak{F}$ is closed under countable union, $M_1 \cup M_2 \in \mathfrak{F}$. For the complement Π^c , define $M = \{m \in \Omega \mid \pi \circ m \in \Pi\}$. By the induction hypothesis, $M \in \mathfrak{F}$. Further observe that $\{m \in \Omega \mid \pi \circ m \in \Pi^c\} = \{m \in \Omega \mid \pi \circ m \notin \Pi\} = \{m \in \Omega \mid \pi \circ m \in \Pi\}^c = M^c$. Then $M^c \in \mathfrak{F}$ follows since $M \in \mathfrak{F}$ and $\mathfrak{F}$ is closed under complement. Now the functions f_Π can be restated as follows:

$$f_\Pi : \text{Paths}^{n-1} \rightarrow [0, 1] : \pi \mapsto \mu_D(\pi, \{m \in \Omega \mid \pi \circ m \in \Pi\})$$

which is measurable with respect to $\mathfrak{F}_{\text{Paths}^{n-1}}$ by Theorem 3.4, where μ_D is restricted to Paths^{n-1} .

By Def. 3.15, we obtain measures on all σ -fields $\mathfrak{F}_{\text{Paths}^n}$ of subsets of paths of length n . This extends to a measure on $(\text{Paths}^\omega, \mathfrak{F}_{\text{Paths}^\omega})$ as follows: First, note that any measurable cylinder can be represented by a base of finite length, i.e. $B_n = \{\pi \in \text{Paths}^\omega \mid \pi[0..n] \in B^n\}$. Now the measures $Pr_{v,D}^n$ on $\mathfrak{F}_{\text{Paths}^n}$ extend to a unique probability measure $Pr_{v,D}^\omega$ on $\mathfrak{F}_{\text{Paths}^\omega}$ by defining $Pr_{v,D}^\omega(B_n) = Pr_{v,D}^n(B^n)$. Although any measurable rectangle with base B^m can equally be represented by a higher-dimensional base (more precisely, if $m < n$ and $B^n = B^m \times \Omega^{n-m}$ then $B_n = B_m$), the Ionescu-Tulcea extension theorem (Thm. 2.19 on page 51) is applicable due to the inductive definition of the measures $Pr_{v,D}^n$ and assures the extension to be well defined and unique.

One important property is still missing: We have not proved yet, that the functions $Pr_{v,D}^\omega$ are indeed probability measures. The next lemma makes up for that:

Lemma 3.6. *$Pr_{v,D}^n$ is a probability measure on $(\text{Paths}^n, \mathfrak{F}_{\text{Paths}^n})$ for all $n \in \mathbb{N}$.*

Proof. By induction on n . v is a probability measure on $(\mathcal{S}, \mathfrak{F}_\mathcal{S})$ and so is $Pr_{v,D}^0$. In the induction step, $n > 0$ and

$$Pr_{v,D}^n(\Pi) = \int_{\text{Paths}^{n-1}} Pr_{v,D}^{n-1}(d\pi) \int_{\Omega} \mathbf{I}_\Pi(\pi \circ m) \mu_D(\pi, dm).$$

By the induction hypothesis, $Pr_{v,D}^{n-1}$ is a probability measure; the same holds for $\mu_D(\pi, \cdot)$ by Lemma 3.5. As the product yields a probability measure again (see Thm. 2.16 on page 46 or [ADD00, 2.6.2]), the claim follows. $\square$

Definition 3.15 inductively *appends* transition triples to the path prefixes of length n to obtain a measure on sets of paths of length $n+1$. In some of our proofs, we make use of the fact that paths can also be constructed reversely: More specifically, we will later need to split a set of paths into a set of prefixes I and a set of suffixes Π . Thus we define the set of path prefixes of length $k > 0$ as $PPref^k = (\mathfrak{F}_\mathcal{S} \times \mathfrak{F}_{Act} \times \mathfrak{B}(\mathbb{R}_{\geq 0}))^k$ and provide a probability measure on its σ -field $\mathfrak{F}_{PPref^k}$:

Definition 3.16 (Prefix measure). Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, v)$ be a CTMDP and D a GM-scheduler on $\mathcal{C}$. For $I \in \mathfrak{F}_{\text{PPref}^k}$ and $k > 0$, define

$$\mu_{v,D}^k(I) = \int_{\text{Paths}^{k-1}} Pr_{v,D}^{k-1}(d\pi) \int_{\text{Act}} D(\pi, d\alpha) \int_{\mathbb{R}_{\geq 0}} \mathbf{I}_I(\pi \xrightarrow{\alpha,t}) \eta_{E(\pi\downarrow, \alpha)}(dt).$$

As $Pr_{v,D}^{k-1}$ is a probability measure, so is $\mu_{v,D}^k$. If $I \in \mathfrak{F}_{\text{PPref}^k}$ and $\Pi \in \mathfrak{F}_{\text{Paths}^n}$, their concatenation is the set $I \times \Pi \in \mathfrak{F}_{\text{Paths}^{k+n}}$; its probability $Pr_{v,D}^{k+n}(I \times \Pi)$ is obtained by multiplying the measure of prefixes $i \in I$ with the suffixes in Π :

Lemma 3.7. Let $\Pi \in \mathfrak{F}_{\text{Paths}^n}$ and $I \in \mathfrak{F}_{\text{PPref}^k}$. If $i = s_0 \xrightarrow{\alpha_0, t_0} \dots \xrightarrow{\alpha_{k-2}, t_{k-2}} s_{k-1} \xrightarrow{\alpha_{k-1}, t_{k-1}}$ is a path prefix from I , define $v_i = \mathbf{P}(s_{k-1}, \alpha_{k-1}, \cdot)$ and $D_i(\pi, \cdot) = D(i \circ \pi, \cdot)$. Then

$$Pr_{v,D}^{k+n}(I \times \Pi) = \int_{\text{PPref}^k} \mu_{v,D}^k(di) \int_{\text{Paths}^n} \mathbf{I}_{I \times \Pi}(i \circ \pi) Pr_{v_i, D_i}^n(d\pi). \quad (3.13)$$

Proof. By induction on n : Let $\Pi \in \mathfrak{F}_{\text{Paths}^0}$, i.e. $\Pi \subseteq \mathcal{S}$.

$$\begin{aligned} Pr_{v,D}^k(I \times \Pi) &= \int_{\text{Paths}^{k-1}} Pr_{v,D}^{k-1}(d\pi) \int_{\Omega} \mathbf{I}_{I \times \Pi}(\pi \circ m) \mu_D(\pi, dm) \\ &= \int_{\text{Paths}^{k-1}} Pr_{v,D}^{k-1}(d\pi) \int_{\text{Act}} D(\pi, d\alpha) \int_{\mathbb{R}_{\geq 0}} \eta_{E(\pi\downarrow, \alpha)}(dt) \int_{\mathcal{S}} \mathbf{I}_{I \times \Pi}(\pi \xrightarrow{\alpha,t} s') \mathbf{P}(\pi\downarrow, \alpha, ds') \\ &= \int_{(\mathcal{S} \times \text{Act} \times \mathbb{R}_{\geq 0})^k} \mu_{v,D}^k(d(\pi \xrightarrow{\alpha,t})) \int_{\mathcal{S}} \mathbf{I}_{I \times \Pi}(\pi \xrightarrow{\alpha,t} s') \mathbf{P}(\pi\downarrow, \alpha, ds') \\ &= \int_{(\mathcal{S} \times \text{Act} \times \mathbb{R}_{\geq 0})^k} \mu_{v,D}^k(d(\pi \xrightarrow{\alpha,t})) \int_{\text{Paths}^0} \mathbf{I}_{I \times \Pi}(\pi \xrightarrow{\alpha,t} s') Pr_{v_i, D_i}^0(ds'). \end{aligned}$$

In the induction step ($n \rightsquigarrow n+1$), we assume as induction hypothesis that (3.13) holds for n and prove its validity for $n+1$:

$$\begin{aligned} Pr_{v,D}^{k+n+1}(I \times \Pi) &= \int_{\text{Paths}^{k+n}} Pr_{v,D}^{k+n}(d\pi) \int_{\Omega} \mathbf{I}_{I \times \Pi}(\pi \circ m) \mu_D(\pi, dm) \\ &= \int_{\text{Paths}^{k+n}} Pr_{v,D}^{k+n}(d(i \circ \pi')) \int_{\Omega} \mathbf{I}_{I \times \Pi}(i \circ \pi' \circ m) \mu_D(i \circ \pi', dm) \\ &\stackrel{\text{i.h.}}{=} \int_{(\mathcal{S} \times \text{Act} \times \mathbb{R}_{\geq 0})^k} \mu_{v,D}^k(di) \int_{\text{Paths}^n} Pr_{v_i, D_i}^n(d\pi') \int_{\Omega} \mathbf{I}_{I \times \Pi}(i \circ \pi' \circ m) \mu_D(i \circ \pi', dm) \\ &= \int_{(\mathcal{S} \times \text{Act} \times \mathbb{R}_{\geq 0})^k} \mu_{v,D}^k(di) \int_{\text{Paths}^n} Pr_{v_i, D_i}^n(d\pi') \int_{\Omega} \mathbf{I}_{I \times \Pi}(i \circ \pi' \circ m) \mu_{D_i}(\pi', dm) \\ &= \int_{(\mathcal{S} \times \text{Act} \times \mathbb{R}_{\geq 0})^k} \mu_{v,D}^k(di) \int_{\text{Paths}^{n+1}} \mathbf{I}_{I \times \Pi}(i \circ \pi) Pr_{v_i, D_i}^{n+1}(d\pi). \quad \square \end{aligned}$$

Lemma 3.7 justifies to split sets of paths and to measure the components of the resulting Cartesian product; therefore, it abstracts from the inductive definition of $Pr_{v,D}^n$.

A class of pathological paths that are not ruled out by Def. 3.12 are infinite paths whose duration converges to some real constant, i.e. paths that visit infinitely many states in a finite amount of time. For $n = 0, 1, 2, \dots$, an increasing sequence $r_n \in \mathbb{R}_{\geq 0}$ is *Zeno* if it converges to a positive real number.

Example 3.6. The sequence $r_n = \sum_{i=0}^n \frac{1}{2^i}$, $n \in \mathbb{N}$ is *Zeno*, as it converges to 2. $\diamond$

In the remainder of this thesis, we rule out Zeno behaviors. To justify this, let us prove that the probability of a set of paths with Zeno behaviors has probability 0. To prepare for this proof, the next lemma states that the probability that after a certain number of steps, the sojourn time is always less than 1 time unit, is 0:

Lemma 3.8. Let $k \in \mathbb{N}$ and $B = \mathcal{S} \times \Omega^k \times (\text{Act} \times [0, 1] \times \mathcal{S})^\omega$; then $Pr_{v,D}^\omega(B) = 0$.

Proof. The proof goes along the lines of [BHHK03, Prop. 1]:

As $\mathcal{S}$ and Act are finite, we can define $\lambda = \max \{E(s, \alpha) \mid s \in \mathcal{S}, \alpha \in \text{Act}\}$. For $n \geq 0$, let $B^n = \mathcal{S} \times \Omega^k \times (\text{Act} \times [0, 1] \times \mathcal{S})^n$ be a measurable base and B_n the induced infinite measurable rectangle. By induction on n , we show that $Pr_{v,D}^\omega(B_n) \leq (1 - e^{-\lambda})^n$:

1. In the induction base, let $n = 0$. Then $Pr_{v,D}^\omega(B_0) = Pr_{v,D}^k(\mathcal{S} \times \Omega^k) = 1 = (1 - e^{-\lambda})^0$.
2. As induction hypothesis, let $Pr_{v,D}^\omega(B_n) \leq (1 - e^{-\lambda})^n$. For B_{n+1} we obtain:

$$\begin{aligned}
 Pr_{v,D}^\omega(B_{n+1}) &= Pr_{v,D}^{n+k+1}(B^n \times \text{Act} \times [0, 1] \times \mathcal{S}) \\
 &= \int_{B^n} \mu_D(\pi, \text{Act} \times [0, 1] \times \mathcal{S}) Pr_{v,D}^{n+k}(d\pi) \\
 &= \int_{B^n} \left(\sum_{\alpha \in \text{Act}} D(\pi, \{\alpha\}) \cdot P(\pi \downarrow, \alpha, \mathcal{S}) \cdot \int_{[0,1]} E(\pi \downarrow, \alpha) e^{-E(\pi \downarrow, \alpha)t} dt \right) Pr_{v,D}^{n+k}(d\pi) \\
 &= \int_{B^n} \sum_{\alpha \in \text{Act}} D(\pi, \{\alpha\}) \cdot P(\pi \downarrow, \alpha, \mathcal{S}) \cdot (1 - e^{-E(\pi \downarrow, \alpha)}) Pr_{v,D}^{n+k}(d\pi) \\
 &\leq (1 - e^{-\lambda}) \cdot \underbrace{\int_{B^n} \sum_{\alpha \in \text{Act}} D(\pi, \{\alpha\}) \cdot P(\pi \downarrow, \alpha, \mathcal{S}) Pr_{v,D}^{n+k}(d\pi)}_{\leq 1} \\
 &\leq (1 - e^{-\lambda}) \cdot \int_{B^n} Pr_{v,D}^{n+k}(d\pi) = (1 - e^{-\lambda}) \cdot Pr_{v,D}^{n+k}(B^n) \\
 &= (1 - e^{-\lambda}) \cdot Pr_{v,D}^\omega(B_n) \leq (1 - e^{-\lambda})^{n+1}.
 \end{aligned}$$

Now $B_0 \supseteq B_1 \supseteq \dots$ and the B_n converge to B , i.e. $B_n \downarrow B$; hence $Pr_{v,D}^\omega(B_n) \rightarrow Pr_{v,D}^\omega(B)$ by Lemma 2.2 (cf. page 16). Further $\lim_{n \rightarrow \infty} Pr_{v,D}^\omega(B_n) \leq \lim_{n \rightarrow \infty} (1 - e^{-\lambda})^n = 0$. As $Pr_{v,D}^\omega$ is a measure (and hence nonnegative), it follows that $Pr_{v,D}^\omega(B) = 0$. $\square$

With this result we can prove the following theorem which justifies to generally rule out Zeno behavior:

Theorem 3.5 (Converging paths theorem). *The probability measure of the set of converging paths is zero.*

Proof. Let $\text{ConvPaths} = \{s_0 \xrightarrow{\alpha_0, t_0} s_1 \xrightarrow{\alpha_1, t_1} \dots \mid \sum_{i=0}^n t_i \text{ converges}\}$. For $\pi \in \text{ConvPaths}$, the sequence $\sum_{i=0}^{\infty} t_i$ converges; thus t_i converges to 0 and there exists $k \in \mathbb{N}$ such that $t_i \leq 1$ for all $i \geq k$. Hence $\text{ConvPaths} \subseteq \bigcup_{k=0}^{\infty} \mathcal{S} \times \Omega^k \times (\text{Act} \times [0, 1] \times \mathcal{S})^{\omega}$. By Lemma 3.8, $Pr_{v,D}^{\omega}(\mathcal{S} \times \Omega^k \times (\text{Act} \times [0, 1] \times \mathcal{S})^{\omega}) = 0$ for all $k \in \mathbb{N}$. Thus we obtain

$$Pr_{v,D}^{\omega}\left(\bigcup_{k=0}^{\infty} \mathcal{S} \times \Omega^k \times (\text{Act} \times [0, 1] \times \mathcal{S})^{\omega}\right) \leq \sum_{k=0}^{\infty} Pr_{v,D}^{\omega}(\mathcal{S} \times \Omega^k \times (\text{Act} \times [0, 1] \times \mathcal{S})^{\omega}) = 0.$$

But then ConvPaths is a subset of a set of measure zero; hence, on $\mathfrak{F}_{\text{Paths}^{\omega}}$ completed³ with respect to $Pr_{v,D}^{\omega}$ we obtain $Pr_{v,D}^{\omega}(\text{ConvPaths}) = 0$. $\square$

3.4 Conclusion

Markov chain theory is an extremely broad field in mathematics. In this chapter, we only discussed the preliminaries that are essential for the remainder of the thesis. More details about CTMCs and DTMCs can be found in the textbooks [KS76, Kul95]. More details about MDPs can be found in [Bel57, How71, Ber95] and in the textbook [Put94].

Compared to the other models presented in this chapter, CTMDPs have received less attention. As do the seminal papers of Miller [Mil68b, Mil68a], most of the results that are known for CTMDPs concentrate on optimizing reward-based measures such as the finite horizon expected state-based reward, the infinite horizon discounted state-based reward or the long run expected average reward. Details about the results that are known in mathematics can be found in [Put94] and in the survey paper [GHLPR06].

Lately, CTMDPs are considered in the field of game theory, where the model has become known as a continuous-time stochastic $1\frac{1}{2}$ player game. However, the results mostly concentrate on time-abstract schedulers [BFK⁺09]. The same holds for the results in [BHKH05], which are closely related to those of this thesis:

In [BHKH05], the authors provide an algorithm to optimize time-bounded reachability probabilities for time-abstract schedulers on a subclass of CTMDPs. This thesis extends these approaches in different respects. Most notably, we lift the restriction to certain subclasses of CTMDPs and consider strictly better time-dependent schedulers. These contributions are described in detail in the following chapters.

³We may assume $\mathfrak{F}_{\text{Paths}^{\omega}}$ to be complete, see Def. 2.4.

4 Schedulers in CTMDPs

Nothing is more difficult, and therefore more precious, than to be able to decide.

(Napoléon Bonaparte)

Schedulers in CTMDPs and other variants of randomly timed games can roughly be classified as to whether they use timing information or not. In the literature, the analysis of CTMDPs is mostly focused on determining optimal schedulers for criteria such as the expected total reward, the expected long-run average reward (cf. the survey [GHLPR06]) and unbounded reachability probabilities [Put94]. For such comparatively simple criteria, time-abstract schedulers suffice. Stated differently, providing the scheduler with information on the amount of time that has passed does not improve its decisions for such properties. When analyzing such criteria, it therefore suffices to either fully abstract from the timing information in the CTMDP or to abstract from it at least partly by transforming the CTMDP into an equivalent discrete-time MDP. The latter process is commonly referred to as uniformization [Put94, p. 562],[GHLPR06].

In comparison to the properties stated above, the focus of this thesis is mostly on time bounded reachability objectives such as the maximum probability to hit a given set of goal states during a finite time-interval. As we will see in this chapter, the maximum achievable probability of such events strongly depends on whether the underlying scheduler class uses timing information or not.

In the previous chapter, we have introduced the class of generic measurable schedulers. It is complete in a sense, as the corresponding *GM*-schedulers may use the complete information about the trajectory that led into the current state. For example, a *GM*-scheduler can access the state history and the sojourn time in each individual state of the history.

In this chapter, we investigate schedulers more closely and define a hierarchy of positional and history-dependent schedulers which refines the notion of measurable schedulers from Sec. 3.3.2. As it turns out, an important distinguishing criterion is the level of detail of timing information the schedulers may exploit, e.g. the delay in the last state, the total time that was spent during the trajectory that led into the current state, or all individual state residence times.

In general, the delay that has to pass in a state s before the CTMDP jumps to a successor state s' is determined by the action that is selected by the scheduler when entering state s' . In the second part of this chapter, we therefore investigate under which conditions this resolution of nondeterminism may be deferred: More precisely, we identify the subclass

of *locally uniform* CTMDPs and show how its schedulers delay their decision up to the point at which the current state s is left.

Rather than focusing on a specific objective, we consider this delayed nondeterminism for arbitrary measurable events. The core of our study is a transformation — called local uniformization — on CTMDPs which unifies the speed of outgoing transitions per state. Whereas classical uniformization [Gra91, GM84, Jen53] adds self-loops to achieve this, local uniformization uses auxiliary copy-states. In this way, we enforce that schedulers in the original and uniformized CTMDP have (for important scheduler classes) the same power, whereas classical loop-based uniformization permits a scheduler to change its decision when re-entering a state through the added self-loop.

Therefore, locally uniform CTMDPs permit to defer the resolution of nondeterminism, i.e., they dissolve the intrinsic dependency between state residence times and schedulers, and can be viewed as MDPs with exponentially distributed state residence times. This characterization provides the basis for Chapter 5, where we develop an approximation algorithm which computes time-bounded reachability probabilities in locally uniform CTMDPs.

Organization of this chapter. Section 4.1 proposes a hierarchy of scheduler classes and refines the notion of generic measurable schedulers from Sec. 3.3.2. In Sec. 4.2, we define local uniformization and prove its correctness. Section 4.3 summarizes the main results and Sec. 4.4 proves that deferring nondeterministic choices induces strictly tighter bounds on quantitative properties.

4.1 A hierarchy of scheduler classes

In Sec. 3.3.2, we have defined the probability of measurable sets of paths with respect to *GM*-schedulers. However, this does not fully describe a CTMDP, as a single scheduler represents only one way to resolve the CTMDP's nondeterministic choices. Therefore, instead of a single scheduler, we consider scheduler *classes* that group schedulers according to the information that they use for making a decision:

Given an event $\Pi \in \mathfrak{F}_{Paths^\omega}$, a scheduler class induces a set of probabilities — one for each scheduler in the respective class — which reflects the CTMDP's possible behaviors.

In this chapter, we propose a variety of scheduler classes (see the lattice depicted in Fig. 4.1) and investigate which of them preserve the minimum and maximum probabilities under local uniformization.

We start our discussion and recall the notion of *GM*-schedulers: As proved in [WJ06], they are the most general class definable on arbitrary CTMDPs. More precisely, the authors prove that all probability measures that conform to a CTMDP's set of valid paths are induced by some *GM*-scheduler. The intuition is as follows: If paths π_1 and π_2 end in state s , a *GM*-scheduler $D : Paths^* \times \mathfrak{F}_{Act} \rightarrow [0, 1]$ may yield different distributions $D(\pi_1, \cdot)$ and $D(\pi_2, \cdot)$ over the next action, depending on the entire histories π_1 and π_2 .

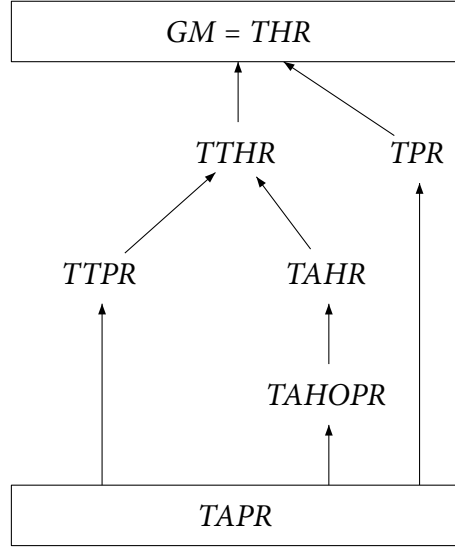


Figure 4.1: A hierarchy of scheduler classes.

Note that π_1 and π_2 contain the state sequence that was traversed, the sojourn time in each of those states and the action that was chosen to move from one state to another. Hence, we also refer to GM -schedulers as *time- and history-dependent randomized* schedulers.

On the contrary, a scheduler D is *time-abstract* and *positional* (a $TAPR$ -scheduler), if $D(\pi_1, \cdot) = D(\pi_2, \cdot)$ for all $\pi_1, \pi_2 \in Paths^*$ that end in the same state. As $D(\pi, \cdot)$ only depends on the current state, it can be specified as a mapping $D : \mathcal{S} \rightarrow Distr(Act)$.

Example 4.1. For $TAPR$ scheduler D with $D(s_0) = \{\alpha \mapsto 1\}$ and $D(s_1) = \{\beta \mapsto 1\}$, the induced stochastic process of the CTMDP in Fig. 4.2(a) is the CTMC depicted in Fig. 4.2(b). Note however, that in general, randomized schedulers do not yield CTMCs as the induced sojourn times are hyper-exponentially distributed. Hence, a continuous-time Markov decision process with an associated randomized scheduler is a slight misnomer, as a hyper-exponentially distributed sojourn time does not obey the Markov property, in general. However, this can safely be ignored, as we will see in the next chapters that considering deterministic schedulers (which obviously induce exponentially distributed sojourn times) suffices to optimize time-bounded reachability properties. $\diamond$

For $TAHOPR$ -schedulers, the decision may depend on the current state s and the length of π_1 and π_2 (*hop-counting schedulers*); accordingly, they are isomorphic to mappings $D : \mathcal{S} \times \mathbb{N} \rightarrow Distr(Act)$. Moreover, D is a *time-abstract history-dependent* scheduler ($TAHR$), if $D(\pi_1, \cdot) = D(\pi_2, \cdot)$ for all histories $\pi_1, \pi_2 \in Paths^*$ with $abs(\pi_1) = abs(\pi_2)$: Given history π , $TAHR$ -schedulers may decide based on the sequence of states and actions in $abs(\pi)$. In [BHKH05], the authors show that $TAHOPR$ - and $TAHR$ -schedulers induce the same probability bounds for timed reachability which are tighter than the bounds induced by the class of $TAPR$ -schedulers.

Time-dependent scheduler classes generally induce probability bounds that exceed

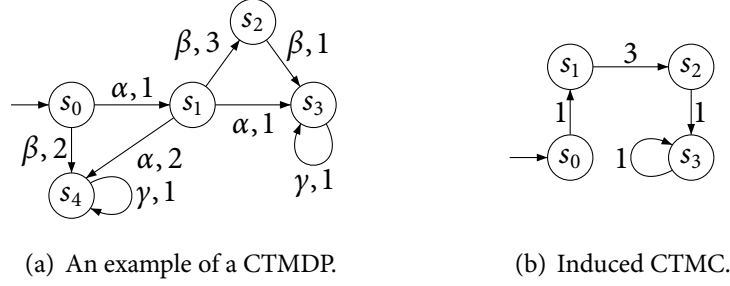


Figure 4.2: An example of a CTMDP and its induced CTMC (under a *TAPD*-scheduler).

those of the corresponding time-abstract classes [BHKH05]. As they are the main focus of this thesis, we discuss them in greater detail here:

If we move from state s_{i-1} to state s_i , a *timed positional* scheduler (*TPR*) yields a distribution over $Act(s_i)$ which depends on the current state s_i and the time it took to go from state s_{i-1} to state s_i ; thus, the class of *TPR*-schedulers extends *TAPR*-schedulers with information on the delay of the last transition.

Similarly, *total time history-dependent* schedulers (*TTHR*) extend *TAHR*-schedulers with information on the time that passed up to the current state: If $D \in TTHR$ and $\pi_1, \pi_2 \in Paths^*$ are histories with $abs(\pi_1) = abs(\pi_2)$ and $\Delta(\pi_1) = \Delta(\pi_2)$, then $D(\pi_1, \cdot) = D(\pi_2, \cdot)$. Here, we use $\Delta(\pi) = \sum_{i=0}^n t_i$ to denote the total time that is spent on a finite path $\pi = s_0 \xrightarrow{\alpha_0, t_0} s_1 \xrightarrow{\alpha_1, t_1} \dots \xrightarrow{\alpha_{n-1}, t_{n-1}} s_n \in Paths^*$. From the definition of *TTHR*, it follows that $TTHR \subseteq GM$. Intuitively, a *TTHR*-schedulers may depend on the accumulated time (that is, on $\Delta(\pi)$), but not on sojourn times in individual states of the history. Hence, for general events, the probability bounds of *TTHR*-schedulers are less strict than those of *GM*-schedulers. However, this does not hold for time-bounded reachability probabilities. To optimize them, an even simpler class of time-dependent schedulers suffices:

For the properties that we investigate in this thesis, the class of *total time positional* schedulers (*TTPR*) is of great importance: A *TTPR*-scheduler is given as a mapping $D : \mathcal{S} \times \mathbb{R}_{\geq 0} \rightarrow Distr(Act)$. Intuitively, it expects the current state in its first argument; the second argument is the total amount of time that has passed before the current state was entered. Hence, *TTPR*-schedulers are similar to *TTHR*-schedulers but abstract from the state-history: For two histories π_1 and π_2 , $D(\pi_1, \cdot) = D(\pi_2, \cdot)$ if π_1 and π_2 end in the same state and if the total amount of time that was spent on π_1 and π_2 is the same, that is, if $\Delta(\pi_1) = \Delta(\pi_2)$.

TTPR-schedulers are of particular interest, as they induce optimal probability bounds with respect to time- and interval bounded reachability objectives: To see this, consider the probability to reach a set of goal states $G \subseteq \mathcal{S}$ within t time units. If state s is reached via $\pi \in Paths^*$ (without visiting G), the maximal probability to enter G is given by a scheduler which maximizes the probability to reach G from state s within the remaining $t - \Delta(\pi)$ time units. Obviously, a *TTPR* scheduler is sufficient in this case. In Chapter 5, we

will come back to this issue (cf. Thm. 5.2 on page 124) and formally prove this claim for a slightly different class of schedulers. However, the proof carries over to *TTPR*-schedulers, trivially.

A further remark is in order here: In [BHKH05] it is proved that *TAHOPD*-schedulers (i.e. *deterministic TAHOPR*-schedulers) suffice for optimizing time-bounded reachability objectives *under all time-abstract schedulers*. This is similar to the continuous-time case, where for time-dependent schedulers it is sufficient to measure the total amount of time that has passed. In particular, information about the state- or action-history (as it is provided by *TAHR*- and *TTHR*-schedulers) is proved to be unnecessary.

Example 4.2. Reconsider the CTMDP depicted in Fig. 4.2(a) and assume that we aim at maximizing the probability to move from state s_0 to state s_3 within a given time bound $z \in \mathbb{R}_{\geq 0}$. Obviously, an optimal *TTPR* scheduler has to choose action α in state s_0 : If it chose β , the CTMDP would move to state s_4 and stay there forever. Thus, we may assume that state s_1 is entered via action α after a sojourn in state s_0 of duration $t_0 \in \mathbb{R}_{\geq 0}$.

Being in state s_1 , a nondeterministic choice between actions α and β occurs: If α is chosen, state s_1 is left with exit rate $E(s_1, \alpha) = \mathbf{R}(s_1, \alpha, s_3) + \mathbf{R}(s_1, \alpha, s_4) = 3$. However, the probability $\mathbf{P}(s_1, \alpha, s_3) = \frac{\mathbf{R}(s_1, \alpha, s_3)}{E(s_1, \alpha)}$ to enter state s_3 (instead of state s_4) is only $\frac{1}{3}$. If action β is chosen, the situation is different: Although the rate for leaving state s_1 under action β is the same (i.e. $E(s_1, \beta) = \mathbf{R}(s_1, \beta, s_2) = 3$), we do not enter the goal state s_3 directly. Instead, the transition from state s_2 to state s_3 with rate $\mathbf{R}(s_2, \beta, s_3) = 1$ induces an additional delay. However, note that if action β is chosen in state s_1 , we reach state s_3 with probability 1.

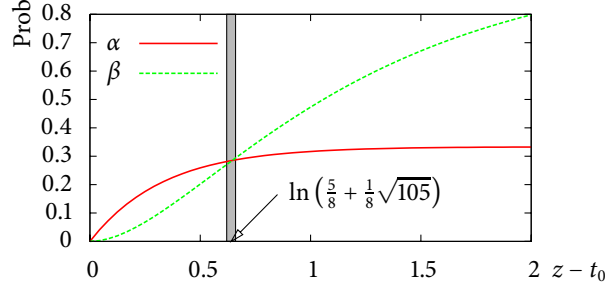
Obviously, the optimal decision in state s_1 depends on the time $z - t_0$ that remains to reach s_3 when t_0 time units have been spent in state s_0 , already. With this reasoning, we obtain an optimal *TTPR*-scheduler D as follows: Define $D(s_0, 0) = \{\alpha \mapsto 1\}$ and $D(s_1, t_0) = \{\alpha \mapsto 1\}$ if $t_0 \geq z - \ln\left(\frac{5}{8} + \frac{1}{8}\sqrt{105}\right)$ and $D(s_1, t_0) = \{\beta \mapsto 1\}$, otherwise.

The derivation for D is as follows: The probability to move within the remaining $x = z - t_0$ time units from state s_1 to state s_3 with action α is given by the function $a(x) = \frac{1}{3}(1 - e^{-3x})$. For action β , the corresponding function $b(x)$ is given by the convolution to go to state s_3 via state s_2 . Hence $b(x) = \int_0^x (3e^{-3t_1} \int_0^{x-t_1} e^{-t_2} dt_2) dt_1$. Fig. 4.3 depicts the two cumulative distribution functions. Now, let $d \in \mathbb{R}_{\geq 0}$ be the unique solution of the equation $a(x) = b(x)$; then $d = \ln\left(\frac{5}{8} + \frac{1}{8}\sqrt{105}\right)$. Obviously, if more than d time units remain, i.e. if $z - t_0 > d$, the optimal decision in state s_1 is action β . On the other hand, if $z - t_0 \leq d$, it is more profitable to choose action α .

For now, we note that (a) time-abstract schedulers obviously do not suffice to obtain the maximum probability and (b) that the scheduler D is a deterministic *TTPD*-scheduler. $\diamond$

With the preceding informal description of the scheduler classes that are mentioned in Fig. 4.1, we define them formally as follows:

Definition 4.1 (Scheduler classes). Let $\mathcal{C}$ be a CTMDP and D a GM-scheduler on $\mathcal{C}$. If

Figure 4.3: Reachability in $z - t$ time units.

π and π' range over $\text{Paths}^*(\mathcal{C})$, the scheduler classes are defined as follows:

$D \in \text{TAPR}$	$\iff \pi \downarrow = \pi' \downarrow \Rightarrow D(\pi) = D(\pi')$
$D \in \text{TAHOPR}$	$\iff (\pi \downarrow = \pi' \downarrow \wedge \pi = \pi') \Rightarrow D(\pi) = D(\pi')$
$D \in \text{TAHR}$	$\iff \text{abs}(\pi) = \text{abs}(\pi') \Rightarrow D(\pi) = D(\pi')$
$D \in \text{TTHR}$	$\iff (\text{abs}(\pi) = \text{abs}(\pi') \wedge \Delta(\pi) = \Delta(\pi')) \Rightarrow D(\pi) = D(\pi')$
$D \in \text{TTPR}$	$\iff (\pi \downarrow = \pi' \downarrow \wedge \Delta(\pi) = \Delta(\pi')) \Rightarrow D(\pi) = D(\pi')$
$D \in \text{TPR}$	$\iff (\pi \downarrow = \pi' \downarrow \wedge \delta(\pi, \pi - 1) = \delta(\pi', \pi' - 1)) \Rightarrow D(\pi) = D(\pi').$

Def. 4.1 justifies to restrict the domain of the schedulers to the information the respective class exploits. In this way, we obtain the characterization in Table 4.1.

In the next section, we come to a transformation on CTMDPs that unifies the speed of outgoing transitions and thereby allows us to defer the resolution of nondeterministic choices: Intuitively, if the sojourn time in a state does not depend on the scheduler, the decision needs not be taken when entering that state, but may be delayed up to the point when the state is left.

4.2 Local uniformization

Generally, the exit rate of a state depends on the action that is chosen by the scheduler in that state. Intuitively, this dependency requires that the scheduler selects the action to continue with directly upon entering a state: Imagine a state s with $\text{Act}(s) = \{\alpha, \beta\}$ such that $E(s, \alpha) \neq E(s, \beta)$: If the nondeterministic choice between α and β was not resolved immediately when entering state s , it is unclear whether the delay in state s is distributed according to $E(s, \alpha)$ or according to $E(s, \beta)$.

For general CTMDPs, we assume that schedulers decide directly each time the CTMDP enters a new state. In particular, if state s is entered at time t and action $\alpha \in \text{Act}(s)$ is chosen by the associated scheduler D , we do not consider the case that D decides for a different action at some later time $t + \varepsilon$ during the sojourn period in state s .

	scheduler class	scheduler signature
time abstract	positional (<i>TAPR</i>)	$D : \mathcal{S} \rightarrow \text{Distr}(\text{Act})$
	hop-counting (<i>TAHOPR</i>)	$D : \mathcal{S} \times \mathbb{N} \rightarrow \text{Distr}(\text{Act})$
	time abstract history dependent (<i>TAHR</i>)	$D : \text{Paths}_{abs}^* \rightarrow \text{Distr}(\text{Act})$
time dependent	timed history dependent (<i>GM</i>)	full timed history $D : \text{Paths}^* \rightarrow \text{Distr}(\text{Act})$
	total time history dependent (<i>TTHR</i>)	sequence of states & total time $D : \text{Paths}_{abs}^* \times \mathbb{R}_{\geq 0} \rightarrow \text{Distr}(\text{Act})$
	total time positional (<i>TTPR</i>)	last state & total time $D : \mathcal{S} \times \mathbb{R}_{\geq 0} \rightarrow \text{Distr}(\text{Act})$
	timed positional (<i>TPR</i>)	last state & delay of last transition $D : \mathcal{S} \times \mathbb{R}_{\geq 0} \rightarrow \text{Distr}(\text{Act})$

Table 4.1: Proposed scheduler classes for CTMDPs.

However, such schedulers are interesting as they may correct decisions that have been made earlier during the sojourn in the current state: For example, such a scheduler could switch to another action if the sojourn takes longer than a given threshold.

In this chapter, we make a first step towards such scheduler classes. Therefore, we identify a strict subclass of CTMDPs where the states' sojourn time distributions are independent of the action that is chosen in the current state. For this subclass, we are able to disentangle the sojourn time distribution and the scheduling decision. More precisely, we define *locally uniform* CTMDPs which require that all exit-rates are state-wise constant for the available actions:

Definition 4.2 (Local uniformity). A CTMDP $(\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ is locally uniform iff there exists $u : \mathcal{S} \rightarrow \mathbb{R}_{>0}$ such that $E(s, \alpha) = u(s)$ for all $s \in \mathcal{S}$ and $\alpha \in \text{Act}(s)$.

In locally uniform CTMDPs, each state s has a unique exit rate $u(s)$; hence, its sojourn time distribution does not depend on the action that is chosen by the scheduler. In this way, locally uniform CTMDPs allow to delay the scheduler's decision until the current state is left. As an implication, we can define a new class of schedulers, which decides only upon *leaving* the current state. Such schedulers allow to resolve the nondeterministic choice when the sojourn in the current state is over. Hence, they are referred to as *late schedulers* to distinguish them from the *early schedulers* that are defined for general CTMDPs.

As we will see in Sec. 4.4, late schedulers profit from the fact that they can defer their decision to the end of a state's sojourn time: In particular, they can incorporate the time that was spent in the current state into their decision, which is why they strictly outper-

form early schedulers (cf. Sec. 4.4).

Moreover, note that late schedulers on locally uniform CTMDPs are equivalent to schedulers that can take back their decisions during the sojourn in a given state: To see this, note that in a locally uniform CTMDP, the decision that determines the CTMDP's stochastic behavior is the one that is taken precisely when leaving the current state. All previous decisions do not influence the associated stochastic process.

Due to their interesting properties, this section investigates locally uniform CTMDPs more closely. Therefore, we postpone the discussion about late schedulers to Sec. 4.4 and Chapter 5, where we consider them in more detail. As the prerequisite for late schedulers are locally uniform CTMDPs, let us first define a transformation on general CTMDPs — called *local uniformization* — which achieves local uniformity and investigate its properties with respect to early schedulers:

Definition 4.3 (Local uniformization). Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, v)$ be a CTMDP and define $u(s) = \max \{E(s, \alpha) \mid \alpha \in \text{Act}(s)\}$ for all $s \in \mathcal{S}$. Then $\bar{\mathcal{C}} = (\bar{\mathcal{S}}, \text{Act}, \bar{\mathbf{R}}, \bar{v})$ is the locally uniform CTMDP induced by $\mathcal{C}$, where $\bar{\mathcal{S}} = \mathcal{S} \cup \mathcal{S}_{cp}$, $\mathcal{S}_{cp} = \{s^\alpha \mid E(s, \alpha) < u(s)\}$ and

$$\bar{\mathbf{R}}(s, \alpha, s') = \begin{cases} \mathbf{R}(s, \alpha, s') & \text{if } s, s' \in \mathcal{S} \\ \mathbf{R}(t, \alpha, s') & \text{if } s = t^\alpha \wedge s' \in \mathcal{S} \\ u(s) - E(s, \alpha) & \text{if } s \in \mathcal{S} \wedge s' = s^\alpha \\ 0 & \text{otherwise.} \end{cases}$$

Further, $\bar{v}(s) = v(s)$ if $s \in \mathcal{S}$ and 0, otherwise.

Local uniformization is done for each state s separately with uniformization rate $u(s)$. If the exit rate of s under action α is less than $u(s)$, we introduce a copy-state s^α and an α -transition which carries the missing rate $\mathbf{R}(s, \alpha, s^\alpha) = u(s) - E(s, \alpha)$. Regarding s^α , only the outgoing α -transitions of s carry over to s^α . Hence s^α is deterministic in the sense that $\text{Act}(s^\alpha) = \{\alpha\}$.

Example 4.3. Consider the fragment CTMDP in Fig. 4.4(a), where $\lambda = \sum \lambda_i$ and $\lambda_i, \mu > 0$ for $i = 0, 1, 2$. It is not locally uniform as $E(s_0, \alpha) = \lambda$ and $E(s_0, \beta) = \lambda + \mu$. By applying our transformation we obtain the locally uniform CTMDP in Fig. 4.4(b) as follows: We set $u(s_0) = \lambda + \mu$ and introduce the copy-state s_0^α . As $E(s_0, \alpha) < u(s_0)$, we add a new α -transition from state s_0 to its copy-state s_0^α with rate μ . Further, all α -transitions of state s_0 (and only those) carry over to state s_0^α ; hence, α -transitions lead from state s_0^α to states s_1 and s_2 with rates λ_1 and λ_2 , respectively. Accordingly, the α -self-loop in state s_0 in Fig. 4.4(a) induces a new α -transition in Fig. 4.4(b) which leads from state s_0^α back to state s_0 . $\diamond$

Local uniformization of $\mathcal{C}$ introduces new states and transitions in $\bar{\mathcal{C}}$. The paths in $\bar{\mathcal{C}}$ reflect this and differ from those of $\mathcal{C}$; more precisely, they may contain sequences of

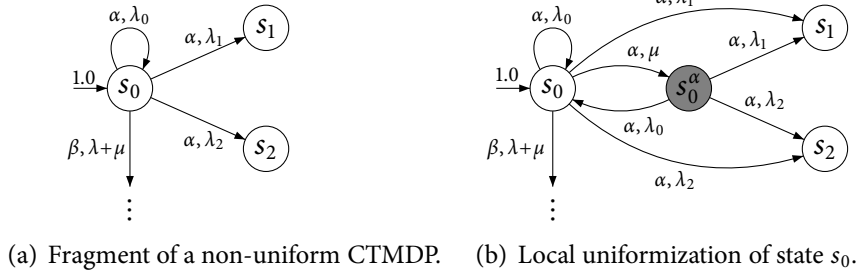


Figure 4.4: How to obtain locally uniform CTMDPs by introducing copy states.

transitions $s \xrightarrow{\alpha, t} s^\alpha \xrightarrow{\alpha, t'} s'$ where s^α is a copy-state. Intuitively, if we identify s and s^α , this corresponds to a single transition $s \xrightarrow{\alpha, t+t'} s'$ in $\mathcal{C}$. To formalize this correspondence, we derive a mapping *merge* on all *valid* paths $\bar{\pi} \in \text{Paths}^*(\bar{\mathcal{C}})$ with $\bar{\pi}[0], \bar{\pi}\downarrow \in \mathcal{S}$: If $|\bar{\pi}| = 0$, $\text{merge}(\bar{\pi}) = \bar{\pi}[0]$. Otherwise, let

$$\text{merge}(s \xrightarrow{\alpha, t} \bar{\pi}) = \begin{cases} s \xrightarrow{\alpha, t} \text{merge}(\bar{\pi}) & \text{if } \bar{\pi}[0] \in \mathcal{S} \\ s \xrightarrow{\alpha, t+t'} \text{merge}(\bar{\pi}') & \text{if } \bar{\pi} = s^\alpha \xrightarrow{\alpha, t'} \bar{\pi}'. \end{cases}$$

Note that the function *merge* is defined only for valid paths, that is, for paths $\bar{\pi}$ whose transitions correspond to existing transitions in the underlying CTMDP $\bar{\mathcal{C}}$. Ignoring invalid paths is justified by the fact, that the set of invalid paths always has probability measure 0 (cf. Def. 3.14), independent of the scheduler.

Naturally, *merge* extends to infinite paths if we do not require $\bar{\pi}\downarrow \in \mathcal{S}$; further, merging a set of paths $\bar{\Pi}$ is defined element-wise and denoted $\text{merge}(\bar{\Pi})$.

Example 4.4. Let $\bar{\pi} = s_0 \xrightarrow{\alpha_0, t_0} s_0^{\alpha_0} \xrightarrow{\alpha_0, t'_0} s_1 \xrightarrow{\alpha_1, t_1} s_2 \xrightarrow{\alpha_2, t_2} s_2^{\alpha_2} \xrightarrow{\alpha_2, t'_2} s_3$ be a path in $\bar{\mathcal{C}}$. Then $\text{merge}(\bar{\pi}) = s_0 \xrightarrow{\alpha_0, t_0+t'_0} s_1 \xrightarrow{\alpha_1, t_1} s_2 \xrightarrow{\alpha_2, t_2+t'_2} s_3$. $\diamond$

Intuitively, the function *merge* collapses the copy states that are introduced in the locally uniform CTMDP $\bar{\mathcal{C}}$ and maps to valid paths in the underlying (not locally uniform) CTMDP $\mathcal{C}$. For the reverse direction, we map sets of paths in $\mathcal{C}$ to sets of paths in $\bar{\mathcal{C}}$. To do so, note that any single path in $\mathcal{C}$ corresponds to a countably infinite set of paths in $\bar{\mathcal{C}}$: Let $s_0 \xrightarrow{\alpha_0, t_0} s_1$ be a path in $\mathcal{C}$; it corresponds to the set $\{\bar{\pi} = s_0 \xrightarrow{\alpha_0, t} s_0^{\alpha_0} \xrightarrow{\alpha_0, t'} s_1 \mid t + t' = t_0\}$ of paths in $\bar{\mathcal{C}}$. We formalize this extension to paths in $\bar{\mathcal{C}}$ as follows:

If $\Pi \subseteq \text{Paths}(\mathcal{C})$, we define

$$\text{extend}(\Pi) = \{\bar{\pi} \in \text{Paths}(\bar{\mathcal{C}}) \mid \text{merge}(\bar{\pi}) \in \Pi\}.$$

To conclude this section, let us state some natural properties of the functions *merge* and *extend* which prove useful to establish the formal results in the remainder of this chapter:

Lemma 4.1. *Let $\mathcal{C}$ be a CTMDP and $\Pi_1, \Pi_2, \dots \subseteq \text{Paths}(\mathcal{C})$. Then the following propositions hold:*

1. $\Pi_1 \subseteq \Pi_2 \Rightarrow \text{extend}(\Pi_1) \subseteq \text{extend}(\Pi_2)$,
2. $\Pi_1 \cap \Pi_2 = \emptyset \Rightarrow \text{extend}(\Pi_1) \cap \text{extend}(\Pi_2) = \emptyset$ and
3. $\bigcup \text{extend}(\Pi_k) = \text{extend}\left(\bigcup \Pi_k\right)$.

Proof. We prove each claim separately:

1. $\Pi_1 \subseteq \Pi_2 \Rightarrow \text{extend}(\Pi_1) \subseteq \text{extend}(\Pi_2)$ follows directly from the definition of $\text{extend}(\Pi)$: To see this, note that if $\Pi_1 \subseteq \Pi_2$, then it holds

$$\{\bar{\pi} \in \text{Paths}(\bar{\mathcal{C}}) \mid \text{merge}(\bar{\pi}) \in \Pi_1\} \subseteq \{\bar{\pi} \in \text{Paths}(\bar{\mathcal{C}}) \mid \text{merge}(\bar{\pi}) \in \Pi_2\}.$$

2. We prove the claim by contraposition: Therefore, assume that $\Pi_1 \cap \Pi_2 = \emptyset$ but $\bar{\pi} \in \text{extend}(\Pi_1) \cap \text{extend}(\Pi_2)$. Then

$$\bar{\pi} \in \{\bar{\pi}' \in \text{Paths}(\bar{\mathcal{C}}) \mid \text{merge}(\bar{\pi}') \in \Pi_1 \wedge \text{merge}(\bar{\pi}') \in \Pi_2\}.$$

But $\Pi_1 \cap \Pi_2 = \emptyset$. Hence we obtain the desired contradiction.

3. For any set $I \subseteq \mathbb{N}$ we have that

$$\begin{aligned} \bigcup_{k \in I} \text{extend}(\Pi_k) &= \bigcup_{k \in I} \{\bar{\pi} \in \text{Paths}(\bar{\mathcal{C}}) \mid \text{merge}(\bar{\pi}) \in \Pi_k\} \\ &= \{\bar{\pi} \in \text{Paths}(\bar{\mathcal{C}}) \mid \text{merge}(\bar{\pi}) \in \bigcup_{k \in I} \Pi_k\} = \text{extend}\left(\bigcup_{k \in I} \Pi_k\right). \quad \square \end{aligned}$$

In the following, we investigate which classes of early schedulers induce the same probability measures for paths in a CTMDP $\mathcal{C}$ and the corresponding set of paths in $\bar{\mathcal{C}}$. Thus, we identify the scheduler classes for which local uniformization is a measure preserving transformation.

For the proof, we proceed stepwise and first adopt a local view: In Sec. 4.2.1, we show that the probability of a single step in $\mathcal{C}$ in which the nondeterministic choice has already been resolved equals the probability of the corresponding steps in $\bar{\mathcal{C}}$. The results are used in Sec. 4.2.2 to define a scheduler $\bar{D}$ on $\bar{\mathcal{C}}$ that corresponds to a given scheduler D on $\mathcal{C}$ and induces the same probabilities.

4.2.1 One-step correctness of local uniformization

Consider the CTMDP in Fig. 4.4(a), where $\lambda = \sum \lambda_i$ and $\lambda_i > 0$ for $i = 0, 1, 2$. Assume that action α is chosen in state s_0 ; then $\frac{\mathbf{R}(s_0, \alpha, s_i)}{\mathbf{E}(s_0, \alpha)} = \frac{\lambda_i}{\lambda}$ is the probability to move to state s_i . Hence the probability to reach state s_i in time interval $[0, t]$ is

$$\frac{\lambda_i}{\lambda} \int_0^t \eta_\lambda(dt_1). \quad (4.1)$$

Let us compute the same probability for $\bar{\mathcal{C}}$ depicted in Fig. 4.4(b): The probability to go from s_0 to s_i directly (with action α) is $\frac{\bar{\mathbf{R}}(s_0, \alpha, s_i)}{\bar{\mathbf{E}}(s_0, \alpha)} = \frac{\lambda_i}{\lambda + \mu}$; however, with probability $\frac{\bar{\mathbf{R}}(s_0, \alpha, s_0^\alpha)}{\bar{\mathbf{E}}(s_0, \alpha)} = \frac{\mu}{\lambda + \mu} \cdot \frac{\lambda_i}{\lambda}$ we instead move to state s_0^α and only then to s_i . In this case, the probability that in the time interval $[0, t]$, an α -transition executes in state s_0 , followed by one of s_0^α is $\int_0^t (\lambda + \mu) e^{-(\lambda + \mu)t_1} \int_0^{t-t_1} \lambda e^{-\lambda t_2} dt_2 dt_1$. Hence, we reach state s_i with action α in at most t time units with probability

$$\frac{\lambda_i}{\lambda + \mu} \int_0^t \eta_{\lambda + \mu}(dt_1) + \frac{\mu}{\lambda + \mu} \cdot \frac{\lambda_i}{\lambda} \int_0^t \eta_{\lambda + \mu}(dt_1) \int_0^{t-t_1} \eta_\lambda(dt_2). \quad (4.2)$$

It is easy to verify that (4.1) and (4.2) are equal:

Lemma 4.2 (Local correctness). *Let $\mathcal{C}$ and $\bar{\mathcal{C}}$ be the CTMDPs depicted in Fig. 4.4. For $i \in \{0, \dots, 2\}$, $\lambda_i, \mu > 0$ and $t \in \mathbb{R}_{\geq 0}$ it holds*

$$\frac{\lambda_i}{\lambda} \int_0^t \eta_\lambda(dt) = \frac{\lambda_i}{\lambda + \mu} \int_0^t \eta_{\lambda + \mu}(dt) + \frac{\mu}{\lambda + \mu} \cdot \frac{\lambda_i}{\lambda} \int_0^t \eta_{\lambda + \mu}(dt_1) \int_0^{t-t_1} \eta_\lambda(dt_2). \quad (4.3)$$

Proof. We can rewrite the right-hand side in Eq. (4.3) as follows:

$$\begin{aligned} & \frac{\lambda_i}{\lambda + \mu} \int_0^t (\lambda + \mu) e^{-(\lambda + \mu)t_1} dt_1 + \frac{\mu}{\lambda + \mu} \cdot \frac{\lambda_i}{\lambda} \int_0^t (\lambda + \mu) e^{-(\lambda + \mu)t_1} \int_0^{t-t_1} \lambda e^{-\lambda t_2} dt_2 dt_1 \\ &= \lambda_i \int_0^t e^{-(\lambda + \mu)t_1} dt_1 + \frac{\mu \cdot \lambda_i}{\lambda} \int_0^t e^{-(\lambda + \mu)t_1} (1 - e^{-\lambda(t-t_1)}) dt_1 \\ &= \lambda_i \int_0^t e^{-(\lambda + \mu)t_1} dt_1 + \frac{\mu \cdot \lambda_i}{\lambda} \int_0^t e^{-(\lambda + \mu)t_1} dt_1 - \frac{\mu \cdot \lambda_i}{\lambda} \int_0^t e^{-(\lambda + \mu)t_1 - \lambda(t-t_1)} dt_1. \end{aligned}$$

Note that the first two integrals are equal. This yields

$$= \lambda_i \left(1 + \frac{\mu}{\lambda}\right) \int_0^t e^{-(\lambda + \mu)t_1} dt_1 - \frac{\mu \cdot \lambda_i}{\lambda} \int_0^t e^{-\mu t_1 - \lambda t} dt_1$$

By rewriting the term $\lambda_i \left(1 + \frac{\mu}{\lambda}\right)$, we obtain the factor $(\lambda + \mu)$ and the exponential density for rate $(\lambda + \mu)$:

$$\begin{aligned}
&= \frac{\lambda_i}{\lambda} \int_0^t (\lambda + \mu) e^{-(\lambda+\mu)t_1} dt_1 - \frac{\mu \cdot \lambda_i}{\lambda} \int_0^t e^{-\mu t_1 - \lambda t} dt_1 \\
&= \frac{\lambda_i}{\lambda} (1 - e^{-(\lambda+\mu)t}) - \frac{\mu \cdot \lambda_i}{\lambda} e^{-\lambda t} \int_0^t e^{-\mu t_1} dt_1 \\
&= \frac{\lambda_i}{\lambda} \left(1 - e^{-(\lambda+\mu)t} - \mu e^{-\lambda t} \int_0^t e^{-\mu t_1} dt_1\right) \\
&= \frac{\lambda_i}{\lambda} (1 - e^{-(\lambda+\mu)t} - e^{-\lambda t} (1 - e^{-\mu t})) \\
&= \frac{\lambda_i}{\lambda} (1 - e^{-(\lambda+\mu)t} - e^{-\lambda t} + e^{-(\lambda+\mu)t}) \\
&= \frac{\lambda_i}{\lambda} (1 - e^{-\lambda t}).
\end{aligned}$$

□

Thus the probability to reach a (non-copy) successor state in $\{s_0, s_1, s_2\}$ is the same for $\mathcal{C}$ and $\overline{\mathcal{C}}$. It can be computed by replacing λ_i with $\sum \lambda_i$ in Eq. (4.1) and Eq. (4.2). Further, note that the result of Lemma 4.2 extends naturally to finitely many successor states $\{s_0, s_1, \dots, s_n\}$. Moreover, if the interval $[0, t]$ is replaced by an element from the Borel σ -field $\mathfrak{B}(\mathbb{R}_{\geq 0})$ and all integrals are interpreted as Lebesgue-integrals, we obtain a straightforward extension of Lemma 4.2 to the class of Borel measurable sets of time points.

Next, we prove that Equalities (4.1) and (4.2) are preserved even if we integrate over a Borel-measurable function $f : \mathbb{R}_{\geq 0} \rightarrow [0, 1]$. To keep our notation as simple as possible, we only consider the probability to reach an arbitrary non-copy state within a Borel measurable set of time points $T \in \mathfrak{B}(\mathbb{R}_{\geq 0})$. Compared to Lemma 4.2, we therefore replace the rate λ_i to move to the i -th non-copy successor state by the cumulated rate $\lambda = \sum \lambda_i$ to go to any non-copy state:

Lemma 4.3 (One-step timing). *Let $f : \mathbb{R}_{\geq 0} \rightarrow [0, 1]$ be a Borel measurable function and $T \in \mathfrak{B}(\mathbb{R}_{\geq 0})$. Then*

$$\begin{aligned}
\int_T f(t) \eta_\lambda(dt) &= \frac{\lambda}{\lambda + \mu} \int_T f(t) \eta_{\lambda+\mu}(dt) \\
&\quad + \frac{\mu}{\lambda + \mu} \int_{\mathbb{R}_{\geq 0}} \eta_{\lambda+\mu}(dt_1) \int_{T \ominus t_1} f(t_1 + t_2) \eta_\lambda(dt_2).
\end{aligned} \tag{4.4}$$

Proof. As usual when proving properties about Lebesgue integrals of Borel measurable functions, we prove the claim stepwise and work our way up from nonnegative simple functions (cf. Def. 2.15 on page 35) to arbitrary nonnegative Borel measurable functions

(cf. Def. 2.14). First, assume that $f : \mathbb{R}_{\geq 0} \rightarrow [0, 1]$ is a simple function. If $T \subseteq \mathbb{R}_{\geq 0}$, then it is easy to see that $f \circ \mathbf{I}_T$ is again a simple function. With this remark, we can rewrite the Lebesgue integral on the right hand side of Eq. (4.4) and obtain

$$\begin{aligned} & \frac{\lambda}{\lambda + \mu} \int_T f(t) \eta_{\lambda+\mu}(dt) + \frac{\mu}{\lambda + \mu} \int_{\mathbb{R}_{\geq 0}} \eta_{\lambda+\mu}(dt_1) \int_{T \ominus t_1} f(t_1 + t_2) \eta_{\lambda}(dt_2) \\ &= \frac{\lambda}{\lambda + \mu} \int_{\mathbb{R}_{\geq 0}} f(t) \cdot \mathbf{I}_T(t) \eta_{\lambda+\mu}(dt) \\ & \quad + \frac{\mu}{\lambda + \mu} \int_{\mathbb{R}_{\geq 0}} \eta_{\lambda+\mu}(dt_1) \int_{\mathbb{R}_{\geq 0}} f(t_1 + t_2) \cdot \mathbf{I}_T(t_1 + t_2) \eta_{\lambda}(dt_2). \end{aligned}$$

Note that in order to rewrite the innermost Lebesgue integral, we further make use of the fact that $t_2 \in T \ominus t_1 \iff t_1 + t_2 \in T$. Applying Fubini's theorem (Thm. 2.17 on page 47), we can switch to a two-dimensional product. In this way, we continue:

$$\begin{aligned} &= \frac{\lambda}{\lambda + \mu} \int_{\mathbb{R}_{\geq 0}} f(t) \cdot \mathbf{I}_T(t) \eta_{\lambda+\mu}(dt) \\ & \quad + \frac{\mu}{\lambda + \mu} \int_{\mathbb{R}_{\geq 0} \times \mathbb{R}_{\geq 0}} f(t_1 + t_2) \cdot \mathbf{I}_T(t_1 + t_2) (\eta_{\lambda+\mu} \times \eta_{\lambda})(d(t_1, t_2)) \end{aligned}$$

The assumption that $f(t)$ is a simple function implies that also $f \circ \mathbf{I}_T : \mathbb{R}_{\geq 0} \rightarrow [0, 1]$ and $f' : \mathbb{R}_{\geq 0}^2 \rightarrow [0, 1] : (t_1, t_2) \mapsto f(t_1 + t_2)$ are simple functions. Now let $\{x_1, x_2, \dots, x_r\}$ be the (finitely many) values that $f \circ \mathbf{I}_T$ takes in $\mathbb{R}_{\geq 0}$ and define $A_j = (f \circ \mathbf{I}_T)^{-1}(x_j)$ and $A'_j = f'^{-1}(x_j)$ for all $j = 1, 2, \dots, r$. With these choices, we can continue to rewrite our integral as follows:

$$\begin{aligned} &= \frac{\lambda}{\lambda + \mu} \sum_{j=1}^r x_j \cdot \eta_{\lambda+\mu}(A_j) + \frac{\mu}{\lambda + \mu} \sum_{j=1}^r x_j \cdot (\eta_{\lambda+\mu} \times \eta_{\lambda})(A'_j) \\ &= \sum_{j=1}^r \left(\frac{\lambda}{\lambda + \mu} \cdot x_j \cdot \eta_{\lambda+\mu}(A_j) + \frac{\mu}{\lambda + \mu} \cdot x_j \cdot (\eta_{\lambda+\mu} \times \eta_{\lambda})(A'_j) \right) \\ &= \sum_{j=1}^r x_j \left(\frac{\lambda}{\lambda + \mu} \cdot \eta_{\lambda+\mu}(A_j) + \frac{\mu}{\lambda + \mu} \cdot (\eta_{\lambda+\mu} \times \eta_{\lambda})(A'_j) \right) \\ &= \sum_{j=1}^r x_j \left(\frac{\lambda}{\lambda + \mu} \cdot \eta_{\lambda+\mu}(A_j) + \frac{\mu}{\lambda + \mu} \cdot \int_{\mathbb{R}_{\geq 0} \times \mathbb{R}_{\geq 0}} \mathbf{I}_{A'_j}(t_1 + t_2) (\eta_{\lambda+\mu} \times \eta_{\lambda})(d(t_1, t_2)) \right). \end{aligned}$$

Now we can apply Thm. 2.17 reversely and come back to an iterated integration:

$$\begin{aligned} &= \sum_{j=1}^r x_j \left(\frac{\lambda}{\lambda + \mu} \cdot \eta_{\lambda+\mu}(A_j) + \frac{\mu}{\lambda + \mu} \cdot \int_{\mathbb{R}_{\geq 0}} \eta_{\lambda+\mu}(dt_1) \int_{\mathbb{R}_{\geq 0}} \mathbf{I}_{A'_j}(t_1 + t_2) \eta_{\lambda}(dt_2) \right) \\ &= \sum_{j=1}^r x_j \left(\frac{\lambda}{\lambda + \mu} \cdot \eta_{\lambda+\mu}(A_j) + \frac{\mu}{\lambda + \mu} \cdot \int_{\mathbb{R}_{\geq 0}} \eta_{\lambda+\mu}(dt_1) \int_{A_j \ominus t_1} \eta_{\lambda}(dt_2) \right) \end{aligned}$$

$$\stackrel{(*)}{=} \sum_{j=1}^r x_j \cdot \int_{A_j} \eta_\lambda(dt) = \sum_{j=1}^r x_j \cdot \eta_\lambda(A_j) = \int_T f(t) \eta_\lambda(dt).$$

Here, the equality (*) follows by extending Lemma 4.2 from intervals to Borel-measurable subsets of $\mathbb{R}_{\geq 0}$ which can be done easily by replacing the Riemann integrals over intervals in the proof of Lemma 4.2 by the corresponding Lebesgue integral over measurable subsets of $\mathbb{R}_{\geq 0}$.

Further, if $f : \mathbb{R}_{\geq 0} \rightarrow [0, 1]$ is Borel measurable, then Thm. 2.11 (on page 36) implies that there exists a sequence of nonnegative simple functions f_n such that $f_n(t) \rightarrow f(t)$ for all $t \in \mathbb{R}_{\geq 0}$. Further, Eq. (4.4) holds for all f_n . With the monotone convergence theorem (Thm. 2.13 on page 38), we obtain

$$\begin{aligned} \int_T f(t) \eta_\lambda(t) dt &= \lim_{n \rightarrow \infty} \int_T f_n(t) \eta_\lambda(dt) \\ &= \lim_{n \rightarrow \infty} \left(\frac{\lambda}{\lambda + \mu} \int_T f_n(t) \eta_{\lambda+\mu}(dt) + \frac{\mu}{\lambda + \mu} \int_{\mathbb{R}_{\geq 0}} \eta_{\lambda+\mu}(dt_1) \int_{T \ominus t_1} f_n(t_1 + t_2) \eta_\lambda(dt_2) \right) \\ &= \frac{\lambda}{\lambda + \mu} \lim_{n \rightarrow \infty} \int_T f_n(t) \eta_{\lambda+\mu}(dt) + \frac{\mu}{\lambda + \mu} \lim_{n \rightarrow \infty} \int_{\mathbb{R}_{\geq 0}} \eta_{\lambda+\mu}(dt_1) \int_{T \ominus t_1} f_n(t_1 + t_2) \eta_\lambda(dt_2) \\ &= \frac{\lambda}{\lambda + \mu} \lim_{n \rightarrow \infty} \int_T f_n(t) \eta_{\lambda+\mu}(t) dt \\ &\quad + \frac{\mu}{\lambda + \mu} \lim_{n \rightarrow \infty} \int_{\mathbb{R}_{\geq 0} \times \mathbb{R}_{\geq 0}} f_n(t_1 + t_2) \cdot \mathbf{I}_T(t_1 + t_2) (\eta_{\lambda+\mu} \times \eta_\lambda) (d(t_1, t_2)) \\ &= \frac{\lambda}{\lambda + \mu} \int_T f(t) \eta_{\lambda+\mu}(dt) \\ &\quad + \frac{\mu}{\lambda + \mu} \int_{\mathbb{R}_{\geq 0} \times \mathbb{R}_{\geq 0}} f(t_1 + t_2) \cdot \mathbf{I}_T(t_1 + t_2) (\eta_{\lambda+\mu} \times \eta_\lambda) (d(t_1, t_2)) \\ &= \frac{\lambda}{\lambda + \mu} \int_T f(t) \eta_{\lambda+\mu}(dt) + \frac{\mu}{\lambda + \mu} \int_{\mathbb{R}_{\geq 0}} \eta_{\lambda+\mu}(dt_1) \int_{T \ominus t_1} f_n(t_1 + t_2) \eta_\lambda(dt_2). \quad \square \end{aligned}$$

The equality of the terms (4.1) and (4.2) proves that the probability of a single step in $\mathcal{C}$ equals the probability of one or two transitions (depending on the copy-state) in $\overline{\mathcal{C}}$.

In the next section, we lift this argument to sets of paths in $\mathcal{C}$ and $\overline{\mathcal{C}}$. Further, note that we did not consider nondeterministic choices yet. This gap will also be bridged in the next section, where we infer a scheduler $\overline{D}$ from a given scheduler D that makes use of the strong relation between the CTMDP $\mathcal{C}$ and its locally uniform counterpart $\overline{\mathcal{C}}$.

4.2.2 Local uniformization is measure preserving

In this section, we prove that for any GM-scheduler $D \in GM(\mathcal{C})$ and for each CTMDP $\mathcal{C}$ there exists a GM-scheduler $\overline{D} \in GM(\overline{\mathcal{C}})$ in $\overline{\mathcal{C}}$ such that the induced probabilities for the sets of paths Π and $extend(\Pi)$ are equal.

However, as $\bar{\mathcal{C}}$ differs from $\mathcal{C}$, we cannot use D to infer probabilities on $\bar{\mathcal{C}}$ directly. Instead, given a history $\bar{\pi}$ in $\bar{\mathcal{C}}$, we define $\bar{D}(\bar{\pi}, \cdot)$ such that it mimics the decision that D takes in $\mathcal{C}$ for history $\text{merge}(\bar{\pi})$. This is formalized as follows: For all $\bar{\pi} \in \text{Paths}^*(\bar{\mathcal{C}})$, define the GM-scheduler $\bar{D}$ such that

$$\bar{D}(\bar{\pi}, \cdot) = \begin{cases} D(\pi, \cdot) & \text{if } \bar{\pi}[0], \bar{\pi} \downarrow \in \mathcal{S} \wedge \text{merge}(\bar{\pi}) = \pi \\ \{\alpha \mapsto 1\} & \text{if } \bar{\pi} \downarrow = s^\alpha \in \mathcal{S}_{\text{cp}} \\ \gamma_{\bar{\pi}} & \text{otherwise,} \end{cases}$$

where $\gamma_{\bar{\pi}}$ is an arbitrary distribution over $\text{Act}(\bar{\pi} \downarrow)$: If merge is applicable to $\bar{\pi}$ (i.e. if $\bar{\pi}$ is a valid path in $\bar{\mathcal{C}}$ and $\bar{\pi}[0]$ and $\bar{\pi} \downarrow$ are non-copy states), then $\bar{D}(\bar{\pi}, \cdot)$ is the distribution that D yields for path $\text{merge}(\bar{\pi})$ in $\mathcal{C}$; further, if $\bar{\pi} \downarrow = s^\alpha$ then $\text{Act}(s^\alpha) = \{\alpha\}$ and thus $\bar{D}$ chooses action α . Finally, $\bar{\mathcal{C}}$ contains paths that start in a copy-state s^α . But as $\bar{v}(s^\alpha) = 0$ for all $s^\alpha \in \mathcal{S}_{\text{cp}}$, they do not contribute any probability, independent of $\bar{D}(\bar{\pi}, \cdot)$. For such paths, as well as for invalid paths, the scheduler decision $\gamma_{\bar{\pi}}$ can be chosen arbitrary without altering the probability measure.

Based on the definition of the scheduler D , we are now going to prove that the probability measure that $\bar{D}$ induces on $\bar{\mathcal{C}}$ for the event $\text{extend}(\Pi)$ is equal to the probability of Π in $\mathcal{C}$ under scheduler D .

Therefore, consider a measurable base $B \in \mathfrak{F}_{\text{Paths}^n}$ of the form $B = S_0 \times A_0 \times T_0 \times \dots \times S_n$ in $\mathcal{C}$. Then B corresponds to the set $\text{extend}(B)$ of paths in $\bar{\mathcal{C}}$. As $\text{extend}(B)$ contains paths of different lengths, we resort to its induced (infinite) cylinder $\text{Cyl}(\text{extend}(B))$ and prove that its probability equals that of B . To clarify notation, note that we use $\text{Cyl}(B^n) = B_n$ to denote the infinite cylinder $B_n \subseteq \text{Paths}^\omega$ that is induced by a finite-dimensional base $B^n \subseteq \text{Paths}^n$ (cf. Sec. 2.5.4 on page 49).

Lemma 4.4 (Measure preservation under local uniformization). *Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ be a CTMDP, D a GM-scheduler on $\mathcal{C}$ and $B = S_0 \times A_0 \times T_0 \times \dots \times S_n \in \mathfrak{F}_{\text{Paths}^n(\mathcal{C})}$. Further, let $\bar{\mathcal{C}} = (\bar{\mathcal{S}}, \text{Act}, \bar{\mathbf{R}}, \bar{\nu})$ be the locally uniform CTMDP induced by $\mathcal{C}$. Then there exists a GM-scheduler $\bar{D}$ such that*

$$\text{Pr}_{\nu, D}^n(B) = \bar{\text{Pr}}_{\bar{\nu}, \bar{D}}^\omega(\text{Cyl}(\text{extend}(B))), \quad (4.5)$$

where $\bar{\text{Pr}}_{\bar{\nu}, \bar{D}}^\omega$ is the probability measure induced by $\bar{D}$ and $\bar{\nu}$ on $\mathfrak{F}_{\text{Paths}^\omega(\bar{\mathcal{C}})}$.

Proof. To shorten notation, let $\bar{B} = \text{extend}(B)$ and $\bar{\mathcal{C}} = \text{Cyl}(\bar{B})$. We prove the claim by induction on the length n of the measurable base B :

In the induction base ($n = 0$), it holds that $B = S_0$. Therefore $\text{Pr}_{\nu, D}^0(B) = \sum_{s \in B} \nu(s) = \sum_{s \in \bar{B}} \bar{\nu}(s) = \bar{\text{Pr}}_{\bar{\nu}, \bar{D}}^0(\bar{B}) = \bar{\text{Pr}}_{\bar{\nu}, \bar{D}}^\omega(\bar{\mathcal{C}})$ and Eq. (4.5) follows.

In the induction step, we extend B with a set of initial path prefixes $I = S_0 \times A_0 \times T_0$ (see Def. 3.16 on page 82) of length one and consider the base $I \times B$ which contains paths

of length $n + 1$:

$$\begin{aligned}
Pr_{v,D}^{n+1}(I \times B) &= \int_I Pr_{v_i,D_i}^n(B) \mu_{v,D}^1(di) & (* \text{ by Lemma 3.7 } *) \\
&= \int_I \overline{Pr}_{\overline{v_i},\overline{D_i}}^\omega(\overline{C}) \mu_{v,D}^1(di) & (* \text{ by ind. hyp. } *) \\
&= \sum_{s \in S_0} v(s) \sum_{\alpha \in A_0} D(s, \alpha) \int_{T_0} \overline{Pr}_{\overline{v_i},\overline{D_i}}^\omega(\overline{C}) \eta_{E(s,\alpha)}(dt) & (* \text{ where } i = (s, \alpha, t) *) \\
&= \sum_{s \in S_0} \overline{v}(s) \sum_{\alpha \in A_0} \overline{D}(s, \alpha) \int_{T_0} \underbrace{\overline{Pr}_{\overline{v_i},\overline{D_i}}^\omega(\overline{C})}_{f(s,\alpha,t)} \eta_{E(s,\alpha)}(dt). & (* \text{ by Def. of } \overline{v}, \overline{D} *)
\end{aligned}$$

The probabilities $\overline{Pr}_{\overline{v_i},\overline{D_i}}^\omega(\overline{C})$ define a measurable function $f(s, \alpha, \cdot) : \mathbb{R}_{\geq 0} \rightarrow [0, 1]$ where $f(s, \alpha, t) = \overline{Pr}_{\overline{v_i},\overline{D_i}}^\omega(\overline{C})$ if $i = (s, \alpha, t)$. Therefore, we can apply Lemma 4.3 and obtain

$$\begin{aligned}
Pr_{v,D}^{n+1}(I \times B) &= \sum_{s \in S_0} \overline{v}(s) \sum_{\alpha \in A_0} \overline{D}(s, \alpha) \cdot \left[\overline{\mathbf{P}}(s, \alpha, \mathcal{S}) \int_{T_0} f(s, \alpha, t) \eta_{\overline{E}(s,\alpha)}(dt) \right. \\
&\quad \left. + \overline{\mathbf{P}}(s, \alpha, s^\alpha) \int_{\mathbb{R}_{\geq 0}} \eta_{\overline{E}(s,\alpha)}(dt_1) \int_{T_0 \ominus t_1} f(s, \alpha, t_1 + t_2) \eta_{\overline{E}(s^\alpha,\alpha)}(dt_2) \right]. \tag{4.6}
\end{aligned}$$

To rewrite this further, note that any path prefix $i = (s, \alpha, t)$ in $\mathcal{C}$ induces the sets of path prefixes $\overline{I}_1(i) = \{s \xrightarrow{\alpha,t}\}$ and $\overline{I}_2(i) = \{s \xrightarrow{\alpha,t_1} s^\alpha \xrightarrow{\alpha,t_2} \mid t_1 + t_2 = t\}$ in $\overline{\mathcal{C}}$, where the set $\overline{I}_1(i)$ corresponds to those path prefixes that reach a state in $\mathcal{S}$ directly, whereas the prefixes that are contained in the set $\overline{I}_2(i)$ take the detour via a copy-state s^α to a state in $\mathcal{S}$.

As defined in Lemma 3.7, $v_i(s') = \mathbf{P}(s, \alpha, s')$ is the probability to go to state s' when moving along prefix i in $\mathcal{C}$. Similarly, for $\overline{\mathcal{C}}$ we define $\overline{v}_i(s')$ as the probability to be in state $s' \in \mathcal{S}$ after a path prefix $\overline{i} \in \overline{I}_1(i) \cup \overline{I}_2(i)$: If $\overline{i} \in \overline{I}_1(i)$ then we move to a state $s' \in \mathcal{S}$ directly and do not visit a copy-state s^α . Thus, $\overline{v}_i(s') = \overline{\mathbf{P}}(s, \alpha, s')$ for $\overline{i} \in \overline{I}_1(i)$. Further, $\mathbf{P}(s, \alpha, s')$ in $\mathcal{C}$ equals the conditional probability $\frac{\overline{\mathbf{P}}(s,\alpha,s')}{\overline{\mathbf{P}}(s,\alpha,\mathcal{S})}$ to enter s' in $\overline{\mathcal{C}}$ given that we move there directly. Therefore, if $\overline{i} \in \overline{I}_1(i)$, it holds that $\overline{v}_i(s') = \overline{\mathbf{P}}(s, \alpha, s') = \overline{\mathbf{P}}(s, \alpha, \mathcal{S}) \cdot v_i(s')$.

If instead $\overline{i} \in \overline{I}_2(i)$, then $\overline{i}$ has the form $s \xrightarrow{\alpha,t_1} s^\alpha \xrightarrow{\alpha,t_2}$; hence, the transition from state s to the copy-state s^α has already been taken. Therefore $\overline{v}_i(s') = \overline{\mathbf{P}}(s^\alpha, \alpha, s')$ is the probability to end up in state s' when leaving the copy-state s^α . By the definition of s^α , this is equal to the probability to move from state s to state s' in $\mathcal{C}$ directly. Hence $\overline{v}_i(s') = v_i(s')$ if $\overline{i} \in \overline{I}_2(i)$.

As defined in Lemma 3.7, $D_i(\pi, \cdot) = D(i \circ \pi, \cdot)$ and $\overline{D}_i(\overline{\pi}, \cdot) = \overline{D}(\overline{i} \circ \overline{\pi}, \cdot)$. From the definition of $\overline{D}$, we obtain that $D_i(\pi, \cdot) = \overline{D}_i(\overline{\pi}, \cdot)$ for all $\overline{i} \in \overline{I}_1(i) \cup \overline{I}_2(i)$ and $\overline{\pi} \in \text{extend}(\pi)$. Hence, it follows that if $i = (s, \alpha, t)$ and $\overline{i} \in \overline{I}_1(i) \cup \overline{I}_2(i)$ it holds

$$\overline{Pr}_{\overline{v_i},\overline{D_i}}^\omega(\overline{C}) = \begin{cases} \overline{\mathbf{P}}(s, \alpha, \mathcal{S}) \cdot \overline{Pr}_{\overline{v_i},\overline{D_i}}^\omega(\overline{C}) & \text{if } \overline{i} \in \overline{I}_1(i) \\ \overline{Pr}_{\overline{v_i},\overline{D_i}}^\omega(\overline{C}) & \text{if } \overline{i} \in \overline{I}_2(i). \end{cases} \tag{4.7}$$

With these remarks, we can rewrite Eq. (4.6) further. Therefore, note that the first summand in Eq. (4.6) corresponds to the set $\bar{I}_1(s, \alpha, t)$ whereas the second summand corresponds to the set $\bar{I}_2(s, \alpha, t_1 + t_2)$. If we apply Eq. (4.7) to the right-hand side of Eq. (4.6), we obtain

$$\begin{aligned} Pr_{v,D}^{n+1}(I \times B) &= \sum_{s \in S_0} \bar{v}(s) \sum_{\alpha \in A_0} \bar{D}(s, \alpha) \int_{T_0} \bar{Pr}_{\bar{v}, \bar{D}}^\omega(\bar{C}) \eta_{\bar{E}(s, \alpha)}(dt) \\ &\quad + \sum_{s \in S_0} \bar{v}(s) \sum_{\alpha \in A_0} \bar{D}(s, \alpha) \cdot \bar{P}(s, \alpha, s^\alpha) \\ &\quad \cdot \int_{\mathbb{R}_{\geq 0}} \eta_{\bar{E}(s, \alpha)}(dt_1) \int_{T_0 \oplus t_1} \bar{Pr}_{\bar{v}, \bar{D}}^\omega(\bar{C}) \eta_{\bar{E}(s^\alpha, \alpha)}(dt_2). \end{aligned}$$

Applying Def. 3.16 allows us to integrate over the sets of path prefixes $\bar{I}_1 = \bigcup_{i \in I} \bar{I}_1(i)$ and $\bar{I}_2 = \bigcup_{i \in I} \bar{I}_2(i)$ which are induced by $I = S_0 \times A_0 \times T_0$ and to obtain

$$\begin{aligned} Pr_{v,D}^{n+1}(I \times B) &= \int_{\bar{I}_1} \bar{Pr}_{\bar{v}, \bar{D}}^\omega(\bar{C}) \bar{\mu}_{\bar{v}, \bar{D}}^1(d\bar{i}) + \int_{\bar{I}_2} \bar{Pr}_{\bar{v}, \bar{D}}^\omega(\bar{C}) \bar{\mu}_{\bar{v}, \bar{D}}^2(d\bar{i}) \\ &= \bar{Pr}_{\bar{v}, \bar{D}}^\omega(\bar{I}_1 \times \bar{C}) + \bar{Pr}_{\bar{v}, \bar{D}}^\omega(\bar{I}_2 \times \bar{C}) \\ &= \bar{Pr}_{\bar{v}, \bar{D}}^\omega(\bar{I} \times \bar{C}) \\ &= \bar{Pr}_{\bar{v}, \bar{D}}^\omega(\bar{I} \times \text{Cyl}(\text{extend}(B))) \\ &= \bar{Pr}_{\bar{v}, \bar{D}}^\omega(\text{Cyl}(\text{extend}(I \times B))). \end{aligned}$$

In this way, the equality $Pr_{v,D}^{n+1}(I \times B) = \bar{Pr}_{\bar{v}, \bar{D}}^\omega(\text{Cyl}(\text{extend}(I \times B)))$ follows, completing the induction step. $\square$

Lemma 4.4 holds for all measurable rectangles $B = S_0 \times A_0 \times T_0 \times \dots \times S_n$; however, we aim at an extension to arbitrary measurable bases $B \in \mathfrak{F}_{\text{Paths}^n(\mathcal{C})}$. To achieve this, we follow the standard arguments in measure theory (cf. Sec. 2.5.4). In essence, we construct a monotone class and use the monotone class theorem to extend our result from the field of finite disjoint unions of measurable rectangles to the class of measurable bases. As the proof technique is interesting in itself, we provide the details here for completeness:

First, let $\mathfrak{G}_{\text{Paths}^n(\mathcal{C})}$ be the class of all finite disjoint unions of measurable rectangles. Then each element of $\mathfrak{G}_{\text{Paths}^n(\mathcal{C})}$ has the form $B_1 \sqcup B_2 \sqcup \dots \sqcup B_n$ with each B_i being a measurable rectangle as defined above. By Lemma 2.10 (cf. page 43), we know that the set $\mathfrak{G}_{\text{Paths}^n(\mathcal{C})}$ forms a *field*.

Lemma 4.5. *Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, v)$ be a CTMDP, D a GM-scheduler on $\mathcal{C}$ and $n \in \mathbb{N}$. Further, let $\bar{\mathcal{C}} = (\bar{\mathcal{S}}, \text{Act}, \bar{\mathbf{R}}, \bar{v})$ be the locally uniform CTMDP induced by $\mathcal{C}$ and let $\bar{D}$ be the scheduler that corresponds to D . Then it holds for all $B \in \mathfrak{G}_{\text{Paths}^n(\mathcal{C})}$:*

$$Pr_{v,D}^n(B) = \bar{Pr}_{\bar{v}, \bar{D}}^\omega(\text{Cyl}(\text{extend}(B))).$$

Proof. As $B \in \mathfrak{G}_{Paths^n(C)}$, it has the form $B = \bigcup_{i=1}^k B_i$ for pairwise disjoint measurable rectangles B_i of length n . Thus

$$\begin{aligned}
 Pr_{v,D}^n(B) &= Pr_{v,D}^n\left(\bigcup_{i=1}^k B_i\right) = \sum_{i=1}^k Pr_{v,D}^n(B_i) & (* \text{ as } B_i \cap B_j = \emptyset \text{ for } i \neq j *) \\
 &= \sum_{i=1}^k \overline{Pr}_{\overline{v},\overline{D}}^\omega\left(Cyl(extend(B_i))\right) & (* \text{ by Lemma 4.4 } *) \\
 &= \overline{Pr}_{\overline{v},\overline{D}}^\omega\left(\bigcup_{i=1}^k Cyl(extend(B_i))\right) & (* \text{ by Lemma 4.1(2) } *) \\
 &= \overline{Pr}_{\overline{v},\overline{D}}^\omega\left(Cyl(extend(B))\right). & (* \text{ by Lemma 4.1(3) } *) \quad \square
 \end{aligned}$$

With the monotone class theorem (Thm. 2.2 on page 22), the preservation property extends from $\mathfrak{G}_{Paths^n}$ to the σ -field $\mathfrak{F}_{Paths^n}$: Here, the definition of a monotone class (cf. Def. 2.5 in Sec. 2.1.2) is applied to a class of subsets of $Paths^n$: A class $\mathfrak{C}$ of subsets of $Paths^n$ is a monotone class iff it is closed under in- and decreasing sequences: if $\Pi_k \in \mathfrak{C}$ and $\Pi \subseteq Paths^n$ such that $\Pi_0 \subseteq \Pi_1 \subseteq \dots$ and $\bigcup_{k=0}^\infty \Pi_k = \Pi$, we write $\Pi_k \uparrow \Pi$ (similarly for $\Pi_k \downarrow \Pi$). Then $\mathfrak{C}$ is a monotone class iff for all $\Pi_k \in \mathfrak{C}$ and $\Pi \subseteq Paths^n$ with $\Pi_k \uparrow \Pi$ or $\Pi_k \downarrow \Pi$ it holds that $\Pi \in \mathfrak{C}$.

Lemma 4.6 (Monotone class). *Let $\mathcal{C} = (\mathcal{S}, Act, \mathbf{R}, v)$ be a CTMDP with GM-scheduler D ; further, let $\overline{\mathcal{C}} = (\overline{\mathcal{S}}, Act, \overline{\mathbf{R}}, \overline{v})$ be $\mathcal{C}$'s induced locally uniform CTMDP and $\overline{D} \in GM(\overline{\mathcal{C}})$ the scheduler induced by D . The set*

$$\mathfrak{C} = \left\{ B \in \mathfrak{F}_{Paths^n(C)} \mid Pr_{v,D}^n(B) = \overline{Pr}_{\overline{v},\overline{D}}^\omega\left(Cyl(extend(B))\right) \right\}$$

is a monotone class.

Proof. We consider increasing and decreasing sequences of sets of paths in $\mathfrak{C}$:

- Assume that $\Pi_i^n \in \mathfrak{C}$ for $i = 1, 2, \dots$, and that the sets Π_i^n form an increasing sequence that converges from below to the limit Π^n , that is, $\Pi_i^n \uparrow \Pi^n$. The fact that σ -fields are closed under limits and that $\Pi_i^n \in \mathfrak{F}_{Paths^n(C)}$ for all $i = 1, 2, \dots$ implies that $\Pi^n \in \mathfrak{F}_{Paths^n(C)}$. Therefore, it remains to show that

$$Pr_{v,D}^n(\Pi^n) = \overline{Pr}_{\overline{v},\overline{D}}^\omega\left(Cyl(extend(\Pi^n))\right).$$

By definition of $\mathfrak{C}$, $Pr_{v,D}^n(\Pi_i^n) = \overline{Pr}_{\overline{v},\overline{D}}^\omega\left(Cyl(extend(\Pi_i^n))\right)$ for all $i \in \mathbb{N}$. Therefore, the limits also agree. More precisely, we have established that

$$\lim_{i \rightarrow \infty} Pr_{v,D}^n(\Pi_i^n) = \lim_{i \rightarrow \infty} \overline{Pr}_{\overline{v},\overline{D}}^\omega\left(Cyl(extend(\Pi_i^n))\right). \quad (4.8)$$

Further, both $Pr_{v,D}^n$ and $\overline{Pr}_{\overline{v},\overline{D}}^\omega$ are measures on $\mathfrak{F}_{Paths^n(\mathcal{C})}$ and $\mathfrak{F}_{Paths^\omega(\overline{\mathcal{C}})}$, respectively. As $\Pi_i^n \uparrow \Pi^n$ is an increasing sequence, it follows by Lemma 4.1 and the definition of Cyl that also $Cyl(\text{extend}(\Pi_{n,i})) \uparrow Cyl(\text{extend}(\Pi_n))$.

From here, we obtain by Lemma 2.2 that

$$\lim_{i \rightarrow \infty} Pr_{v,D}^n(\Pi_i^n) = Pr_{v,D}^n(\Pi^n) \quad \text{and} \quad (4.9)$$

$$\lim_{i \rightarrow \infty} \overline{Pr}_{\overline{v},\overline{D}}^\omega(Cyl(\text{extend}(\Pi_i^n))) = \overline{Pr}_{\overline{v},\overline{D}}^\omega(Cyl(\text{extend}(\Pi^n))). \quad (4.10)$$

Thus, we have proved that

$$\begin{aligned} Pr_{v,D}^n(\Pi^n) &\stackrel{(4.9)}{=} \lim_{i \rightarrow \infty} Pr_{v,D}^n(\Pi_i^n) \stackrel{(4.8)}{=} \lim_{i \rightarrow \infty} \overline{Pr}_{\overline{v},\overline{D}}^\omega(Cyl(\text{extend}(\Pi_i^n))) \\ &\stackrel{(4.10)}{=} \overline{Pr}_{\overline{v},\overline{D}}^\omega(Cyl(\text{extend}(\Pi^n))). \end{aligned}$$

- Now, let $\Pi_i^n \in \mathfrak{C}$ and $\Pi_i^n \downarrow \Pi^n$. This case is analogue, as $\lim_{i \rightarrow \infty} Pr_{v,D}^n(\Pi_i^n) = Pr_{v,D}^n(\Pi^n)$ also holds for decreasing sequences $\Pi_i^n \downarrow \Pi^n$. Hence, the proof goes along the same lines as the one done before for increasing sequences. $\square$

Lemma 4.7 (Extension). *Let $\mathcal{C} = (\mathcal{S}, Act, \mathbf{R}, v)$ be a CTMDP, D a GM-scheduler on $\mathcal{C}$ and $n \in \mathbb{N}$. Further, let $\overline{\mathcal{C}} = (\overline{\mathcal{S}}, Act, \overline{\mathbf{R}}, \overline{v})$ be the locally uniform CTMDP induced by $\mathcal{C}$, and $\overline{D} \in GM(\overline{\mathcal{C}})$ the scheduler that corresponds to D . Then it holds for all measurable bases $B \in \mathfrak{F}_{Paths^n(\mathcal{C})}$ that*

$$Pr_{v,D}^n(B) = \overline{Pr}_{\overline{v},\overline{D}}^\omega(Cyl(\text{extend}(B))).$$

Proof. By Lemma 4.6, $\mathfrak{C}$ is a monotone class and by Lemma 4.5 it follows that $\mathfrak{G}_{Paths^n(\mathcal{C})} \subseteq \mathfrak{C}$. Thus, the monotone class theorem (cf. Thm. 2.2) applies and $\mathfrak{F}_{Paths^n} \subseteq \mathfrak{C}$. Hence $Pr_{v,D}^n(B) = \overline{Pr}_{\overline{v},\overline{D}}^\omega(Cyl(\text{extend}(B)))$ for all $B \in \mathfrak{F}_{Paths^n}$. $\square$

Lemma 4.4 and its measure-theoretic extension to the σ -field are the basis for the major results of this chapter. We discuss them in the following section.

4.3 Preservation results for local uniformization

The first result states the correctness of the construction of the scheduler $\overline{D}$, that is, it asserts that D and $\overline{D}$ assign the same probability to corresponding sets of paths.

Theorem 4.1. *Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ be a CTMDP and D a GM-scheduler on $\mathcal{C}$. Further, let $\bar{\mathcal{C}} = (\bar{\mathcal{S}}, \text{Act}, \bar{\mathbf{R}}, \bar{\nu})$ be the induced locally uniform CTMDP and $\bar{D}$ the scheduler that corresponds to D . Then it holds for all $\Pi \in \mathfrak{F}_{\text{Paths}^\omega}$ that*

$$Pr_{\nu, D}^\omega(\Pi) = \overline{Pr}_{\bar{\nu}, \bar{D}}^\omega(\text{extend}(\Pi)).$$

Proof. Each cylinder $\Pi \in \mathfrak{F}_{\text{Paths}^\omega(\mathcal{C})}$ is induced by a measurable base [ADD00, Thm. 2.7.2]; hence $\Pi = \text{Cyl}(B)$ for some $B \in \mathfrak{F}_{\text{Paths}^n(\mathcal{C})}$ and $n \in \mathbb{N}$. But then, $Pr_{\nu, D}^\omega(\Pi) = Pr_{\nu, D}^n(B)$. Further, it is easy to verify that $\text{extend}(\text{Cyl}(B)) = \text{Cyl}(\text{extend}(B))$. Thus $Pr_{\nu, D}^n(B) = \overline{Pr}_{\bar{\nu}, \bar{D}}^n(\text{extend}(B))$ by Lemma 4.7. $\square$

With Lemma 4.4 and its extension, we are now ready to prove that local uniformization does not alter the CTMDP in a way that we leak probability mass with respect to the most important scheduler classes:

Theorem 4.2. *Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ be a CTMDP and let $\bar{\mathcal{C}} = (\bar{\mathcal{S}}, \text{Act}, \bar{\mathbf{R}}, \bar{\nu})$ be its induced locally uniform CTMDP. For all $\Pi \in \mathfrak{F}_{\text{Paths}^\omega(\mathcal{C})}$ and each scheduler class $\mathfrak{D}$ from the set $\{\text{GM}, \text{TTHR}, \text{TTPR}, \text{TAHR}, \text{TAPR}\}$ it holds that*

$$\sup_{D \in \mathfrak{D}(\mathcal{C})} Pr_{\nu, D}^\omega(\Pi) \leq \sup_{D' \in \mathfrak{D}(\bar{\mathcal{C}})} \overline{Pr}_{\bar{\nu}, D'}^\omega(\text{extend}(\Pi)). \quad (4.11)$$

Proof. By Thm. 4.1, the claim follows for the class of all GM-schedulers, that is, for $\mathfrak{D} = \text{GM}$. For the other classes, it remains to check that the GM-scheduler $\bar{D}$ used in Lemma 4.4 also falls into the respective class. Here, we state the proof for TTPR: If $D : \mathcal{S} \times \mathbb{R}_{\geq 0} \rightarrow \text{Distr}(\text{Act}) \in \text{TTPR}$, define $\bar{D}(s, \Delta) = D(s, \Delta)$ if $s \in \mathcal{S}$ and $\bar{D}(s^\alpha, \Delta) = \{\alpha \mapsto 1\}$ for $s^\alpha \in \mathcal{S}_{\text{cp}}$. Then Lemma 4.4 applies verbatim. $\square$

Note that Thm. 4.2 does not mention the scheduler classes TPR and TAHOPR. This is for good reason: In Thm. 4.4, we will construct a counterexample that disproves Eq. (4.11) for these scheduler classes: Note that although we obtain a GM-scheduler $\bar{D}$ on $\bar{\mathcal{C}}$ for any $D \in \text{TPR}(\mathcal{C}) \cup \text{TAHOPR}(\mathcal{C})$ by Thm. 4.1, $\bar{D}$ is not guaranteed to lie in $\text{TPR}(\bar{\mathcal{C}})$ (or $\text{TAHOPR}(\bar{\mathcal{C}})$, respectively). Hence, Eq. (4.11) does not hold directly for all scheduler classes that are subsets of GM.

For the main result, we identify the scheduler classes, that do not gain probability mass by local uniformization:

Theorem 4.3. *Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ be a CTMDP, $\bar{\mathcal{C}} = (\bar{\mathcal{S}}, \text{Act}, \bar{\mathbf{R}}, \bar{\nu})$ its induced locally uniform CTMDP and $\Pi \in \mathfrak{F}_{\text{Paths}^\omega(\mathcal{C})}$. Then*

$$\sup_{D \in \mathfrak{D}(\mathcal{C})} Pr_{\nu, D}^\omega(\Pi) = \sup_{D' \in \mathfrak{D}(\bar{\mathcal{C}})} \bar{Pr}_{\bar{\nu}, D'}^\omega(\text{extend}(\Pi)) \quad \text{for } \mathfrak{D} \in \{TTPR, TAPR\}.$$

Proof. Theorem 4.2 proves the direction from left to right. For the reverse, let $D' \in TTPR(\bar{\mathcal{C}})$ and define $D \in TTPR(\mathcal{C})$ such that $D(s, \Delta) = D'(s, \Delta)$ for all $s \in \mathcal{S}, \Delta \in \mathbb{R}_{\geq 0}$. Then $\bar{D} = D'$ and $\bar{Pr}_{\bar{\nu}, D'}^\omega(\text{extend}(\Pi)) = Pr_{\nu, D}^\omega(\Pi)$ by Thm. 4.1. Hence the claim for $TTPR$ follows; analogue for $D' \in TAPR(\bar{\mathcal{C}})$. $\square$

Conjecture 4.1. *We conjecture that Thm. 4.3 also holds for GM and TTHR. For $D' \in GM(\bar{\mathcal{C}})$, we aim at defining a scheduler $D \in GM(\mathcal{C})$ that induces the same probabilities on $\mathcal{C}$. However, a history $\pi \in \text{Paths}^*(\mathcal{C})$ corresponds to the uncountable set $\text{extend}(\pi)$ in $\bar{\mathcal{C}}$ such that $D'(\bar{\pi}, \cdot)$ may be different for each $\bar{\pi} \in \text{extend}(\pi)$.*

As D can only decide once on history π , in order to mimic D' on $\bar{\mathcal{C}}$, we propose to weigh each distribution $D'(\bar{\pi}, \cdot)$ with the conditional probability of $d\bar{\pi}$ given $\text{extend}(\pi)$.

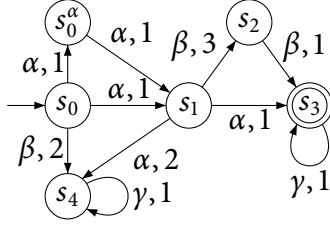
In the following, we disprove Eq. (4.11) for TPR - and $TAHOPR$ -schedulers. Intuitively, TPR -schedulers rely on the sojourn time in the last state; however, local uniformization changes the exit rates of states by adding transitions to copy-states.

Theorem 4.4. *For $\mathfrak{G} \in \{TPR, TAHOPR\}$, there exists a CTMDP $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ and a measurable set of paths $\Pi \in \mathfrak{F}_{\text{Paths}^\omega(\mathcal{C})}$ such that*

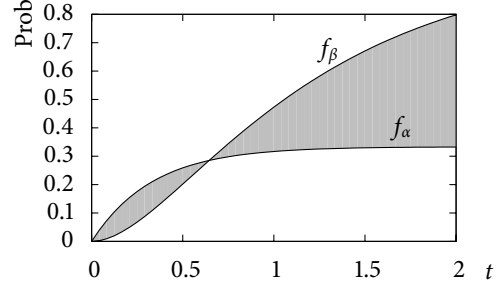
$$\sup_{D \in \mathfrak{G}(\mathcal{C})} Pr_{\nu, D}^\omega(\Pi) > \sup_{D' \in \mathfrak{G}(\bar{\mathcal{C}})} \bar{Pr}_{\bar{\nu}, D'}^\omega(\text{extend}(\Pi)).$$

Proof. We give the proof for TPR : Consider the CTMDPs $\mathcal{C}$ and $\bar{\mathcal{C}}$ in Fig. 4.2(a) and Fig. 4.5(a), respectively.

Let $\Pi \in \mathfrak{F}_{\text{Paths}^\omega(\mathcal{C})}$ be the set of paths in $\mathcal{C}$ that reach state s_3 in 1 time unit and let $\bar{\Pi} = \text{extend}(\Pi)$. To optimize $Pr_{\nu, D}^\omega(\Pi)$ and $\bar{Pr}_{\bar{\nu}, D'}^\omega(\bar{\Pi})$, any scheduler D (resp. D') must choose $\{\alpha \mapsto 1\}$ in state s_0 . Nondeterminism only remains in state s_1 ; here, the optimal distribution over $\{\alpha, \beta\}$ depends on the time t_0 that was spent to reach state s_1 : In $\mathcal{C}$ and $\bar{\mathcal{C}}$, the probability to go from s_1 to s_3 in the remaining $t = 1 - t_0$ time units is $f_\alpha(t) = \frac{1}{3} - \frac{1}{3}e^{-3t}$ for α and $f_\beta(t) = 1 + \frac{1}{2}e^{-3t} - \frac{3}{2}e^{-t}$ for β . Fig. 4.5(b) shows the cumulative distribution functions (cdfs) of f_α and f_β ; as any convex combination of α and β results in a cdf in the shaded area of Fig. 4.5(b), we only need to consider the extreme distributions $\{\alpha \mapsto 1\}$ and



(a) Local uniformization of Fig. 4.2(a).

(b) From state s_1 to state s_3 .Figure 4.5: Timed reachability of state s_3 (starting in s_1) in $\mathcal{C}$ and $\bar{\mathcal{C}}$.

$\{\beta \mapsto 1\}$ for maximal reachability. Let d be the unique solution (in $\mathbb{R}_{>0}$) of $f_\alpha(t) = f_\beta(t)$, i.e. the point where the two cdfs cross. Then $D_{opt}(s_0 \xrightarrow{\alpha, t_0} s_1, \cdot) = \{\alpha \mapsto 1\}$ if $1 - t_0 \leq d$ and $\{\beta \mapsto 1\}$ otherwise, is an optimal GM-scheduler for Π on $\mathcal{C}$ and $D_{opt} \in TPR(\mathcal{C}) \cap TTPR(\mathcal{C})$ as it depends only on the delay of the last transition.

For $\bar{\Pi}$, D' is an optimal GM-scheduler on $\bar{\mathcal{C}}$ if $D'(s_0 \xrightarrow{\alpha, t_0} s_1, \cdot) = D_{opt}(s_0 \xrightarrow{\alpha, t_0} s_1, \cdot)$ as before and $D'(s_0 \xrightarrow{\alpha, t_0} s_0^\alpha \xrightarrow{\alpha, t_1} s_1, \cdot) = \{\alpha \mapsto 1\}$ if $1 - t_0 - t_1 \leq d$ and $\{\beta \mapsto 1\}$ otherwise. Note that by definition, $D' = \overline{D_{opt}}$ and $\overline{D_{opt}} \in TTPR(\bar{\mathcal{C}})$, whereas $D' \notin TPR(\bar{\mathcal{C}})$ as any $TPR(\bar{\mathcal{C}})$ scheduler is independent of t_0 . For history $\pi = s_0 \xrightarrow{\alpha, t_0} s_0^\alpha \xrightarrow{\alpha, t_1} s_1$, the best approximation of t_0 is the expected sojourn time in state s_0 , i.e. $\frac{1}{E(s_0, \alpha)}$. For the induced scheduler $D'' \in TPR(\bar{\mathcal{C}})$, it holds $D''(s_1, t_1) \neq D'(s_0 \xrightarrow{\alpha, t_0} s_0^\alpha \xrightarrow{\alpha, t_1} s_1)$ almost surely. But as $\overline{D_{opt}}$ is optimal, there exists $\varepsilon > 0$ such that $\overline{Pr}_{\bar{v}, D''}^\omega(\bar{\Pi}) = \overline{Pr}_{\bar{v}, \overline{D_{opt}}}^\omega(\bar{\Pi}) - \varepsilon$. Therefore

$$\sup_{D'' \in TPR(\bar{\mathcal{C}})} \overline{Pr}_{\bar{v}, D''}^\omega(\bar{\Pi}) < \overline{Pr}_{\bar{v}, \overline{D_{opt}}}^\omega(\bar{\Pi}) = Pr_{v, D_{opt}}^\omega(\Pi) = \sup_{D \in TPR(\mathcal{C})} Pr_{v, D}^\omega(\Pi).$$

For *TAHOPR*, a similar proof applies that relies on the fact that local uniformization changes the number of transitions needed to reach a goal state. $\square$

This proves that by local uniformization, essential information for *TP* and *TAHOPR* schedulers is lost. In other cases, schedulers from *TAHR* and *TAHOPR* gain information by local uniformization:

Theorem 4.5. *There exists a CTMDP $\mathcal{C} = (\mathcal{S}, Act, \mathbf{R}, v)$ and a set of paths $\Pi \in \mathfrak{F}_{Paths^\omega(\mathcal{C})}$ such that*

$$\sup_{D \in \mathfrak{G}(\mathcal{C})} Pr_{v, D}^\omega(\Pi) < \sup_{D' \in \mathfrak{G}(\bar{\mathcal{C}})} \overline{Pr}_{\bar{v}, D'}^\omega(\text{extend}(\Pi)) \quad \text{for } \mathfrak{G} = \{\text{TAHR}, \text{TAHOPR}\}.$$

Proof. Consider the CTMDPs $\mathcal{C}$ and $\bar{\mathcal{C}}$ in Fig. 4.2(a) and Fig. 4.5(a), resp. Let Π be the time-bounded reachability property of state s_3 within 1 time unit and let $\bar{\Pi} = \text{extend}(\Pi)$. We prove the claim for TAHR: Therefore, we derive $D \in \text{TAHR}(\mathcal{C})$ such that $Pr_{v,D}^\omega(\Pi) = \sup_{D' \in \text{TAHR}(\mathcal{C})} Pr_{v,D'}^\omega(\Pi)$. For this, $D(s_0) = \{\alpha \mapsto 1\}$ must obviously hold. Thus, the only nondeterministic choice occurs in state s_1 for time-abstract history $s_0 \xrightarrow{\alpha} s_1$ where $D(s_0 \xrightarrow{\alpha} s_1) = \mu$, $\mu \in \text{Distr}(\{\alpha, \beta\})$. For initial state s_0 , Fig. 4.6(a) depicts $Pr_{v,D}^\omega(\Pi)$ for all $\mu \in \text{Distr}(\{\alpha, \beta\})$; obviously, $D(s_0 \xrightarrow{\alpha} s_1) = \{\beta \mapsto 1\}$ maximizes $Pr_{v,D}^\omega(\Pi)$. On $\bar{\mathcal{C}}$, we prove that there exists $D' \in \text{TAHR}(\bar{\mathcal{C}})$ such that $Pr_{v,D}^\omega(\Pi) < \bar{Pr}_{v,D'}^\omega(\bar{\Pi})$: To maximize $\bar{Pr}_{v,D'}^\omega(\bar{\Pi})$, define $D'(s_0) = \{\alpha \mapsto 1\}$. Note that D' may yield different distributions for the time-abstract paths $s_0 \xrightarrow{\alpha} s_1$ and $s_0 \xrightarrow{\alpha} s_0^\alpha \xrightarrow{\alpha} s_1$; for $\mu, \mu_c \in \text{Distr}(\{\alpha, \beta\})$ such that $\mu = D'(s_0 \xrightarrow{\alpha} s_1)$ and $\mu_c = D'(s_0 \xrightarrow{\alpha} s_0^\alpha \xrightarrow{\alpha} s_1)$ the probability of $\bar{\Pi}$ under D' is depicted in Fig. 4.6(b) for all $\mu, \mu_c \in \text{Distr}(\{\alpha, \beta\})$. Clearly, $\bar{Pr}_{v,D'}^\omega(\bar{\Pi})$ is maximal if $D'(s_0 \xrightarrow{\alpha} s_1) = \{\beta \mapsto 1\}$ and $D'(s_0 \xrightarrow{\alpha} s_0^\alpha \xrightarrow{\alpha} s_1) = \{\alpha \mapsto 1\}$. Further, Fig. 4.6(b) shows that with this choice of D' , $\bar{Pr}_{v,D'}^\omega(\bar{\Pi}) > Pr_{v,D}^\omega(\Pi)$ and the claim follows. For TAHOPR, the proof applies analogously. $\square$

With these counterexamples, we complete our discussion of local uniformization and come back to the question that was raised at the beginning of Sec. 4.2: The motivation to study locally uniform CTMDPs is to delay the scheduling decision until the current state is left.

As we have seen, for TTPR- and TAPR- schedulers, any given CTMDP can be transformed into a locally uniform one while preserving all measures. Moreover, in this thesis, we are particularly interested in time-bounded reachability objectives; for them, we know that TTPR schedulers are sufficient, that is, we do not need to consider any other class of schedulers to obtain the optimal reachability probabilities.

However, a word of caution is necessary at this point: The results of this chapter might lead to the conclusion, that for time-bounded reachability objectives, one can transform an arbitrary CTMDP into a locally uniform one and investigate it with respect to late schedulers. Albeit possible, there is still an open theoretical problem in this approach:

The results of this chapter do not prove in any way, that local uniformization preserves measures with respect to late schedulers. Obviously, for such a proof, we need to define the semantics of non-locally uniform CTMDPs under late schedulers. However, in this setting, the scheduling decision and the sojourn time distribution become dependent on each other. The natural result are measurable schedulers that decide continuously during the sojourn in the current state. However, the implications of such a definition are ongoing research and outside the scope of this thesis.

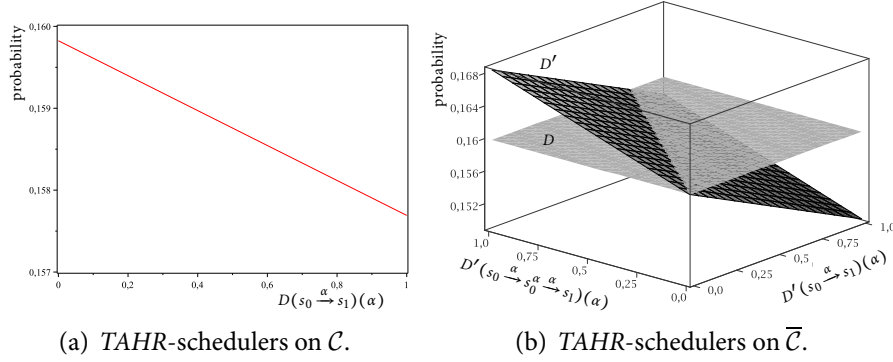


Figure 4.6: Optimal TAHR-schedulers for time-bounded reachability.

4.4 Delaying nondeterministic choices

In this section, we finally discuss late schedulers in more detail. As stated previously, we therefore have to assume that the given CTMDP is locally uniform.

In the following, we show how local uniformity permits to derive the class of *late schedulers* which resolve the nondeterministic choices in the current state only upon leaving that state. Intuitively, a late scheduler may exploit information about the current state's sojourn time for its decision. As a consequence, we prove in this section that late schedulers on locally uniform CTMDPs induce more accurate probability bounds than the class of (early) GM-schedulers.

To begin, assume that $\mathcal{C} = (\mathcal{S}, Act, \mathbf{R}, \nu)$ is a locally uniform CTMDP and D is a GM-scheduler on $\mathcal{C}$. Then $E(s, \alpha) = u(s)$ for all $s \in \mathcal{S}$ and $\alpha \in Act$ (cf. Def. 4.2). This independence of the exit-rate from the action that is chosen implies that the measures $\eta_{E(s, \alpha)}$ in the integral in Def. 3.14 do not depend on α . Thus, we may exchange the order of integration in Eq. (3.11) by applying [ADD00, Thm. 2.6.6]. More precisely, we can rewrite the measure on combined transitions given in Def. 3.14 (see on page 79) to account for the fact that the sojourn time distribution becomes independent from the scheduler. Hence, for locally uniform CTMDPs and late schedulers, the measure $\mu_D(\pi, M)$ as defined in Eq. (3.11) can be restated as follows:

$$\mu_D(\pi, M) = \int_{\mathbb{R}_{\geq 0}} \eta_{u(\pi \downarrow)}(dt) \int_{Act} D(\pi, d\alpha) \int_{\mathcal{S}} \mathbf{I}_M(\alpha, t, s') \mathbf{P}(s, \alpha, ds'). \quad (4.12)$$

Formally, Eq.(4.12) now permits to define *late schedulers* as measurable mappings $D : Paths^*(\mathcal{C}) \times \mathbb{R}_{\geq 0} \times \mathfrak{F}_{Act} \rightarrow [0, 1]$ that extend the class of GM-schedulers by also considering the sojourn time in the current state, that is, in state $\pi \downarrow$. Formally, the class of late schedulers (denoted ML) is defined as the set of all measurable mappings $Paths^*(\mathcal{C}) \times \mathbb{R}_{\geq 0} \times \mathfrak{F}_{Act} \rightarrow [0, 1]$ which satisfy $D(\pi, t, \cdot) \in Distr(Act(\pi \downarrow))$ for all $t \in \mathbb{R}_{\geq 0}$ and $\pi \in Paths^*$.

The details of the adaptation of the probability measures to late schedulers are discussed in Chapter 5 (see also Def. 5.1 on page 116), where we develop an approximation

algorithm which computes the maximum time-bounded reachability probabilities in locally uniform CTMDPs under late schedulers.

Note however, that local uniformity is essential for the derivation of late schedulers: In the general case, the measures $\eta_{E(s,\alpha)}(dt)$ and a late scheduler $D(\pi, t, d\alpha)$ are inter-dependent in t and α ; hence, in Def. 3.14, $\mu_D(\pi, \cdot)$ is not well-defined for late-schedulers. Intuitively, in general CTMDPs the sojourn time t of the current state s depends on D while D depends on t .

Let ML and GM denote the classes of late and GM -schedulers, respectively.

Theorem 4.6 (Comparison of early and late schedulers). *Let GM and ML denote the classes of early and late schedulers. Further, let $\mathcal{C} = (\mathcal{S}, Act, \mathbf{R}, v)$ be a locally uniform CTMDP. Then it holds for all $\Pi \in Paths^\omega(\mathcal{C})$ that*

$$\sup_{D \in GM} Pr_{v,D}^\omega(\Pi) \leq \sup_{D \in ML} Pr_{v,D}^\omega(\Pi). \quad (4.13)$$

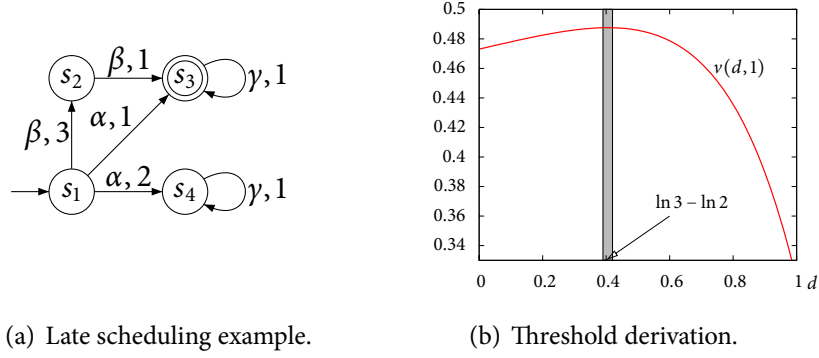
Moreover, Inequality (4.13) is strict in general.

Proof. By definition, $GM \subseteq ML$; to see this, let $D_e : Paths^*(\mathcal{C}) \times \mathfrak{F}_{Act} \rightarrow [0, 1] \in GM$ be an early scheduler and define the late scheduler $D_l : Paths^*(\mathcal{C}) \times \mathbb{R}_{\geq 0} \times \mathfrak{F}_{Act} \rightarrow [0, 1] \in ML$ such that $D_l(\pi, t, \cdot) = D_e(\pi, \cdot)$, where t is the sojourn time in $\pi \downarrow$ that is available to the ML -scheduler. With this construction, any GM -scheduler can be considered as an ML -scheduler which ignores the sojourn time in $\pi \downarrow$. Further, the probability measures induced by D_e and D_l are equal by definition. Thus, Inequality (4.13) follows directly.

Now we come to the second claim and prove that ML -schedulers generally induce strictly larger probability bounds than GM -schedulers: Let $\mathcal{C}$ be the locally uniform CTMDP depicted in Fig. 4.7(a), and let Π be the time-bounded reachability probability for state s_3 and time-bound $z = 1$. As we have seen in Ex. 4.2 on page 89, the optimal choice for an *early* scheduler if state s_1 is entered and 1 time unit remains to reach state s_3 is action β (as $1 > \ln\left(\frac{5}{8} + \frac{1}{8}\sqrt{105}\right)$, cf. Ex. 4.2). Therefore, we obtain the maximum reachability probability for early schedulers:

$$\sup_{D \in GM} Pr_{v,D}^\omega(\Pi) = \int_0^1 \left(3e^{-3t_1} \int_0^{1-t_1} e^{-t_2} dt_2 \right) dt_1 = 1 + \frac{1}{2}e^{-3} - \frac{3}{2}e^{-1} \simeq 0.4731.$$

On the other hand, the optimal late scheduler can be derived as follows: Assume that the sojourn in state s_1 lasts for t_1 time units. If the scheduler chooses α upon leaving s_1 , the CTMDP enters state s_3 with probability $\frac{1}{3}$. On the other hand, action β incurs an additional delay with rate 1, but reaches state s_3 with probability 1. We derive the minimum amount of time $d \in \mathbb{R}_{\geq 0}$ that needs to remain after the sojourn in state s_1 is over, such that the probability induced by choosing β is larger than $\frac{1}{3}$ (i.e. the probability induced by α).



(a) Late scheduling example.

(b) Threshold derivation.

Figure 4.7: Late schedulers outperform early schedulers.

Formally, we seek $d \in \mathbb{R}_{\geq 0}$ such that the *ML*-scheduler D with

$$D(s_1, t) = \begin{cases} \{\beta \mapsto 1\} & \text{if } t \leq d \\ \{\alpha \mapsto 1\} & \text{otherwise} \end{cases}$$

is optimal. For the CTMDP in Fig. 4.7(a) and a fixed $d \in \mathbb{R}_{\geq 0}$, the probability to move from state s_1 to state s_3 within z time units is given by the function v , where

$$\begin{aligned} v(d, z) &= \frac{1}{3} \int_{z-d}^z 3e^{-3t_1} dt_1 + \int_0^{z-d} \left(3e^{-3t_1} \cdot \int_0^{z-t_1} e^{-t_2} dt_2 \right) dt_1 \\ &= \frac{1}{3} \int_{z-d}^z 3e^{-3t_1} dt_1 + \int_0^{z-d} (3e^{-3t_1} - 3e^{-2t_1-z}) dt_1. \end{aligned}$$

Here, the second integral corresponds to the convolution of the delays of the transitions that lead from state s_1 via state s_2 to state s_3 . Intuitively, in the first integral, the sojourn in state s_1 falls into the interval $[z-d, z]$; hence, time is short and action α is chosen. The second integral corresponds to sojourn times $t_1 \in [0, z-d]$, where we favor β over α . To prove the claim, it suffices to consider the time horizon $z = 1$: In this case, Fig. 4.7(b) depicts the probability $v(d, 1)$ for all $0 \leq d \leq z = 1$; analytically, it is easy to derive that $v(d, 1)$ is maximal for $d = d_{\max} = \ln 3 - \ln 2$.

Hence, if the remaining amount of time $z - t_1$ after leaving state s_1 is less than $\ln 3 - \ln 2$, we choose action α ; otherwise we choose action β . This yields the scheduler $D(s_1, t_1, \cdot) = \{\beta \mapsto 1\}$ if $t_1 < 1 + \ln 2 - \ln 3$ and $\{\alpha \mapsto 1\}$, otherwise. Finally, computing the maximum achievable probability under the late scheduler D as derived above yields probability

$$Pr_{v,D}^{\omega}(\Pi) = v(d_{\max}, 1) = 1 + \frac{19}{24}e^{-3} - \frac{3}{2}e^{-1} \simeq 0.4876.$$

Hence, D induces a probability which is approximately 1.45% higher than the maximum probability that can be obtained by early schedulers. Therefore, we have proved that optimal *ML*-schedulers perform strictly better than optimal *GM*-schedulers. $\square$

4.5 Conclusion

In this chapter, we study a hierarchy of early scheduler classes for CTMDPs and investigate their sensitivity for general measures with respect to local uniformization. This transformation is shown to be measure-preserving for *TAPR* and *TTPR* schedulers. Moreover, in contrast to *TPR* and *TAHOPR* schedulers, *GM*, *TTHR* and *TAHR* schedulers cannot lose information to optimize their decisions. *TAHR* and *TAHOPR* schedulers can also gain information. We conjecture that our transformation is also measure-preserving for *TTHR* and *GM* schedulers.

The starting point for considering local uniformization was the observation that locally uniform CTMDPs separate the sojourn time distribution from the scheduler decision which allows us to define strictly more powerful scheduler classes compared to those that are proposed for general CTMDPs.

Hence, it was a natural question to investigate means to uniformize early CTMDPs. However, more research is necessary in this direction, as we did not prove that local uniformization is measure preserving for late schedulers and general CTMDPs.

Moreover, the slightly simpler structure of locally uniform CTMDPs allows us to derive an approximation algorithm that computes time bounded reachability probabilities in locally uniform CTMDPs. This will be the topic of Chapter 5.

5 The analysis of late CTMDPs

The only reason for time is so
that everything doesn't
happen at once.

(Albert Einstein)

In this chapter, we develop a discretization technique which allows us to analyze time-bounded reachability probabilities in *late* CTMDPs. As we have seen in the previous chapters, the sojourn time distribution of the current state in a CTMDP generally depends on the action that is chosen by the associated *GM*-scheduler. This dependency requires the scheduler to decide early, that is, when entering the current state. Therefore, we sometimes refer to the class of *GM*-schedulers and the associated CTMDPs as *early schedulers* and *early CTMDPs*, respectively.

In contrast to general CTMDPs and *GM*-schedulers, Chapter 4 has introduced local uniformization and motivated the use of late schedulers (from the class *ML*): More precisely, we have seen in Sec. 4.4 that for locally uniform CTMDPs *late schedulers* generally outperform the early schedulers from Sec. 3.3.2. This comes as no surprise, as in locally uniform CTMDPs, the states' sojourn time distributions do not depend on the scheduler's choice. Hence, local uniformity allows us to delay the scheduling decision until the current state is left, resulting in the class of late schedulers. However, another result of Sec. 4.4 is that late schedulers are well-defined only for locally uniform CTMDPs.

Up to now, the motivation to consider locally uniform CTMDPs and late schedulers may appear to be merely technical. However, this would be a wrong conclusion, as we will see in the forthcoming chapters that local uniformity is a property that is commonly found in controlled queuing systems (cf. the case study in Sec. 5.4 at the end of this chapter) and stochastic Petri net formalisms such as GSPNs (cf. Chapter 8). Moreover, the ideas and techniques developed in this chapter carry over to interactive Markov chains (cf. Chapter 6) whose Markovian states can be considered locally uniform. Therefore it is fair to say that the ideas presented in this chapter provide the essence of the approximations used throughout this thesis.

From a technical perspective, local uniformity is an extremely useful property when it comes to the analysis of CTMDPs. Therefore, the focus of this chapter is on the analysis of locally uniform CTMDPs under late scheduling disciplines.

Its main contribution is a solution method for the time-bounded reachability problem in locally uniform CTMDPs: We propose a technique to compute the maximum proba-

bility to reach a set G of goal states within a given time bound z under all late schedulers. More precisely, we prove that for time-bounded reachability, it suffices to consider *total time positional deterministic late schedulers* (TTPDL) which base their decision only on the elapsed time and on the current state. Exploiting this result, we characterize the maximum time-bounded reachability probability as the least fixed point of a higher-order operator which involves integration over the time domain. This allows us to reduce the time-bounded reachability problem for locally uniform CTMDPs to the problem of computing step-bounded reachability probabilities in discrete-time MDPs. More precisely, we approximate the behavior of the CTMDP up to an a priori specified error bound $\varepsilon > 0$ by defining its discretized MDP such that its maximum step-bounded reachability probability coincides (up to ε) with the maximum time-bounded reachability probability of the underlying CTMDP.

In this way, we derive a quantifiably correct approximation method that solves the time-bounded reachability problem for locally uniform CTMDPs by reducing it to the step-bounded reachability problem in MDPs. The latter is a well studied problem [Put94] and can be solved efficiently by linear programming techniques, policy iteration [How60] or value iteration algorithms [Bel57, Ber95]. Hence, our approach is also efficient from a complexity theory point of view. More precisely, we rely on the value iteration algorithm and prove that the worst-case time complexity of our approach is in $\mathcal{O}(m \cdot (\lambda z)^2 / \varepsilon)$, where m denotes the number of transitions in the locally uniform CTMDP and λ is its maximal exit rate.

Although we present all results only for maximum time-bounded reachability probabilities, all proofs can easily be adapted to the dual problem of determining the minimum time-bounded reachability probability.

Organization of this chapter. Section 5.1 introduces the probability measures for locally uniform CTMDPs and late schedulers in full detail. In Sec. 5.2, we develop a fixed-point characterization for the maximal time-bounded reachability probability in locally uniform CTMDP. Moreover, we prove that total time positional schedulers suffice to maximize time-bounded reachability objectives. Section 5.3 defines the discretization, which reduces the time-bounded reachability problem in locally uniform CTMDPs to a step-bounded reachability computation in an MDP. The case study in Sec. 5.4 shows the applicability of our approach by analyzing the best- and worst-case finishing probabilities in the famous stochastic job scheduling problem. Finally, Sec. 5.5 concludes the chapter.

5.1 Locally uniform CTMDPs

As a preparation for the development of our approximation, let us recall the definition of locally uniform CTMDPs and introduce their probabilistic semantics in detail. As we have seen already in Sec. 4.4, the motivation for considering locally uniform CTMDPs

is the fact that they allow us to define a special class of late schedulers which generally induce strictly better probability bounds than the standard *GM*-schedulers.

More precisely, in the standard definition of CTMDPs (cf. Def. 3.11), the exit rate of a state depends on the action that is chosen in that state. This is not the case in locally uniform CTMDPs: Here, we require that the exit rate (and hence the sojourn time distribution) in a state is the same for all enabled actions in that state. Accordingly, we consider the subclass of locally uniform CTMDPs. It is characterized by Def. 4.2 (see page 91) which is equivalent to stating that a CTMDP $\mathcal{C} = (S, Act, \mathbf{R}, \nu)$ is *locally uniform* iff $\forall s \in S. \forall \alpha, \beta \in Act(s). E(s, \alpha) = E(s, \beta)$.

Hence local uniformity ensures that the sojourn time in any state does not depend on the action that is chosen in that state. Hence, we may use $E(s) = E(s, \alpha)$ for some $\alpha \in Act(s)$ to denote the exit rate of state s . In the remainder of this chapter, we assume that all CTMDPs are locally uniform and only mention this restriction where necessary.

Example 5.1. Consider the CTMDP $\mathcal{C}$ in Fig. 5.1. It is locally uniform as in state s_0 , the exit rate under action α is $E(s_0, \alpha) = \sum_{s' \in S} \mathbf{R}(s_0, \alpha, s') = \mathbf{R}(s_0, \alpha, s_2) + \mathbf{R}(s_0, \alpha, s_3) = 1 + 2 = 3$ which equals the exit rate $E(s_0, \beta) = \mathbf{R}(s_0, \beta, s_1) = 3$ of state s_0 under action β . Apart from the fact that it is locally uniform, the behavior of the CTMDP $\mathcal{C}$ is as usual: The choice between actions α and β in state s_0 is nondeterministic. If α is chosen, the α -transitions to states s_2 and s_3 compete for execution. The motivation for local uniformity is the fact, that the sojourn time in s_0 becomes independent of the action that is chosen. In any case, it is exponentially distributed with rate $E(s_0) = 3$. $\diamond$

5.1.1 Probability measures in locally uniform CTMDPs

As we already observed in Sec. 4.4, locally uniform CTMDPs allow us to define *ML*-schedulers that cannot be defined for general CTMDPs and which perform strictly better than general *GM*-schedulers. In Sec. 5.1.2 we will come back to this issue and define the semantics of late schedulers in locally uniform CTMDPs in more detail.

A further remark is necessary before we do so: Obviously, locally uniform CTMDPs are a strict subclass of ordinary CTMDPs: Hence, the construction of their associated measurable spaces remains unaltered and all definitions from Sec. 3.3.2 (see page 76) carry over to the current setting. The probability measures defined on those measurable spaces change however when considering late schedulers:

5.1.2 Measurable late schedulers

The restriction to locally uniform CTMDPs allows us to define a new class of schedulers which we refer to as “late” schedulers. In the classical setting (cf. Sec. 3.3.2), the scheduler immediately decides for an action when entering a state. Intuitively, this is a necessity as the state’s sojourn time distribution is determined by the action that is chosen by the scheduler.

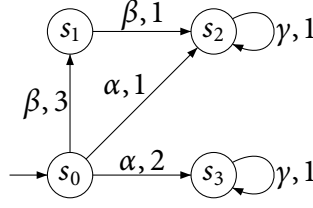


Figure 5.1: Example of a locally uniform CTMDP.

In locally uniform CTMDPs, the setting is different as the state's sojourn time distribution is independent of the selected action. Intuitively, no matter which action the scheduler chooses, the sojourn in the current state remains unaffected. Therefore it is natural to consider schedulers that delay their decision up to the point when the sojourn time has elapsed and the nondeterminism must be resolved in order to obtain the successor-state distribution. This argument leads to the definition of *ML-schedulers*, which postpone their decision up to the point when the current state is left. Thereby, they are able to additionally incorporate the current state's sojourn time into their decision. This is why they expect the sojourn time in the current state as an additional argument:

Definition 5.1 (Measurable late scheduler). A late scheduler for a CTMDP $(\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ is a mapping $D : \text{Paths}^* \times \mathbb{R}_{\geq 0} \times \mathfrak{F}_{\text{Act}} \rightarrow [0, 1]$ where $D(\pi, t, \cdot) \in \text{Distr}(\text{Act}(\pi \downarrow))$ for all $t \in \mathbb{R}_{\geq 0}$ and $\pi \in \text{Paths}^*$. A late scheduler D is a measurable late scheduler (*ML-scheduler*) iff the functions $D(\cdot, \cdot, A) : \text{Paths}^* \times \mathbb{R}_{\geq 0} \rightarrow [0, 1]$ are measurable for all $A \in \mathfrak{F}_{\text{Act}}$.

Similar to the definition of *GM-schedulers* (see Def. 3.13), the measurability condition for *ML-schedulers* states that for all $A \in \mathfrak{F}_{\text{Act}}$ and $B \in \mathfrak{B}([0, 1])$ it must hold that $\{(\pi, t) \mid D(\pi, t, A) \in B\} \in \sigma(\mathfrak{F}_{\text{Paths}^*} \times \mathfrak{B}(\mathbb{R}_{\geq 0}))$.

Intuitively, the behavior of an *ML-scheduler* is described as follows: Let π be a finite path ending in state s with $|\text{Act}(s)| \geq 1$. If state s is left after t units of time, then $D(\pi, t, \cdot)$ is the probability distribution over $\text{Act}(s)$ which resolves the nondeterminism in state s for history π and sojourn time t . For an *ML-scheduler* D , the argument t only refers to the time spent in the current state s . However, D can infer the total time t_π that has passed before taking the decision $D(\pi, t)$ from the sojourn time t and the timing information contained in the trajectory π : Formally, we therefore set $t_\pi = \Delta(\pi) + t$.

Let $ML(\mathcal{C})$ denote the class of *ML-schedulers* for a locally uniform CTMDP $\mathcal{C}$; we omit the reference to $\mathcal{C}$ whenever it is clear from the context. Further, a scheduler $D \in ML$ is *deterministic* if for all $\pi \in \text{Paths}^*$ and $t \in \mathbb{R}_{\geq 0}$, the distribution $D(\pi, t, \cdot)$ is degenerate; otherwise, it is *randomized*. Where appropriate, we use $D(\pi, t)$ to denote the distribution $D(\pi, t, \cdot)$. If $D \in ML$ is deterministic and $D(\pi, t) = \{\alpha \mapsto 1\}$, we identify the distribution $\{\alpha \mapsto 1\}$ and action α .

In the following, we focus on *total time positional late schedulers* [Mil68a, NSK09]

which decide only based on the current state and the total elapsed time t_π , that is, they consider the sum of the time that has elapsed during the trajectory π and the sojourn time in the current state:

Definition 5.2 (Total-time positional late scheduler). Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ be a CT-MDP and $D \in \text{ML}$. The scheduler D is a total-time positional randomized late scheduler (TTPRL) iff for all $\pi_1, \pi_2 \in \text{Paths}^*$ and for all $t_1, t_2 \in \mathbb{R}_{\geq 0}$ it holds that

$$(\pi_1 \downarrow = \pi_2 \downarrow \wedge \Delta(\pi_1) + t_1 = \Delta(\pi_2) + t_2) \Rightarrow D(\pi_1, t_1) = D(\pi_2, t_2).$$

A TTPRL-scheduler yields the same distribution for trajectories π_1 and π_2 if π_1 and π_2 end in the same state (the current state) and if the sums of the time that has passed on path π_1 (resp. path π_2) and the sojourn time t_1 (resp. t_2) of the current state are equal. Therefore, any TTPRL-scheduler D' is isomorphic to a mapping $D : \mathcal{S} \times \mathbb{R}_{\geq 0} \rightarrow \text{Distr}(\text{Act})$, where $D(s, t_\pi) = D'(\pi, t)$ for all paths $\pi \in \text{Paths}^*$ and $t \in \mathbb{R}_{\geq 0}$ with $\Delta(\pi) + t = t_\pi$ and $\pi \downarrow = s$. For the other direction, any measurable mapping $D : \mathcal{S} \times \mathbb{R}_{\geq 0} \rightarrow \text{Distr}(\text{Act})$ induces the TTPRL-scheduler D' with $D'(\pi, t) = D(\pi \downarrow, \Delta(\pi) + t)$. To ease notation and to distinguish between ML and TTPRL-schedulers, in the following we use this one-to-one correspondence and specify TTPRL-schedulers as functions $D : \mathcal{S} \times \mathbb{R}_{\geq 0} \rightarrow \text{Distr}(\text{Act})$. As before, if $D(\pi, t)$ is degenerate for all $\pi \in \text{Paths}^*$ and $t \in \mathbb{R}_{\geq 0}$, the scheduler D is deterministic; accordingly, we use TTPDL to denote the subclass of deterministic TTPRL-schedulers.

5.1.3 Probability measures

Given a CTMDP $\mathcal{C}$, each ML-scheduler D induces a unique stochastic process on $\mathcal{C}$. However, due to the different scheduling discipline of ML-schedulers (compared to GM-schedulers) we have to adapt the definition of the induced probability measures. Therefore, we follow the lines of Sec. 3.3.2 and make adjustments where necessary. As it turns out, we only have to adapt the probability measure $\mu_D(\pi, \cdot)$ for sets of measurable combined transitions (cf. Def. 3.14); all further definitions carry over without modifications.

Recall that paths in a CTMDP can be seen as a finite (or infinite) concatenation of combined transitions; we stick to the notations of Sec. 3.3.2 and use $\Omega = \text{Act} \times \mathbb{R}_{\geq 0} \times \mathcal{S}$ and $\mathfrak{F} = \sigma(\mathfrak{F}_{\text{Act}} \otimes \mathfrak{B}(\mathbb{R}_{\geq 0}) \otimes \mathfrak{F}_{\mathcal{S}})$ to denote the set of combined transitions and their associated σ -field.

Definition 5.3 (Probability of combined transitions). Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ be a CT-MDP and $D \in \text{ML}$. For all $\pi \in \text{Paths}^*$, define the probability measure $\mu_D(\pi, \cdot) : \mathfrak{F} \rightarrow [0, 1]$ where

$$\mu_D(\pi, M) = \int_{\mathbb{R}_{\geq 0}} \eta_{E(\pi \downarrow)}(dt) \int_{\text{Act}} D(\pi, t, d\alpha) \int_{\mathcal{S}} \mathbf{I}_M(\alpha, t, s') \mathbf{P}(s, \alpha, ds').$$

Recall that $\eta_{E(\pi\downarrow)}$ is the exponential distribution of the sojourn time t of the state $\pi\downarrow$ which has rate $E(\pi\downarrow)$; further, $\mathbf{I}_M$ is the characteristic function of $M \in \mathfrak{F}$. In fact, as in Sec. 3.3.2, $\mu_D(\pi, M)$ is the probability to continue with some combined transition in M , given that we hit the current state $\pi\downarrow$ along the trajectory π . However, for late schedulers $D \in ML$, μ_D refers to a slightly different probability measure where the scheduler knows the amount of time that has passed in the current state.

Having the probability measures $\mu_D(\pi, \cdot)$ at hand, we now can define the probabilities of measurable sets of paths in exactly the same way as for early schedulers. We restate the definition here for completeness:

Definition 5.4 (Probability measure). Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ be a CTMDP and $D \in ML$. For $n \geq 0$, we define the probability measures $Pr_{\nu, D}^n$ on the measurable space $(\text{Paths}^n, \mathfrak{F}_{\text{Paths}^n})$ inductively:

$$\begin{aligned} Pr_{\nu, D}^0 : \mathfrak{F}_{\text{Paths}^0} &\rightarrow [0, 1] : \Pi \mapsto \sum_{s \in \Pi} \nu(s) \quad \text{and for } n > 0: \\ Pr_{\nu, D}^n : \mathfrak{F}_{\text{Paths}^n} &\rightarrow [0, 1] : \Pi \mapsto \int_{\text{Paths}^{n-1}} Pr_{\nu, D}^{n-1}(d\pi) \int_{\Omega} \mathbf{I}_{\Pi}(\pi \circ m) \mu_D(\pi, dm). \end{aligned}$$

All other results, especially the extension to measurable cylinders and to the σ -field over infinite paths carry over from Def. 3.15 on page 80.

5.2 A fixed point characterization for time-bounded reachability

In this section, we aim at computing the upper bounds on the probability to reach a set $G \subseteq \mathcal{S}$ of goal states within a given time bound z (denoted $\diamond^{[0, z]} G$) with respect to the class of ML -schedulers.

Definition 5.5 (Maximum time-bounded reachability). Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ be a CTMDP, $G \subseteq \mathcal{S}$, $s \in \mathcal{S}$ and $z \in \mathbb{R}_{\geq 0}$. Then

$$p_{\max}^{\mathcal{C}, G} : \mathcal{S} \times \mathbb{R}_{\geq 0} \rightarrow [0, 1] : (s, z) \mapsto \sup_{D \in ML} Pr_{\nu_s, D}^{\omega}(\diamond^{[0, z]} G)$$

is the maximum time-bounded reachability for the set G of goal states and time bound z .

We omit the superscripts $\mathcal{C}$ and G of $p_{\max}^{\mathcal{C}, G}$ if they are clear from the context. Any scheduler $D \in ML$ induces the reachability probability function $Pr_{\nu_s, D}^{\omega}(\diamond^{[0, \cdot]} G) : \mathbb{R}_{\geq 0} \rightarrow [0, 1]$,

which is continuous by definition. As the following lemma proves, continuity — and thereby measurability with respect to $\mathfrak{B}(\mathbb{R}_{\geq 0})$ — extends to the function $p_{\max}(s, \cdot)$:

Lemma 5.1. *The functions $p_{\max}(s, \cdot) : \mathbb{R}_{\geq 0} \rightarrow [0, 1]$ are continuous and measurable.*

Proof. We have to prove that for all $s \in \mathcal{S}$ and $z \in \mathbb{R}_{\geq 0}$,

$$\lim_{\delta \rightarrow 0^+} p_{\max}(s, z - \delta) = p_{\max}(s, z) = \lim_{\delta \rightarrow 0^+} p_{\max}(s, z + \delta). \quad (5.1)$$

By definition, the reachability probability functions $Pr_{v_s, D}^\omega(\diamond^{[0, \cdot]} G)$ are continuous and monotone; thus, their point-wise supremum $p_{\max}(s, \cdot)$ is also monotone. However, the proof that $p_{\max}(s, \cdot)$ is continuous is not that easy. To see why, note that in general, the pointwise supremum of a countable family of continuous functions is not guaranteed to be continuous. Hence, a more detailed argument is necessary:

To prove that $p_{\max}(s, \cdot)$ is continuous, we proceed by contraposition and assume that there exists $z \in \mathbb{R}_{\geq 0}$ such that (5.1) is violated: Assume that $p_{\max}(s, \cdot)$ is not continuous from the left at point $z \in \mathbb{R}_{\geq 0}$, i.e.

$$\exists \varepsilon > 0. \lim_{\delta \rightarrow 0^+} p_{\max}(s, z - \delta) = p_{\max}(s, z) - \varepsilon. \quad (5.2)$$

Then choose $D \in ML$ such that $Pr_{v_s, D}^\omega(\diamond^{[0, z]} G) = p_{\max}(s, z) - \xi$ for some $\xi \leq \frac{\varepsilon}{2}$. By definition, the function $Pr_{v_s, D}^\omega(\diamond^{[0, \cdot]} G) : \mathbb{R}_{\geq 0} \rightarrow [0, 1]$ is continuous. Further, $Pr_{v_s, D}^\omega(s, z') \leq p_{\max}(s, z')$ for all $z' \in \mathbb{R}_{\geq 0}$ by definition of $p_{\max}$. Therefore, $\lim_{\delta \rightarrow 0^+} Pr_{v_s, D}^\omega(\diamond^{[0, z - \delta]} G) \leq \lim_{\delta \rightarrow 0^+} p_{\max}(s, z - \delta)$. Hence

$$\begin{aligned} p_{\max}(s, z) - \xi &= Pr_{v_s, D}^\omega(\diamond^{[0, z]} G) \\ &= \lim_{\delta \rightarrow 0^+} Pr_{v_s, D}^\omega(\diamond^{[0, z - \delta]} G) \\ &\leq \lim_{\delta \rightarrow 0^+} p_{\max}(s, z - \delta). \end{aligned}$$

But then, $\lim_{\delta \rightarrow 0^+} p_{\max}(s, z - \delta) \geq p_{\max}(s, z) - \xi > p_{\max}(s, z) - \varepsilon$, contradicting (5.2).

Similarly, we prove by contradiction that $p_{\max}(s, \cdot)$ is right-continuous: Assume that $p_{\max}(s, \cdot)$ is not right-continuous, that is, there exists $z \in \mathbb{R}_{\geq 0}$ such that

$$\exists \varepsilon > 0. \lim_{\delta \rightarrow 0^+} p_{\max}(s, z + \delta) = p_{\max}(s, z) + \varepsilon. \quad (5.3)$$

This implies that there exists a scheduler $D \in ML$ that satisfies $\lim_{\delta \rightarrow 0^+} Pr_{v_s, D}^\omega(\diamond^{[0, z + \delta]} G) = \lim_{\delta \rightarrow 0^+} p_{\max}(s, z + \delta) - \xi$ for some $\xi \leq \frac{\varepsilon}{2}$. As before, the function $Pr_{v_s, D}^\omega(\diamond^{[0, \cdot]} G) : \mathbb{R}_{\geq 0} \rightarrow [0, 1]$ is continuous. Further, $Pr_{v_s, D}^\omega(s, z') \leq p_{\max}(s, z')$ for all $z' \in \mathbb{R}_{\geq 0}$ by definition of

p_{max} . Therefore, $Pr_{v_s, D}^\omega(\Diamond^{[0, z]} G) = \lim_{\delta \rightarrow 0^+} Pr_{v_s, D}^\omega(\Diamond^{[0, z+\delta]} G) = \lim_{\delta \rightarrow 0^+} p_{max}(s, z + \delta) - \xi$. Hence

$$\begin{aligned} \lim_{\delta \rightarrow 0^+} p_{max}(s, z + \delta) - \xi &= \lim_{\delta \rightarrow 0^+} Pr_{v_s, D}^\omega(\Diamond^{[0, z+\delta]} G) \\ &= Pr_{v_s, D}^\omega(\Diamond^{[0, z]} G) \\ &\leq p_{max}(s, z). \end{aligned}$$

But then, $\lim_{\delta \rightarrow 0^+} p_{max}(s, z + \delta) \leq p_{max}(s, z) + \xi < p_{max}(s, z) + \varepsilon$, contradicting (5.3).

Thus, $p_{max}(s, \cdot)$ is continuous. As continuity implies measurability [ADD00, p.36], the claim follows. $\square$

The next theorem shows that the function p_{max} is the least fixed point of a higher order operator Ω which is defined on measurable functions $F : \mathcal{S} \times \mathbb{R}_{\geq 0} \rightarrow [0, 1]$. This result is essential for the discretization developed in Sec. 5.3.1. It has been inspired by a similar fixed point characterization which is used in [BHHK03, Thm. 1] to derive the probability of time-bounded until formulas in CTMCs.

Theorem 5.1 (A fixed point characterization for time-bounded reachability). *Let $\mathcal{C} = (\mathcal{S}, Act, \mathbf{R}, v)$ be a CTMDP and $G \subseteq \mathcal{S}$ a set of goal states. Then p_{max} is the least fixed point of the higher-order operator $\Omega : (\mathcal{S} \times \mathbb{R}_{\geq 0} \rightarrow [0, 1]) \rightarrow (\mathcal{S} \times \mathbb{R}_{\geq 0} \rightarrow [0, 1])$ which is defined for $s \in \mathcal{S}$, $z \in \mathbb{R}_{\geq 0}$, and measurable function $F : \mathcal{S} \times \mathbb{R}_{\geq 0} \rightarrow [0, 1]$ such that $\Omega(F)(s, z) = 1$ if $s \in G$ and for $s \notin G$:*

$$\Omega(F)(s, z) = \int_0^z E(s) e^{-E(s)t} \cdot \max_{\alpha \in Act} \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot F(s', z - t) dt. \quad (5.4)$$

Proof. The proof is split in two parts: First, we show that p_{max} is a fixed point of Ω . Second, we prove that p_{max} is the least fixed point of Ω by decomposing the event $\Diamond^{[0, z]} G$ with respect to the number n of transitions that are needed to reach a state in G . By induction on n , we then prove that $p_{max}(s, z) \leq F(s, z)$ for any other fixed point F of Ω and all $s \in \mathcal{S}$ and $z \in \mathbb{R}_{\geq 0}$.

We prove that p_{max} is a fixed point of Ω as follows: If $s \in G$, then $p_{max}(s, z) = 1 = \Omega(p_{max})(s, z)$ and the claim follows. If $s \notin G$, we proceed as follows. Let $\Pi(z, n)$ be the set of all infinite paths $\pi = s_0 \xrightarrow{\alpha_0, t_0} s_1 \xrightarrow{\alpha_1, t_1} \dots$ such that $s_n \in G$ and $s_i \notin G$ for all $i < n$ and $\sum_{i=0}^{n-1} t_i \leq z$. Further, let $p_{max}^n(s, z) = \sup_{D \in ML} Pr_{v_s, D}^\omega(\bigcup_{i=0}^n \Pi(z, i))$ be the least upper bound on the probability to reach G within z time units with at most n transitions.

In a first step, we prove that $p_{max}^{n+1}(s, z) = \Omega(p_{max}^n)(s, z)$. By definition we have:

$$\Omega(p_{max}^n)(s, z) = \int_0^z E(s) e^{-E(s)t} \cdot \max_{\alpha \in Act} \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot p_{max}^n(s', z - t) dt$$

$$= \int_0^z E(s) e^{-E(s)t} \cdot \underbrace{\max_{\alpha \in Act} \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot \sup_{D' \in ML} Pr_{v_{s'}, D'}^\omega \left(\bigcup_{i=0}^n \Pi(z-t, i) \right)}_{c(\alpha)} dt. \quad (5.5)$$

For given state $s \in \mathcal{S}$, a given number of transitions $n \in \mathbb{N}$, a given time bound z and a fixed sojourn time t , we define the function $c : Act \rightarrow [0, 1]$ such that $c(\alpha) = \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot \sup_{D' \in ML} Pr_{v_{s'}, D'}^\omega \left(\bigcup_{i=0}^n \Pi(z-t, i) \right)$. Further, let $\gamma \in Act$ denote a maximal action for state s and time t , i.e. $c(\gamma) = \max_{\alpha \in Act} c(\alpha)$. Obviously, any convex combination of actions does not yield values larger than $c(\gamma)$: More precisely, it holds $c(\gamma) = \sup_{\mu \in \text{Distr}(Act)} \sum_{\alpha \in Act} \mu(\alpha) \cdot c(\alpha)$.

Now, let $D \in ML$, $s \in \mathcal{S}$, $\alpha \in Act$ and $t \in \mathbb{R}_{\geq 0}$. We define the ML scheduler $D_{s, \alpha, t}$ such that $D_{s, \alpha, t}(\pi, t')(\beta) = D(s \xrightarrow{\alpha, t} \pi, t')(\beta)$ for all $\pi \in \text{Paths}^*$, $\beta \in Act$ and $t' \in \mathbb{R}_{\geq 0}$. Hence, $D_{s, \alpha, t}$ yields the same decisions for history π as the original scheduler D does for the history $s \xrightarrow{\alpha, t} \pi$, where we define $s \xrightarrow{\alpha, t} \pi = s \xrightarrow{\alpha, t} s_0 \xrightarrow{\alpha_0, t_0} s_1 \xrightarrow{\alpha_1, t_1} \dots$ if $\pi = s_0 \xrightarrow{\alpha_0, t_0} s_1 \xrightarrow{\alpha_1, t_1} \dots$. Thus, we can rewrite (5.5):

$$\begin{aligned} \Omega(p_{max}^n)(s, z) &= \int_0^z E(s) e^{-E(s)t} \cdot \underbrace{\sup_{\mu \in \text{Distr}(Act)} \sum_{\alpha \in Act} \mu(\alpha)}_{c(\alpha)} \\ &\quad \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot \sup_{D' \in ML} Pr_{v_{s'}, D'}^\omega \left(\bigcup_{i=0}^n \Pi(z-t, i) \right) dt \\ &= \int_0^z E(s) e^{-E(s)t} \cdot \sup_{D \in ML} \sum_{\alpha \in Act} D(s, t)(\alpha) \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \\ &\quad \cdot Pr_{v_{s'}, D_{s, \alpha, t}}^\omega \left(\bigcup_{i=0}^n \Pi(z-t, i) \right) dt \\ &\stackrel{(*)}{=} \sup_{D \in ML} \int_0^z E(s) e^{-E(s)t} \cdot \sum_{\alpha \in Act} D(s, t)(\alpha) \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \\ &\quad \cdot Pr_{v_{s'}, D_{s, \alpha, t}}^\omega \left(\bigcup_{i=0}^n \Pi(z-t, i) \right) dt \\ &= \sup_{D \in ML} Pr_{v_s, D}^\omega \left(\bigcup_{i=0}^{n+1} \Pi(z, i) \right) = p_{max}^{n+1}(s, z). \end{aligned}$$

Note that in the above derivation, we swap the supremum $\sup_{D \in ML}$ and the integral to obtain equality (*). In this case this can be done, as each late scheduler $D \in ML$ is a function which expects the integration variable t as an argument: To see this, fix some $t \in [0, z]$ and let $D^{t,1}, D^{t,2}, \dots$ be a sequence of schedulers that converges to the supremum, that is

$$\sup_{D \in ML} \sum_{\alpha \in Act} D(s, t)(\alpha) \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot Pr_{v_{s'}, D_{s, \alpha, t}}^\omega \left(\bigcup_{i=0}^n \Pi(z-t, i) \right) dt$$

$$= \lim_{i \rightarrow \infty} \sum_{\alpha \in Act} D^{t,i}(s, t)(\alpha) \cdot \sum_{s' \in S} \mathbf{P}(s, \alpha, s') \cdot Pr_{v_{s'}, D^{t,i}}^{\omega} \left(\bigcup_{i=0}^n \Pi(z - t, i) \right) dt.$$

If we define a sequence of *ML*-schedulers $\hat{D}^i$ such that $\hat{D}^i(s, t) = D^{t,i}(s, t)$ for all $t \in [0, z]$, then the probabilities induced by the sequence $\hat{D}^i$ converge pointwise to the supremum by construction. Hence, equality (*) follows.

Thus $p_{max}^{n+1}(s, z) = \Omega(p_{max}^n)(s, z)$; further, Prop. 5.1 (Prop. 5.1 is given below on page 123) states that $\lim_{n \rightarrow \infty} p_{max}^n(s, z) = p_{max}(s, z)$ for all $s \in S$ and $z \in \mathbb{R}_{\geq 0}$. Therefore

$$\begin{aligned} p_{max}(s, z) &= \lim_{n \rightarrow \infty} p_{max}^n(s, z) = \lim_{n \rightarrow \infty} p_{max}^{n+1}(s, z) \\ &= \lim_{n \rightarrow \infty} \Omega(p_{max}^n)(s, z) = \Omega\left(\lim_{n \rightarrow \infty} p_{max}^n\right)(s, z) = \Omega(p_{max})(s, z), \end{aligned}$$

proving that p_{max} is a fixed point of Ω . In the above derivation step, note that by definition of Ω one can show that $\lim_{n \rightarrow \infty} \Omega(p_{max}^n)(s, z) = \Omega(\lim_{n \rightarrow \infty} p_{max}^n)(s, z)$:

$$\begin{aligned} \lim_{n \rightarrow \infty} \Omega(p_{max}^n)(s, z) &= \lim_{n \rightarrow \infty} \int_0^z E(s) e^{-E(s)t} \cdot \max_{\alpha \in Act} \sum_{s' \in S} \mathbf{P}(s, \alpha, s') \cdot p_{max}^n(s', z - t) dt \\ &= \int_0^z E(s) e^{-E(s)t} \cdot \max_{\alpha \in Act} \sum_{s' \in S} \mathbf{P}(s, \alpha, s') \cdot \lim_{n \rightarrow \infty} p_{max}^n(s', z - t) dt \\ &= \Omega(\lim_{n \rightarrow \infty} p_{max}^n)(s, z). \end{aligned}$$

It remains to show that p_{max} is the least fixed point of Ω . From the first part, we know that p_{max} is a fixed point of Ω and that $p_{max}^{n+1}(s, z) = \Omega(p_{max}^n)(s, z)$. Now, let $F : S \times \mathbb{R}_{\geq 0} \rightarrow [0, 1]$ be another fixed point of Ω . By induction on n , we show that $p_{max}^n(s, z) \leq F(s, z)$ for all $n \in \mathbb{N}$. For the base case, $p_{max}^0(s, z) = 1 = \Omega(F(s, z)) = F(s, z)$ if $s \in G$ and $p_{max}^0(s, z) = 0 \leq F(s, z)$, otherwise. Further,

$$\begin{aligned} p_{max}^{n+1}(s, z) &= \Omega(p_{max}^n)(s, z) \\ &= \int_0^z E(s) e^{-E(s)t} \cdot \max_{\alpha \in Act} \sum_{s' \in S} \mathbf{P}(s, \alpha, s') \cdot p_{max}^n(s', z - t) dt \\ &\quad (\text{* by the induction hypothesis *}) \\ &\leq \int_0^z E(s) e^{-E(s)t} \cdot \max_{\alpha \in Act} \sum_{s' \in S} \mathbf{P}(s, \alpha, s') \cdot F(s', z - t) dt \\ &= \Omega(F(s, z)) = F(s, z). \end{aligned}$$

Hence, $F(s, z) \geq \lim_{n \rightarrow \infty} p_{max}^n(s, z) = p_{max}(s, z)$ and the claim follows. $\square$

In the proof of Thm. 5.1 we need to exchange the order of taking the limit and the supremum. This is justified by the following proposition:

Proposition 5.1. *Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ be a CTMDP, $s \in \mathcal{S}$, $G \subseteq \mathcal{S}$ and $z \in \mathbb{R}_{\geq 0}$. Further, let $\Pi(z, i)$ be defined as in the proof of Thm. 5.1. Then*

$$\lim_{n \rightarrow \infty} \sup_{D \in ML} Pr_{\nu_s, D}^\omega \left(\bigcup_{i=0}^n \Pi(z, i) \right) = \sup_{D \in ML} Pr_{\nu_s, D}^\omega \left(\diamond^{[0, z]} G \right).$$

Proof. Recall that $\Pi(z, i) = \{ \pi \in \text{Paths}^\omega \mid \pi[i] \in G \wedge \forall k < i. \pi[k] \notin G \wedge \sum_{k=0}^{i-1} \delta(\pi, k) \leq z \}$. Let $\Pi_n := \bigcup_{i=0}^n \Pi(z, i)$; then $\Pi_n \subseteq \Pi_{n+1}$ and $\Pi_n \uparrow \diamond^{[0, z]} G$. By [ADD00, Thm. 1.2.7(a)], it holds for all $D \in ML$ that $Pr_{\nu_s, D}^\omega(\Pi_n) \rightarrow Pr_{\nu_s, D}^\omega(\diamond^{[0, z]} G)$ for $n \rightarrow \infty$. As this reasoning applies to all $D \in ML$, it holds that $\sup \{ Pr_{\nu_s, D}^\omega(\Pi_n) \mid D \in ML \} \rightarrow \sup \{ Pr_{\nu_s, D}^\omega(\diamond^{[0, z]} G) \mid D \in ML \}$ for $n \rightarrow \infty$. Therefore we can conclude that $\lim_{n \rightarrow \infty} \sup \{ Pr_{\nu_s, D}^\omega(\Pi_n) \mid D \in ML \} = \sup_{D \in ML} Pr_{\nu_s, D}^\omega(\diamond^{[0, z]} G)$. $\square$

Let us come back to Thm. 5.1. Intuitively, the term $E(s)e^{-E(s)t}$ on the right-hand side of Eq. 5.4 corresponds to the density of the sojourn time in state s ; accordingly, if state s is left at time t , we multiply with the maximum probability (with respect to all actions $\alpha \in \text{Act}$) to reach a goal state in G via action α within the remaining $z - t$ time units.

5.2.1 Optimal TTPDL schedulers

Given the fixed point characterization of Thm. 5.1, we now define a TTPDL scheduler which induces the probabilities $p_{\max}$. Note that the fact that this is possible has an important implication: Obviously, the additional information available to ML-schedulers is irrelevant for achieving maximum time-bounded reachability probabilities!

A scheduler $D \in ML$ is *optimal* for the set of goal states G and time bound z iff for all $D' \in ML$ and $s \in \mathcal{S}$ it holds that $Pr_{\nu_s, D'}^\omega(\diamond^{[0, z]} G) \leq Pr_{\nu_s, D}^\omega(\diamond^{[0, z]} G)$. Further, for $\varepsilon > 0$, $D \in ML$ is ε -*optimal* for G and z iff $|Pr_{\nu_s, D}^\omega(\diamond^{[0, z]} G) - p_{\max}(s, z)| \leq \varepsilon$ for all $s \in \mathcal{S}$. Note that up to now, it is not clear whether an optimal scheduler exists. We answer this question in the affirmative by first defining a TTPDL scheduler D^z and then proving that D^z is optimal (cf. Thm. 5.2):

Definition 5.6 (The scheduler D^z). *Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ be a CTMDP, $G \subseteq \mathcal{S}$ a set of goal states and $z \in \mathbb{R}_{\geq 0}$ a time bound. Given an arbitrary (fixed) total order $<$ on Act , we define the TTPDL scheduler D^z such that for all $s \in \mathcal{S}$ and $t_\pi \leq z$:*

$$D^z(s, t_\pi) = \min_{<} \{ \alpha \in \text{Act}(s) \mid \forall \beta \in \text{Act}(s). f(s, z - t_\pi, \beta) \leq f(s, z - t_\pi, \alpha) \},$$

where $f(s, z', \gamma) = \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \gamma, s') \cdot p_{\max}(s', z')$. If $t_\pi > z$, set $D^z(s, t_\pi) = \min_{<} \text{Act}(s)$.

Here $f(s, z - t_\pi, \beta)$ denotes the maximum probability to reach a state in G within the remaining $z - t_\pi$ time units via action β for the case that t_π time units have expired on the path that led to state s and in state s itself. However, multiple actions α may exist that maximize $f(s, z - t_\pi, \alpha)$. Hence, we fix some total order $<$ to ensure uniqueness of D^z . Note that Def. 5.6 implies that $D^z(s, t_\pi + t) = D^{z-t_\pi}(s, t)$ for all $s \in \mathcal{S}$, $t, z \in \mathbb{R}_{\geq 0}$ and $t_\pi \leq z$.

Exploiting the measurability of $p_{\max}$ (cf. Lemma 5.1), we show that D^z is measurable:

Lemma 5.2. *The schedulers D^z are measurable for all $z \in \mathbb{R}_{\geq 0}$.*

Proof. Let $z \in \mathbb{R}_{\geq 0}$ be a time bound and let $<$ be the total order on Act as given in Def. 5.6. Then D^z is defined by

$$D^z(s, t_\pi) = \min_{<} \{ \alpha \in Act(s) \mid \forall \beta \in Act(s). f(s, z - t_\pi, \beta) \leq f(s, z - t_\pi, \alpha) \}$$

and depends only on the function

$$f(s, z', \gamma) = \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \gamma, s') \cdot p_{\max}(s', z') = \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \gamma, s') \cdot \sup_{D' \in ML} Pr_{v_{s'}, D'}^\omega(\Diamond^{[0, z']} G).$$

By Lemma 5.1, the function $p_{\max}(s, \cdot)$ is continuous; this implies that $p_{\max}(s, \cdot)$ is measurable with respect to the Lebesgue-measure on $\mathfrak{B}(\mathbb{R}_{\geq 0})$. Hence, the functions $f(s', \cdot, \gamma) : \mathbb{R}_{\geq 0} \rightarrow [0, 1]$ are measurable. Now $D^z(s, t_\pi) = \alpha$ iff $f(s, z - t_\pi, \alpha) = \max_{\beta \in Act} f(s, z - t_\pi, \beta)$ and α is minimal with respect to $<$. Measurability of D^z now follows from the fact, that the maximum of measurable functions is again measurable and that by $<$, the minimal action is uniquely determined. $\square$

With the measurability of D^z , we are now able to prove that the scheduler D^z indeed maximizes the probability of reaching G within at most z time units for any initial state s :

Theorem 5.2 (Optimality). *Let $\mathcal{C} = (\mathcal{S}, Act, \mathbf{R}, v)$ be a CTMDP, $G \subseteq \mathcal{S}$ a set of goal states, $s \in \mathcal{S}$ an initial state and $z \in \mathbb{R}_{\geq 0}$ a time bound. Then*

$$Pr_{v_s, D^z}^\omega(\Diamond^{[0, z]} G) = p_{\max}(s, z).$$

Proof. For the proof, we define total time step counting positional late schedulers which are a superclass of *TTPL* schedulers that also considers the number of transitions taken before reaching the current state: A scheduler $D \in ML$ is a *total time step counting positional late scheduler* (*TTSCPRL*) iff $\forall \pi_1, \pi_2 \in Paths^*. \forall t_1, t_2 \in \mathbb{R}_{\geq 0}. (\pi_1 \downarrow = \pi_2 \downarrow \wedge |\pi_1| = |\pi_2| \wedge \Delta(\pi_1) + t_1 = \Delta(\pi_2) + t_2) \Rightarrow D(\pi_1, t_1) = D(\pi_2, t_2)$. Hence, any *TTSCPRL* scheduler $D \in ML$ can be expressed equivalently as a function $D' : \mathcal{S} \times \mathbb{N} \times \mathbb{R}_{\geq 0} \rightarrow Distr(Act)$, where

$D'(\pi \downarrow, |\pi|, \Delta(\pi) + t) = D(\pi, t)$ for all $\pi \in \text{Paths}^*$ and $t \in \mathbb{R}_{\geq 0}$. Note that *TTSCPRL* schedulers extend *TTPRL* schedulers, as they additionally depend in their second argument on the number of transitions that have occurred up to the current state. A *TTSCPRL* scheduler D is *deterministic* (*TTSCPDL*) iff $\forall s \in \mathcal{S}. \forall c \in \mathbb{N}. \forall t_\pi \in \mathbb{R}_{\geq 0}. \exists \alpha \in \text{Act}. D(s, c, t_\pi) = \{\alpha \mapsto 1\}$. To ease notation, we assume that *TTSCPDL* schedulers are given as mappings of the form $D : \mathcal{S} \times \mathbb{N} \times \mathbb{R}_{\geq 0} \rightarrow \text{Act}$.

For the proof, we define the *TTSCPDL* schedulers $D_n^z : \mathcal{S} \times \mathbb{N} \times \mathbb{R}_{\geq 0} \rightarrow \text{Act}$ with respect to the total order $<$ on *Act* used in Def. 5.6 such that

$$D_n^z(s, c, t_\pi) = \min_{<} \{ \alpha \in \text{Act}(s) \mid \forall \beta \in \text{Act}(s). \\ f'(s, n - c - 1, z - t_\pi, \beta) \leq f'(s, n - c - 1, z - t_\pi, \alpha) \},$$

where $f'(s, n', z', \gamma) = \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \gamma, s') \cdot \sup_{D' \in \text{ML}} Pr_{v_{s'}, D'}^\omega(\bigcup_{i=0}^{n'} \Pi(z', i))$. Hence $D_n^z(s, c, t_\pi)$ is the optimal action if $n - c - 1$ steps and $z - t_\pi$ time units remain to reach a goal state in G .

Now, let $p_{\max}^n(s, z) = \sup_{D \in \text{ML}} Pr_{v_s, D}^\omega(\bigcup_{i=0}^n \Pi(z, i))$ be defined as in the proof of Thm. 5.1. Further, we define $q_{\max}^n(s, z) = Pr_{v_s, D_n^z}^\omega(\bigcup_{i=0}^n \Pi(z, i))$ and $q_{\max}(s, z) = Pr_{v_s, D^z}^\omega(\diamond^{[0, z]} G)$. Thus, we aim at proving that $p_{\max} = q_{\max}$; as a first step, we show by induction on n that $p_{\max}^n = q_{\max}^n$:

1. In the induction base, we distinguish two cases: If $s \in G$, then $p_{\max}^0(s, z) = 1 = q_{\max}^0(s, z)$; otherwise, $p_{\max}^0(s, z) = 0 = q_{\max}^0(s, z)$. Hence, $p_{\max}^0 = q_{\max}^0$.
2. To prove the induction step, we use the fact (cf. the proof of Thm. 5.1) that $p_{\max}^{n+1} = \Omega(p_{\max}^n)$. As induction hypothesis, assume that $p_{\max}^n = q_{\max}^n$. Then

$$\begin{aligned} p_{\max}^{n+1}(s, z) &= \Omega(p_{\max}^n)(s, z) && (* \text{ as shown in the proof of Thm. 5.1 } *) \\ &= \int_0^z E(s) e^{-E(s)t} \cdot \max_{\alpha \in \text{Act}} \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot p_{\max}^n(s', z - t) dt \\ &&& (* \text{ definition of } D_{n+1}^z *) \\ &= \int_0^z E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, D_{n+1}^z(s, 0, t), s') \cdot p_{\max}^n(s', z - t) dt \\ &&& (* \text{ applying the induction hypothesis } *) \\ &= \int_0^z E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, D_{n+1}^z(s, 0, t), s') \cdot q_{\max}^n(s', z - t) dt \\ &&& (* \text{ definition of } q_{\max}^n *) \\ &= \int_0^z E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, D_{n+1}^z(s, 0, t), s') \cdot Pr_{v_{s'}, D_{n+1}^z}^\omega(\bigcup_{i=0}^n \Pi(z, i)) dt \\ &&& (* \text{ see Remark 5.1 below } *) \\ &= \int_0^z E(s) e^{-E(s)t} \sum_{s' \in \mathcal{S}} \mathbf{P}(s, D_{n+1}^z(s, 0, t), s') Pr_{v_{s'}, D_{n+1}^z}^\omega(\bigcup_{i=0}^n \Pi(z, i)) dt \end{aligned}$$

(* by definition of $Pr_{v_s, D_{n+1}^z}^\omega$ *)

$$= Pr_{v_s, D_{n+1}^z}^\omega \left(\bigcup_{i=0}^{n+1} \Pi(z, i) \right) = q_{max}^{n+1}(s, z).$$

Remark 5.1. In the derivations above, we use $D_{n+1}^z(\cdot, \cdot^{+1}, \cdot^{+t})$ to denote the TTSCPDL scheduler that is given by $D_{n+1}^z(\cdot, \cdot^{+1}, \cdot^{+t}) : \mathcal{S} \times \mathbb{N} \times \mathbb{R}_{\geq 0} \rightarrow \text{Act}$ with $D_{n+1}^z(\cdot, \cdot^{+1}, \cdot^{+t})(s, c, t_\pi) = D_{n+1}^z(s, c+1, t+t_\pi)$. Note that from the definition of D_n^z and the function f' , it follows directly that $D_{n+1}^z(s, c+1, t+t_\pi) = D_n^{z-t}(s, c, t_\pi)$ for all $s \in \mathcal{S}$, $c \in \mathbb{N}$, $t \leq z$ and $t_\pi \in \mathbb{R}_{\geq 0}$.

With the above induction, we have shown that $p_{max}^n = q_{max}^n$ for all $n \in \mathbb{N}$. Now it remains to prove that $q_{max}^n \rightarrow q_{max}$ for $n \rightarrow \infty$. Therefore, note that

$$\begin{aligned} \lim_{n \rightarrow \infty} f'(s, n, z', \gamma) &= \lim_{n \rightarrow \infty} \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \gamma, s') \cdot \sup_{D' \in ML} Pr_{v_{s'}, D'}^\omega \left(\bigcup_{i=0}^n \Pi(z', i) \right) \quad (* \text{ def. } f'^*) \\ &= \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \gamma, s') \cdot \lim_{n \rightarrow \infty} \sup_{D' \in ML} Pr_{v_{s'}, D'}^\omega \left(\bigcup_{i=0}^n \Pi(z', i) \right) \\ &= \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \gamma, s') \cdot \sup_{D' \in ML} Pr_{v_{s'}, D'}^\omega \left(\Diamond^{[0, z']} G \right) \quad (* \text{ by Prop. 5.1} *) \\ &= f(s, z', \gamma). \end{aligned}$$

As D^z and D_n^z are defined with respect to functions f and f' , respectively, it follows that for $n \rightarrow \infty$, $D_n^z(s, c, t_\pi) = D^z(s, t_\pi)$ for all $c \in \mathbb{N}$, $s \in \mathcal{S}$ and $t_\pi \in \mathbb{R}_{\geq 0}$. Thus for $n \rightarrow \infty$:

$$q_{max}^n(s, z) = Pr_{v_s, D_n^z}^\omega \left(\bigcup_{i=0}^n \Pi(z, i) \right) \rightarrow Pr_{v_s, D^z}^\omega \left(\Diamond^{[0, z]} G \right) = q_{max}(s, z).$$

Now the claim follows as we have for all $s \in \mathcal{S}$ and $z \in \mathbb{R}_{\geq 0}$:

$$p_{max}(s, z) = \lim_{n \rightarrow \infty} p_{max}^n(s, z) = \lim_{n \rightarrow \infty} q_{max}^n(s, z) = q_{max}(s, z). \quad \square$$

The proof of the theorem is quite technical. Therefore, we give another, slightly more intuitive but formally not completely correct argument and explain why the technical details (such as the introduction of TTSCPDL schedulers) in the formal proof of Thm. 5.2 are indeed necessary:

By Thm. 5.1, it holds for all $s \in \mathcal{S}$ and $z \in \mathbb{R}_{\geq 0}$ that

$$\begin{aligned} p_{max}(s, z) &= \Omega(p_{max})(s, z) \\ &= \int_0^z E(s) e^{-E(s)t} \cdot \max_{\alpha \in \text{Act}} \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot p_{max}(s', z-t) dt \\ &= \int_0^z E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, D^z(s, t), s') \cdot p_{max}(s', z-t) dt. \end{aligned}$$

Applying this equality recursively to the term $p_{max}(s', z-t)$ shows that D^z induces the probability $p_{max}(s, z)$ for the event $\Diamond^{[0, z]} G$ and initial state s . To see this, note that

$D^{z-t_\pi}(s, t) = D^z(s, t_\pi + t)$ for all $t_\pi \leq z$ and $t \in \mathbb{R}_{\geq 0}$; hence, the scheduler D^{z-t} at time t' which is used in the next recursion step (i.e. within $p_{\max}(s', z - t')$) equals D^z at time $t + t'$. Hence the above equation yields a recursive definition of $p_{\max}$ which depends only on D^z .

However, the above reasoning uses an inductive argument on z , although the domain of z (i.e. the positive reals) is not well-founded. Therefore, in the formal proof of Thm. 5.2 we use induction on the number $n \in \mathbb{N}$ of transitions available to reach G within time z and resort to step counting *TTSCPD*L schedulers.

A direct consequence of Thm. 5.2 is the existence of optimal schedulers. Further:

Corollary 5.1. *TTPDL schedulers suffice to maximize time-bounded reachability probabilities.*

5.2.2 Piecewise-constant schedulers

In Def. 5.5, the upper bound $p_{\max}$ on the maximum time-bounded reachability probability of a set G of goal states is defined with respect to the class of *ML*-schedulers. Corollary 5.1 allows us to only consider the subclass of *TTPDL* schedulers to compute $p_{\max}$, i.e. we restrict to schedulers of the form $D : \mathcal{S} \times \mathbb{R}_{\geq 0} \rightarrow \text{Act}$. However, *TTPDL* schedulers are still continuous in their second argument. To obtain schedulers with a finite representation, we now introduce *piecewise-constant TTPDL* schedulers.

They prove to be useful for the scheduler synthesis that we discuss in Sec. 5.3.4. As we will see, a byproduct of our discretization technique is an ε -optimal τ -scheduler which approximates the optimal reachability probability up to an a priori specified error ε and which changes its decisions only in between time-intervals of length τ .

Definition 5.7 (Piecewise-constant TTPDL scheduler). Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ be a CTMDP and $D : \mathcal{S} \times \mathbb{R}_{\geq 0} \rightarrow \text{Act}$ a TTPDL scheduler. D is piecewise-constant iff for all $s \in \mathcal{S}$ and $\alpha \in \text{Act}(s)$ there exist disjoint intervals $A_{s,\alpha}^0, A_{s,\alpha}^1, A_{s,\alpha}^2, \dots \subseteq \mathbb{R}_{\geq 0}$ such that for all $t_\pi \in \mathbb{R}_{\geq 0}$: $D(s, t_\pi) = \alpha \iff t_\pi \in \bigcup_{i=0}^\infty A_{s,\alpha}^i$. A piecewise-constant scheduler D is non-Zeno if $\left| \{A_{s,\alpha}^i \mid \inf A_{s,\alpha}^i < z\} \right| < \infty$ for all $z \in \mathbb{R}_{\geq 0}$, $s \in \mathcal{S}$ and $\alpha \in \text{Act}$.

We use *PCDL* to denote the set of all piecewise-constant and non-Zeno *TTPDL* schedulers. Intuitively, for a state $s \in \mathcal{S}$ and a given time-bound z , a *PCDL*-scheduler changes its decision for an action only finitely many times: The intervals $A_{s,\alpha}^i$ in Def. 5.7 describe the time-periods, in which the scheduler chooses action α constantly if the current state is s . The non-Zeno assumption implies that only finitely many decision epochs occur up to time z .

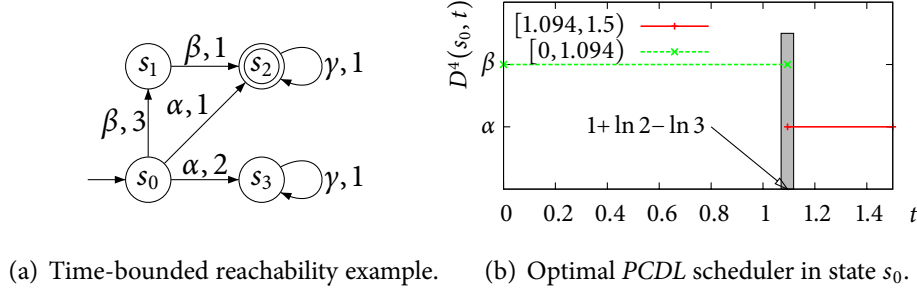


Figure 5.2: Maximizing time-bounded reachability objectives with PCDL schedulers.

Theorem 5.3 (PCDL schedulers maximize time-bounded reachability probabilities).

Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ be a CTMDP, $G \subseteq \mathcal{S}$ a set of goal states, $s \in \mathcal{S}$ an initial state and $z \in \mathbb{R}_{\geq 0}$ a time bound. Then

$$p_{\max}(s, z) = \sup_{D \in \text{PCDL}} Pr_{\nu_s, D}^{\omega} (\diamond^{[0, z]} G).$$

Proof. The proof relies on a measure theoretic argument: As D^z is measurable and deterministic, the sets $A_{s, \alpha} = \{t_\pi \in \mathbb{R}_{\geq 0} \mid D^z(s, t_\pi) = \alpha\}$ are Borel measurable for all $s \in \mathcal{S}$ and $\alpha \in \text{Act}$. The approximation theorem (cf. Thm. 2.4 on page 24) then permits to approximate each set $A_{s, \alpha}$ arbitrarily closely by a finite number of intervals which give rise to a PCDL scheduler.

Therefore, let $s \in \mathcal{S}$, $\alpha \in \text{Act}$ and define $A_{s, \alpha} = D^z(s, \cdot)^{-1}(\alpha)$. By definition, D^z is a measurable scheduler. Hence $A_{s, \alpha} \in \mathfrak{B}$. Now let $\mathfrak{B}_0$ be a field of subsets of $\mathbb{R}_{\geq 0}$ that generates the σ -field $\mathfrak{B}$, i.e. let $\sigma(\mathfrak{B}_0) = \mathfrak{B}$. Given $\varepsilon > 0$, we can apply Thm. 2.4 to approximate the set $A_{s, \alpha}$ by a set $B_{s, \alpha} \in \mathfrak{B}_0$ up to an error of ε . More precisely, let $\theta : \mathfrak{B} \rightarrow \mathbb{R}_{\geq 0}$ be the Lebesgue measure defined by the distribution function $\Theta(y) = y$ for $y \in \mathbb{R}_{\geq 0}$. Thus, we use the Lebesgue measure θ to measure the “length” of measurable subsets of $\mathbb{R}_{\geq 0}$. If $A \triangle B = (A \setminus B) \cup (B \setminus A)$ denotes set difference, Thm. 2.4 assures that $B_{s, \alpha}$ exists such that $\theta(A_{s, \alpha} \triangle B_{s, \alpha}) < \varepsilon$.

For $\mathfrak{B}_0$, we choose the set of finite disjoint unions of right semi-closed intervals; as $\mathfrak{B}_0$ is a field and $\sigma(\mathfrak{B}_0) = \mathfrak{B}$, this is a valid choice (see also Lemma 2.6 and Def. 2.7). As $B_{s, \alpha} \in \mathfrak{B}_0$, there exist $n_{s, \alpha} \in \mathbb{N}$ and disjoint intervals $B_{s, \alpha}^0, \dots, B_{s, \alpha}^{n_{s, \alpha}}$ such that $B_{s, \alpha} = \bigcup_{i=0}^{n_{s, \alpha}} B_{s, \alpha}^i$. Now we are ready to construct a scheduler D_ε^z which approximates D^z up to an error of ε as follows: $D_\varepsilon^z(s, t_\pi) = \alpha \iff t_\pi \in \bigcup_{i=0}^{n_{s, \alpha}} B_{s, \alpha}^i$. By definition, D_ε^z is a piecewise constant and a non-Zeno scheduler. Thus $D_\varepsilon^z \in \text{PCDL}$ for all $\varepsilon > 0$; further, from the fact that $\theta(\{t_\pi \in \mathbb{R}_{\geq 0} \mid D^z(s, t_\pi) \neq D_\varepsilon^z(s, t_\pi)\}) < \varepsilon$, we obtain for the probability measures on combined transitions (cf. Def. 5.3) that $\lim_{\varepsilon \rightarrow 0} \mu_{D_\varepsilon^z}(\pi, \cdot) = \mu_{D^z}(\pi, \cdot)$ for all $\pi \in \text{Paths}^*$. This

extends inductively (cf. Def. 5.4) to the probability measure on infinite paths, i.e.

$$\lim_{\varepsilon \rightarrow 0} Pr_{v_s, D_\varepsilon^\omega}^\omega \left(\Diamond^{[0, z]} G \right) = Pr_{v_s, D^z}^\omega \left(\Diamond^{[0, z]} G \right) = p_{\max}(s, z).$$

Now the claim follows, as $D_\varepsilon^z \in TTPDL$ for all $\varepsilon > 0$. $\square$

The ε -optimal schedulers that we compute in the discretization algorithm in Sec. 5.3.1 yield a special subclass of *PCDL* schedulers, where the time intervals on which the scheduling decision remains constant all have the same length $\tau > 0$. To formally reason about such schedulers, we introduce τ -schedulers as a special subclass of *PCDL* schedulers:

Definition 5.8 (τ -scheduler). Let $\mathcal{C} = (\mathcal{S}, Act, \mathbf{R}, \nu)$ be a CTMDP, $\tau \in \mathbb{R}_{>0}$ and $D \in PCDL$. Then D is a τ -scheduler iff for all $s \in \mathcal{S}$ and $k \in \mathbb{N}$:

$$\exists \alpha \in Act(s). \forall t_\pi \in [k\tau, (k+1)\tau). D(s, t_\pi) = \alpha.$$

Any *PCDL* scheduler is a τ -scheduler if its choices remain constant on intervals of length at least τ . As it turns out, the probabilities induced by *PCDL* and by τ -schedulers converge for small τ :

Theorem 5.4 (Limiting τ -scheduler). Let $\mathcal{C} = (\mathcal{S}, Act, \mathbf{R}, \nu)$ be a CTMDP, $G \subseteq \mathcal{S}$ a set of goal states, $z \in \mathbb{R}_{\geq 0}$ a time bound and $s \in \mathcal{S}$ an initial state. For any scheduler $D \in PCDL$, there exist τ -schedulers D_τ such that

$$\lim_{\tau \rightarrow 0} Pr_{v_s, D_\tau}^\omega \left(\Diamond^{[0, z]} G \right) = Pr_{v_s, D}^\omega \left(\Diamond^{[0, z]} G \right).$$

Proof. As $D \in PCDL$, there exist $n_{s, \alpha} \in \mathbb{N}$ and disjoint intervals $B_{s, \alpha}^0, \dots, B_{s, \alpha}^{n_{s, \alpha}}$ for all $s \in \mathcal{S}$ and $\alpha \in Act$ such that $D(s, t_\pi) = \alpha$ iff $t_\pi \in B_{s, \alpha}^i$ for some $i \leq n_{s, \alpha}$. If $\tau \rightarrow 0$, we can approximate those intervals arbitrarily closely, that is, there exist schedulers D_τ such that $D_\tau(s, \cdot)^{-1}(\alpha) \rightarrow D(s, \cdot)^{-1}(\alpha)$. Similar to the proof of Thm. 5.3, this implies that $\lim_{\tau \rightarrow 0} \mu_{D_\tau} = \mu_D$ and therefore

$$\lim_{\tau \rightarrow 0} Pr_{v_s, D_\tau}^\omega \left(\Diamond^{[0, z]} G \right) = Pr_{v_s, D}^\omega \left(\Diamond^{[0, z]} G \right),$$

proving the claim. $\square$

Example 5.2. Recall the locally uniform CTMDP $\mathcal{C}$ that was used to introduce late schedulers in Sec. 4.4. It is depicted again in Fig. 5.2(a). The ε -optimal scheduler¹ that maximizes

¹The scheduler depicted in Fig. 5.2(b) is the result that is computed by our implementation when maximizing the time-bounded reachability probability for state s_2 with time-bound $z = 4$.

the time-bounded reachability probability for the set $G = \{s_2\}$ of goal states and for time bound $z = 1.5$ is depicted in Fig. 5.2(b). As expected, its decisions coincide with the theoretical derivation that we made in the proof of Thm. 4.6 for the optimal ML-scheduler (see page 109). $\diamond$

5.3 Computing time-bounded reachability probabilities

In the preceding section we have established the theory which is necessary for the main contribution of this chapter. In particular, we will make use of the fixed-point characterization in Thm. 5.1 and the fact (provided by Thm. 5.2) that we may restrict ourselves to *TTPDL* schedulers. With these preliminaries, we are now ready to reduce the problem of computing maximum time-bounded reachability in CTMDPs to the problem of maximizing the step-bounded reachability probability in (discrete-time) MDPs.

The latter is a well-studied problem which can be solved efficiently, e.g. by value iteration algorithms [Ber95]. The discretization that we use for our reduction is defined such that it is exact up to an a priori given error bound $\varepsilon > 0$; hence, the results can be made arbitrarily precise. We study the complexity of our approach and show how to synthesize ε -optimal schedulers automatically.

5.3.1 Discretizing time in locally uniform CTMDPs

As before, let $\mathcal{C}$ be a locally uniform CTMDP, $G \subseteq \mathcal{S}$ a set of goal states, $s \in \mathcal{S}$ an initial state and $z \in \mathbb{R}_{\geq 0}$ a time bound. We aim at computing $p_{\max}(s, z)$ up to an a priori fixed error $\varepsilon > 0$. If $s \in G$, this is trivial as $p_{\max}(s, z) = 1$ for all $z \in \mathbb{R}_{\geq 0}$. To compute $p_{\max}(s, z)$ for $s \notin G$, we use the fixed point characterization of $p_{\max}$ from Thm. 5.1. More precisely, consider the first sub-interval $[0, \tau]$ of the integral in Eq. (5.4) separately and split the whole integral accordingly:

$$\begin{aligned} p_{\max}(s, z) &= \Omega(p_{\max})(s, z) \\ &= \int_0^\tau E(s) e^{-E(s)t} \cdot \max_{\alpha \in \text{Act}} \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot p_{\max}(s', z - t) dt \\ &\quad + \int_\tau^z E(s) e^{-E(s)t} \cdot \max_{\alpha \in \text{Act}} \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot p_{\max}(s', z - t) dt. \end{aligned} \tag{5.6}$$

Now, let $A(s, z)$ and $B(s, z)$ denote the first, resp. second summand in Eq. (5.6). Shifting the range of integration in $B(s, z)$ by $(-\tau)$, the next Lemma derives a straightforward recursive representation of the probability $B(s, z)$ which can easily be used for our discretization purposes:

Lemma 5.3. For all $s \in S$, $z \in \mathbb{R}_{\geq 0}$ and $\tau \in [0, z]$ it holds that

$$B(s, z) = e^{-E(s)\tau} \cdot p_{\max}(s, z - \tau). \quad (5.7)$$

Proof. We proceed as follows:

$$\begin{aligned} B(s, z) &= \int_{\tau}^z E(s) e^{-E(s)t} \cdot \max_{\alpha \in Act} \sum_{s' \in S} \mathbf{P}(s, \alpha, s') \cdot p_{\max}(s', z - t) dt \\ &= \int_0^{z-\tau} E(s) e^{-E(s)(t+\tau)} \cdot \max_{\alpha \in Act} \sum_{s' \in S} \mathbf{P}(s, \alpha, s') \cdot p_{\max}(s', z - (t + \tau)) dt \\ &= \int_0^{z-\tau} E(s) e^{-E(s)t} e^{-E(s)\tau} \cdot \max_{\alpha \in Act} \sum_{s' \in S} \mathbf{P}(s, \alpha, s') \cdot p_{\max}(s', z - (t + \tau)) dt \\ &= e^{-E(s)\tau} \cdot \underbrace{\int_0^{z-\tau} E(s) e^{-E(s)t} \cdot \max_{\alpha \in Act} \sum_{s' \in S} \mathbf{P}(s, \alpha, s') \cdot p_{\max}(s', (z - \tau) - t) dt}_{p_{\max}(s, z - \tau)} \\ &= e^{-E(s)\tau} \cdot p_{\max}(s, z - \tau). \quad \square \end{aligned}$$

Note that the factor $e^{-E(s)\tau}$ in Eq. (5.7) is the probability that no transition occurs in state s in the first τ time units. Hence, $B(s, z)$ is the maximum probability of the event that starting from state s , the set G is reached within z time units while no transition occurs in the time interval $[0, \tau]$. To be more precise, let $\#_{[0, \tau]} : Paths^{\omega} \rightarrow \mathbb{N}$ be the random variable which describes the number of transitions that occur in time interval $[0, \tau]$. Then, it holds that $B(s, z) = \sup_{D \in ML} Pr_{v_s, D}^{\omega}(\Diamond^{[0, z]} G \cap \#_{[0, \tau]} = 0)$. With the same reasoning, the first summand $A(s, z)$ of (5.6) is the maximum probability to reach G within time z with at least one transition taking place in $[0, \tau]$. Hence,

$$A(s, z) = \sup_{D \in ML} Pr_{v_s, D}^{\omega}(\Diamond^{[0, z]} G \cap \#_{[0, \tau]} \geq 1).$$

Now, decompose the underlying event of $A(s, z)$ into disjoint subsets according to the number of transitions that occur in time interval $[0, \tau]$:

$$(\Diamond^{[0, z]} G \cap \#_{[0, \tau]} \geq 1) = \bigcup_{n=1}^{\infty} (\Diamond^{[0, z]} G \cap \#_{[0, \tau]} = n).$$

Accordingly, let $A_n(s, z)$ be the maximum probability to reach G in z time units with exactly n transitions occurring in the first time slice $[0, \tau]$. In this way, we maximize the probability of each event $(\Diamond^{[0, z]} G \cap \#_{[0, \tau]} = n)$ separately:

$$A_n(s, z) = \sup_{D \in ML} Pr_{v_s, D}^{\omega}(\Diamond^{[0, z]} G \cap \#_{[0, \tau]} = n). \quad (5.8)$$

To relate $A(s, z)$ with the probabilities $A_n(s, z)$, observe that

$$\begin{aligned}
 A(s, z) &= \sup_{D \in ML} Pr_{v_s, D}^\omega \left(\diamond^{[0, z]} G \cap \#_{[0, \tau]} \geq 1 \right) \\
 &= \sup_{D \in ML} Pr_{v_s, D}^\omega \left(\bigcup_{n=1}^{\infty} \left(\diamond^{[0, z]} G \cap \#_{[0, \tau]} = n \right) \right) \\
 &\leq \sum_{n=1}^{\infty} \left(\sup_{D \in ML} Pr_{v_s, D}^\omega \left(\diamond^{[0, z]} G \cap \#_{[0, \tau]} = n \right) \right) \\
 &= \sum_{n=1}^{\infty} A_n(s, z).
 \end{aligned} \tag{5.9}$$

The next major step is to derive an analytic expression for the probability $A_1(s, z)$:

Lemma 5.4. *Let $\mathcal{C} = (S, Act, \mathbf{R}, v)$ be a CTMDP, $G \subseteq S$ a set of goal states, $s \in S$ an initial state, $z \in \mathbb{R}_{\geq 0}$ a time bound and $\tau > 0$ a step duration. For $A_1(s, z)$ as defined in Eq. (5.8) it holds*

$$A_1(s, z) = \int_0^\tau E(s) e^{-E(s)t} \cdot \max_{\alpha \in Act} \sum_{s' \in S} \mathbf{P}(s, \alpha, s') \cdot e^{-E(s')(\tau-t)} \cdot p_{\max}(s', z - \tau) dt. \tag{5.10}$$

Note that $A_1(s, z)$ is the maximum probability to reach G within z time units and that exactly one transition occurs within time interval $[0, \tau]$. This is reflected in the integral in Lemma 5.4: Here, the integration variable t corresponds to the precise point in time when state s is left; further, if we move to state s' after t units of time, we stay in the successor state s' for at least $(\tau - t)$ time units (i.e. the time that remains in the first step duration) with probability $e^{-E(s')(\tau-t)}$. Finally, we multiply with $p_{\max}(s', z - \tau)$, i.e. with the maximum achievable probability to reach G in the remaining $(z - \tau)$ time units, starting in state s' .

Proof. Let $E = (\diamond^{[0, z]} G \cap \#_{[0, \tau]} = 1)$ be the event that corresponds to the probability $A_1(s, z)$. Given an ML scheduler D , the measure of the event E differs from the time-bounded reachability event $\diamond^{[0, z]} G$ in the additional requirement that exactly one transition occurs in time interval $[0, \tau]$. Hence, we obtain the probability

$$\begin{aligned}
 Pr_{v_s, D}^\omega \left(\diamond^{[0, z]} G \cap \#_{[0, \tau]} = 1 \right) &= \int_0^\tau E(s) e^{-E(s)t} \cdot \sum_{\alpha \in Act} D(s, t)(\alpha) \\
 &\quad \cdot \sum_{s' \in S} \mathbf{P}(s, \alpha, s') \cdot e^{-E(s')(\tau-t)} \cdot Pr_{v_{s'}, D(s \xrightarrow{\alpha, t} s', +(\tau-t))}^\omega \left(\diamond^{[0, z-\tau]} G \right) dt. \tag{5.11}
 \end{aligned}$$

The term $e^{-E(s')(\tau-t)}$ in Eq. (5.11) is the probability that after leaving state s at time point t and entering the successor state s' , no transition occurs for the next $(\tau - t)$ time units.

The *ML*-scheduler $D(s \xrightarrow{\alpha, t} \cdot, \cdot^{+(\tau-t)})$ is defined such that if $\pi = s' \xrightarrow{\alpha', t'} \pi''$ for some $\pi'' \in \text{Paths}^*$, then $D(s \xrightarrow{\alpha, t} \cdot, \cdot^{+(\tau-t)})(\pi, t) = D(s \xrightarrow{\alpha, t} \pi', t)$, where $\pi' = s' \xrightarrow{\alpha', t' + (\tau-t)} \pi''$.

Hence, if state s is left at time t and no transition occurs in the successor state s' within the following $\tau - t$ time units, then $D(s \xrightarrow{\alpha, t} \cdot, \cdot^{+(\tau-t)})$ decides on the remaining path as D does on the suffix of the complete path. Note that due to the memoryless property of the exponential distribution, we may split the sojourn time in state s' in two parts: First, the sojourn in state s' before τ and the remaining sojourn time. Hence Eq. (5.11) expresses the probability to reach G from state s within time bound z and that *exactly one* transition occurs in time interval $[0, \tau]$.

With these preliminaries, we introduce the *ML*-scheduler D_1^z , which induces the maximum probability for the event E . Similar to the scheduler D^z , it is deterministic; however, it is not fully positional: To ease its definition, let $g(s, \alpha, t) \in [0, 1]$ be the maximum probability to reach G in z time units, if state s has been left at time t and action α has been chosen and no transition occurs in the remaining $\tau - t$ time units:

$$g(s, \alpha, t) = \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot e^{-E(s')(\tau-t)} \cdot \sup_{D' \in \text{ML}} \text{Pr}_{v_{s'}, D'}^\omega (\diamond^{[0, z-\tau]} G).$$

We obtain $D_1^z : \text{Paths}^* \times \mathbb{R}_{\geq 0} \rightarrow \text{Act}$ as follows: If $|\pi| = 0$, then $\pi = s$ for some $s \in \mathcal{S}$ and $D_1^z(s, t) = \min_{\prec} \{\alpha \in \text{Act}(s) \mid \forall \beta \in \text{Act}(s). g(s, \beta, t) \leq g(s, \alpha, t)\}$. Otherwise, we know that at least one transition has occurred. Hence, we define D_1^z such that it optimizes the probability to reach G in the remaining time $z - (\Delta(\pi) + t)$. Therefore we set $D_1^z(\pi, t) = D^z(\pi \downarrow, \Delta(\pi) + t)$ if $|\pi| > 0$.

Now we prove that D_1^z is optimal w.r.t. E by contraposition: Assume that there exists $D' \in \text{ML}$ such that $\text{Pr}_{v_s, D'}^\omega(E) > \text{Pr}_{v_s, D_1^z}^\omega(E)$. By the following derivation, this leads to a contradiction:

$$\begin{aligned} \text{Pr}_{v_s, D'}^\omega(E) &= \int_0^\tau E(s) e^{-E(s)t} \cdot \sum_{\alpha \in \text{Act}} D'(s, t)(\alpha) \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot e^{-E(s')(\tau-t)} \\ &\quad \cdot \text{Pr}_{v_{s'}, D'}^\omega(s \xrightarrow{\alpha, t} \cdot, \cdot^{+(\tau-t)}) (\diamond^{[0, z-\tau]} G) dt \\ &\leq \int_0^\tau E(s) e^{-E(s)t} \cdot \sum_{\alpha \in \text{Act}} D'(s, t)(\alpha) \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot e^{-E(s')(\tau-t)} \\ &\quad \cdot \sup_{D \in \text{ML}} \text{Pr}_{v_{s'}, D}^\omega (\diamond^{[0, z-\tau]} G) dt \\ &\leq \int_0^\tau E(s) e^{-E(s)t} \\ &\quad \cdot \underbrace{\max_{\alpha \in \text{Act}} \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot e^{-E(s')(\tau-t)} \cdot \sup_{D \in \text{ML}} \text{Pr}_{v_{s'}, D}^\omega (\diamond^{[0, z-\tau]} G)}_{g(s, \alpha, t)} dt \\ &= \int_0^\tau E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, D_1^z(s, t), s') \cdot e^{-E(s')(\tau-t)} \cdot \sup_{D \in \text{ML}} \text{Pr}_{v_{s'}, D}^\omega (\diamond^{[0, z-\tau]} G) dt \end{aligned}$$

$$\begin{aligned}
&= \int_0^\tau E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, D_1^z(s, t), s') \cdot e^{-E(s')(\tau-t)} \cdot Pr_{v_{s'}, D^{z-\tau}}^\omega (\diamond^{[0, z-\tau]} G) dt \\
&= Pr_{v_s, D_1^z}^\omega(E).
\end{aligned}$$

Hence, the scheduler D_1^z yields the maximum probability for the event E and we obtain

$$\begin{aligned}
A_1(s, z) &= \sup_{D \in ML} Pr_{v_s, D}^\omega (\diamond^{[0, z]} G \cap \#_{[0, \tau]} = 1) \\
&= \int_0^\tau E(s) e^{-E(s)t} \cdot \sum_{\alpha \in Act} D_1^z(s, t)(\alpha) \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot e^{-E(s')(\tau-t)} \\
&\quad \cdot Pr_{v_{s'}, D_1^z(s \xrightarrow{\alpha, t}, \cdot, +(\tau-t))}^\omega (\diamond^{[0, z-\tau]} G) dt \\
&= \int_0^\tau E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, D_1^z(s, t), s') \cdot e^{-E(s')(\tau-t)} \\
&\quad \cdot Pr_{v_{s'}, D^z(\cdot, \cdot, +\tau)}^\omega (\diamond^{[0, z-\tau]} G) dt \\
&= \int_0^\tau E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, D_1^z(s, t), s') \cdot e^{-E(s')(\tau-t)} \\
&\quad \cdot Pr_{v_{s'}, D^{z-\tau}}^\omega (\diamond^{[0, z-\tau]} G) dt \\
&= \int_0^\tau E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, D_1^z(s, t), s') \cdot e^{-E(s')(\tau-t)} \cdot p_{\max}(s', z - \tau) dt \\
&= \int_0^\tau E(s) e^{-E(s)t} \cdot \max_{\alpha \in Act} \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot e^{-E(s')(\tau-t)} \cdot p_{\max}(s', z - \tau) dt,
\end{aligned}$$

completing the proof. $\square$

Now we approximate the probability $A(s, z)$ from below via a new probability $X(s, z)$, which is closely related to $A_1(s, z)$: More precisely, we obtain $X(s, z)$ by bounding the probability $e^{-E(s')(\tau-t)}$ in Eq. (5.10) from above by 1. Hence $A_1(s, z) \leq X(s, z)$; moreover, by a continuity argument we can prove that $X(s, z) \leq A(s, z)$.

With these two inequalities and the definition of $X(s, z)$ we establish an error bound for our discretization. Formally, the following sandwich lemma proves that $X(s, z)$ converges to $A(s, z)$ for $\tau \rightarrow 0$:

Lemma 5.5 (One-step approximation). *Let $\mathcal{C} = (\mathcal{S}, Act, \mathbf{R}, v)$ be a CTMDP, $G \subseteq \mathcal{S}$ a set of goal states, $\lambda = \max_{s \in \mathcal{S}} E(s)$, $s \in \mathcal{S}$ an initial state, $z \in \mathbb{R}_{\geq 0}$ a time bound and $\tau > 0$ a step duration. If we define*

$$X(s, z) = \int_0^\tau E(s) e^{-E(s)t} \cdot \max_{\alpha \in Act} \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot p_{\max}(s', z - \tau) dt, \quad (5.12)$$

then $X(s, z)$ approximates $A(s, z)$ in the following sense:

$$X(s, z) \leq A(s, z) \leq X(s, z) + \frac{(\lambda\tau)^2}{2}. \quad (5.13)$$

Proof. To establish the lower bound in Eq. (5.13), it suffices to note that

$$A(s, z) \stackrel{(5.6)}{=} \int_0^\tau E(s) e^{-E(s)t} \cdot \max_{\alpha \in Act} \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot \underbrace{p_{\max}(s', z - t)}_{\geq p_{\max}(s', z - \tau)} dt.$$

By definition, for all $s' \in \mathcal{S}$, the function $p_{\max}(s', \cdot)$ is monotonically increasing in its second argument, that is, for increasing time bounds z , the maximum reachability probability $p_{\max}(s', z)$ increases. Reversely, the function $p_{\max}(s', z - t)$ is monotonically decreasing for increasing t .

Hence $t < \tau$ implies that $p_{\max}(s', z - t) \geq p_{\max}(s', z - \tau)$ and we obtain

$$A(s, z) \geq \int_0^\tau E(s) e^{-E(s)t} \cdot \max_{\alpha \in Act} \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot p_{\max}(s', z - \tau) dt \stackrel{(5.12)}{=} X(s, z).$$

For the upper bound in Eq. (5.13), let us first investigate the relation between $X(s, z)$ and $A_1(s, z)$. By Lemma 5.4, we derive:

$$\begin{aligned} A_1(s, z) &= \int_0^\tau E(s) e^{-E(s)t} \cdot \max_{\alpha \in Act} \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot \underbrace{e^{-E(s')(z-t)}}_{\leq 1} \cdot p_{\max}(s', z - \tau) dt \\ &\leq \int_0^\tau E(s) e^{-E(s)t} \cdot \max_{\alpha \in Act} \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot p_{\max}(s', z - \tau) dt \\ &\stackrel{(5.12)}{=} X(s, z). \end{aligned}$$

Therefore, we have proved that $X(s, z)$ is an upper bound for $A_1(s, z)$; formally:

$$A_1(s, z) \leq X(s, z). \quad (5.14)$$

In the next step, we also obtain an upper bound for the sum $\sum_{n=2}^\infty A_n(s, z)$: To see how this works, recall that for an exponential distribution with rate λ and a time interval $[0, \tau]$, the Poisson distribution $\rho(n, \lambda\tau) = e^{-\lambda\tau} \cdot \frac{(\lambda\tau)^n}{n!}$ expresses the probability that n transitions occur within $[0, \tau]$. This allows us to derive an upper bound, first for each $A_n(s, z)$ separately:

$$\begin{aligned} A_n(s, z) &= \sup_{D \in ML} Pr_{v_s, D}^\omega (\diamond^{[0, z]} G \cap \#_{[0, \tau]} = n) \leq \sup_{D \in ML} Pr_{v_s, D}^\omega (\#_{[0, \tau]} = n) \\ &\leq \rho(n, \lambda\tau) = e^{-\lambda\tau} \cdot \frac{(\lambda\tau)^n}{n!}. \end{aligned} \quad (5.15)$$

Moreover, by maximality of λ , the probability that more than n transitions occur in any state $s \in \mathcal{S}$ within time interval $[0, \tau]$ is at most

$$\sum_{i=n+1}^\infty \rho(i, \lambda\tau) = e^{-\lambda\tau} \sum_{i=n+1}^\infty \frac{(\lambda\tau)^i}{i!} = e^{-\lambda\tau} \cdot R_n(\lambda\tau), \quad (5.16)$$

where $R_n(x) = \sum_{i=n+1}^{\infty} \frac{x^i}{i!}$ is the remainder term of the Taylor expansion of $f(x) = e^x$ for the point 0. By Taylor's theorem, there exists $\xi \in [0, \lambda\tau]$ such that

$$R_n(\lambda\tau) = \frac{f^{(n+1)}(\xi)}{(n+1)!} \cdot (\lambda\tau)^{n+1} = \frac{e^\xi}{(n+1)!} \cdot (\lambda\tau)^{n+1}, \quad (5.17)$$

where $f^{(n+1)}$ denotes the $(n+1)$ -th derivative of f .

Summarizing the above reasoning, we obtain an upper bound for the error that is induced by approximating $A(s, z)$ by only considering $X(s, z)$:

$$\begin{aligned} A(s, z) &\stackrel{(5.9)}{\leq} \sum_{n=1}^{\infty} A_n(s, z) \stackrel{(5.14)}{\leq} X(s, z) + \sum_{n=2}^{\infty} A_n(s, z) \stackrel{(5.15)}{\leq} X(s, z) + \sum_{n=2}^{\infty} \rho(n, \lambda\tau) \\ &\stackrel{(5.16)}{=} X(s, z) + e^{-\lambda\tau} \cdot R_1(\lambda\tau). \end{aligned}$$

By Taylor's theorem and Eq. (5.17), there exists $\xi \in [0, \lambda\tau]$ such that $R_1(\lambda\tau) = \frac{e^\xi}{2} \cdot (\lambda\tau)^2$. For an upper bound, choose ξ maximal in $[0, \lambda\tau]$. Then

$$A(s, z) \leq X(s, z) + e^{-\lambda\tau} \cdot R_1(\lambda\tau) \leq X(s, z) + e^{-\lambda\tau} \cdot \frac{e^{\lambda\tau}}{2} (\lambda\tau)^2 = X(s, z) + \frac{(\lambda\tau)^2}{2}.$$

Thus we have $A(s, z) \leq X(s, z) + \frac{(\lambda\tau)^2}{2}$, completing the proof for the upper bound. $\square$

By Eq. (5.13), we can approximate $A(s, z)$ from below via $X(s, z)$, allowing for an error of at most $\frac{(\lambda\tau)^2}{2}$. Thus, for $\tau \rightarrow 0^+$ it holds that $X(s, z) = A(s, z)$. We use the one-step error bound $\frac{(\lambda\tau)^2}{2}$ later in Thm. 5.5 to derive the overall error bound for our analysis.

5.3.2 Reduction to step-bounded reachability in MDPs

Based on $X(s, z)$ and $B(s, z)$, we are now ready to derive a *discretization* for $p_{\max}(s, z)$ in a locally uniform CTMDP $\mathcal{C}$ with respect to a *step duration* τ :

Definition 5.9 (Discretization). Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ be a CTMDP, and let $\tau > 0$ be a step duration. The induced MDP $\mathcal{C}_\tau = (\mathcal{S}, \text{Act}, \mathbf{P}_\tau, \nu)$ is defined such that for all $s, s' \in \mathcal{S}$ and $\alpha \in \text{Act}(s)$:

$$\mathbf{P}_\tau(s, \alpha, s') = \begin{cases} (1 - e^{-E(s)\tau}) \cdot \mathbf{P}(s, \alpha, s') & \text{if } s \neq s' \\ (1 - e^{-E(s)\tau}) \cdot \mathbf{P}(s, \alpha, s') + e^{-E(s)\tau} & \text{if } s = s'. \end{cases}$$

Further, for all $\alpha \notin \text{Act}(s)$, we define $\mathbf{P}_\tau(s, \alpha, s') = 0$.

In the MDP $\mathcal{C}_\tau$, each step corresponds to one time slice of length τ in the CTMDP $\mathcal{C}$. For a single step and a fixed successor state $s' \neq s$, $\mathbf{P}_\tau(s, \alpha, s')$ equals the probability that

a transition to s' occurs within τ time units, given that α is chosen. In case that $s' = s$, the first summand of $\mathbf{P}_\tau(s, \alpha, s)$ is the probability to take the loop back to s ; the second summand denotes the probability that no transition occurs within time τ and thus $s = s'$.

Let $p_{\max}^{\mathcal{C}_\tau}(s, k)$ be the maximum probability to reach G starting from state s in at most k discrete steps in the (discrete time) MDP $\mathcal{C}_\tau$. Obviously $p_{\max}^{\mathcal{C}_\tau}(s, k) = 1$ if $s \in G$ and $p_{\max}^{\mathcal{C}_\tau}(s, 0) = 0$ if $s \notin G$. Further, for $s \notin G$ and $k > 0$, $p_{\max}^{\mathcal{C}_\tau}(s, k)$ is defined recursively:

$$p_{\max}^{\mathcal{C}_\tau}(s, k) = \max_{\alpha \in \text{Act}} \sum_{s' \in \mathcal{S}} \mathbf{P}_\tau(s, \alpha, s') \cdot p_{\max}^{\mathcal{C}_\tau}(s', k-1). \quad (5.18)$$

The next theorem proves that the probability to reach G from state s within at most $k = \frac{z}{\tau}$ steps in the discrete-time MDP $\mathcal{C}_\tau$ converges from below (for $\tau \rightarrow 0$) to the corresponding time-bounded reachability probability in the CTMDP $\mathcal{C}$:

Theorem 5.5. *Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ be a CTMDP, $\lambda = \max_{s \in \mathcal{S}} E(s)$, $G \subseteq \mathcal{S}$ a set of goal states, $z \in \mathbb{R}_{\geq 0}$ a time bound and $k \in \mathbb{N}_{>0}$ the number of discretization steps, such that $\tau = \frac{z}{k}$. Then it holds for all $s \in \mathcal{S}$:*

$$p_{\max}^{\mathcal{C}_\tau}(s, k) \leq p_{\max}^{\mathcal{C}}(s, z) \leq p_{\max}^{\mathcal{C}_\tau}(s, k) + \frac{(\lambda z)^2}{2k}. \quad (5.19)$$

The proof is by induction on the number k of discretization steps, where the lower and upper bounds are established for each step of length τ using Lemma 5.4 and Lemma 5.5.

Proof. Recall that $p_{\max}^{\mathcal{C}}(s, z) = A(s, z) + B(s, z)$ and $X(s, z) \leq A(s, z) \leq X(s, z) + \frac{(\lambda z)^2}{2}$ by Eq. (5.13). We prove Eq. (5.19) by induction on k :

1. For $k = 1$, we have $z = \tau$. If $s \in G$, then $p_{\max}^{\mathcal{C}_\tau}(s, 1) = 1 = p_{\max}^{\mathcal{C}}(s, \tau)$, proving (5.19); if $k = 1$ and $s \notin G$, the lower bound in (5.19) holds as $p_{\max}^{\mathcal{C}_\tau}(s, 1) = \max_{\alpha \in \text{Act}} (1 - e^{-E(s)\tau}) \cdot \mathbf{P}(s, \alpha, G) = X(s, \tau) \leq p_{\max}^{\mathcal{C}}(s, \tau)$. For the upper bound, note that $s \notin G$ implies $B(s, \tau) = 0$. Thus $p_{\max}^{\mathcal{C}}(s, \tau) = A(s, \tau) + B(s, \tau) = A(s, \tau)$. By Lemma 5.5, we know that $A(s, \tau) \leq X(s, \tau) + \frac{(\lambda \tau)^2}{2}$. Moreover, $X(s, \tau) = p_{\max}^{\mathcal{C}_\tau}(s, \tau)$ by definition. Therefore $p_{\max}^{\mathcal{C}}(s, \tau) \leq p_{\max}^{\mathcal{C}_\tau}(s, \tau) + \frac{(\lambda \tau)^2}{2}$.
2. For the induction step, together with Lemma 5.5 (which provides $X(s, z)$) and Lemma 5.3 (the analytic expression for $B(s, z)$) we have

$$\begin{aligned} X(s, z) + B(s, z) &= \left[\int_0^\tau E(s) e^{-E(s)t} \max_{\alpha \in \text{Act}} \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot p_{\max}^{\mathcal{C}}(s', z - \tau) dt \right] \\ &\quad + \left[e^{-E(s)\tau} \cdot p_{\max}^{\mathcal{C}}(s, z - \tau) \right] \end{aligned}$$

$$\begin{aligned}
&= \left[\max_{\alpha \in Act} \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot p_{\max}^C(s', z - \tau) \cdot \int_0^\tau E(s) e^{-E(s)t} dt \right] \\
&\quad + \left[e^{-E(s)\tau} \cdot p_{\max}^C(s, z - \tau) \right] \\
&= \max_{\alpha \in Act} \left[\left(1 - e^{-E(s)\tau} \right) \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, \alpha, s') \cdot p_{\max}^C(s', z - \tau) \right] \quad (5.20) \\
&\quad + \left[e^{-E(s)\tau} \cdot p_{\max}^C(s, z - \tau) \right]
\end{aligned}$$

By definition of $\mathbf{P}_\tau(s, \alpha, s')$ (where the second summand in Eq. (5.20) corresponds to the special case of $s = s'$), we derive from Eq. (5.20):

$$X(s, z) + B(s, z) = \max_{\alpha \in Act} \sum_{s' \in \mathcal{S}} \mathbf{P}_\tau(s, \alpha, s') \cdot p_{\max}^C(s', z - \tau). \quad (5.21)$$

First we consider the lower bound on the left part of Eq. (5.19): By induction hypothesis, it holds that $p_{\max}^{C_\tau}(s', k - 1) \leq p_{\max}^C(s', z - \tau)$ for all $s' \in \mathcal{S}$. Then

$$\begin{aligned}
p_{\max}^C(s, z) &\geq X(s, z) + B(s, z) \\
&\stackrel{(5.21)}{=} \max_{\alpha \in Act} \sum_{s' \in \mathcal{S}} \mathbf{P}_\tau(s, \alpha, s') \cdot p_{\max}^C(s', z - \tau) \\
&\stackrel{i.h.}{\geq} \max_{\alpha \in Act} \sum_{s' \in \mathcal{S}} \mathbf{P}_\tau(s, \alpha, s') \cdot p_{\max}^{C_\tau}(s', k - 1) = p_{\max}^{C_\tau}(s, k).
\end{aligned}$$

The proof for the upper bound is as follows: By Lemma 5.5 it holds that $A(s, z) \leq X(s, z) + \frac{(\lambda\tau)^2}{2}$. Together with Eq. (5.21) we derive

$$\begin{aligned}
p_{\max}^C(s, z) &= A(s, z) + B(s, z) \\
&\leq X(s, z) + \frac{(\lambda\tau)^2}{2} + B(s, z) \\
&\stackrel{(5.21)}{=} \frac{(\lambda\tau)^2}{2} + \max_{\alpha \in Act} \sum_{s' \in \mathcal{S}} \mathbf{P}_\tau(s, \alpha, s') \cdot p_{\max}^C(s', z - \tau).
\end{aligned}$$

Applying the induction hypothesis, we obtain

$$\begin{aligned}
p_{\max}^C(s, z) &\leq \frac{(\lambda\tau)^2}{2} + \max_{\alpha \in Act} \sum_{s' \in \mathcal{S}} \mathbf{P}_\tau(s, \alpha, s') \left(p_{\max}^{C_\tau}(s', k - 1) + \frac{(\lambda(z - \tau))^2}{2(k - 1)} \right) \\
&= \frac{(\lambda\tau)^2}{2} + \frac{(\lambda(z - \tau))^2}{2(k - 1)} + \max_{\alpha \in Act} \sum_{s' \in \mathcal{S}} \mathbf{P}_\tau(s, \alpha, s') \cdot p_{\max}^{C_\tau}(s', k - 1). \quad (5.22)
\end{aligned}$$

From here, we complete the induction step: Therefore, rewrite the summands $\frac{(\lambda\tau)^2}{2}$ and $\frac{(\lambda(z - \tau))^2}{2(k - 1)}$ in the right part of Eq. (5.22) further:

$$\frac{(\lambda\tau)^2}{2} + \frac{(\lambda(z - \tau))^2}{2(k - 1)} = \frac{(\lambda\tau)^2 k(k - 1) + (\lambda(z - \tau))^2 k}{2k(k - 1)} \quad (* \text{ as } k = \frac{z}{\tau} *)$$

$$\begin{aligned}
&= \frac{(\lambda\tau)^2 \cdot \frac{z}{\tau} \cdot \frac{z-\tau}{\tau} + (\lambda(z-\tau))^2 \cdot \frac{z}{\tau}}{2k \cdot \frac{z-\tau}{\tau}} \\
&= \frac{\lambda^2 z(z-\tau) + \lambda^2 (z-\tau)^2 \cdot \frac{z}{\tau}}{2k \cdot \frac{z-\tau}{\tau}} \\
&= \frac{\lambda^2 \tau z + \lambda^2 z(z-\tau)}{2k} = \frac{\lambda^2 (\tau z + z^2 - \tau z)}{2k} = \frac{(\lambda z)^2}{2k}.
\end{aligned}$$

In this way, the right part of Eq. (5.22) can be simplified to $p_{\max}^{\mathcal{C}_\tau}(s, k) + \frac{(\lambda z)^2}{2k}$. $\square$

Example 5.3. Consider the CTMDP $\mathcal{C}$ in Fig. 5.3(a). To compute the maximum probability to reach $G = \{s_2\}$ within z time units up to a precision of ε , choose $k \in \mathbb{N}$ such that $\varepsilon \geq \frac{(\lambda z)^2}{2k}$, where $\lambda = \max_{s \in \mathcal{S}} E(s) = 3$. The step duration $\tau = \frac{z}{k}$ induces the discretized MDP $\mathcal{C}_\tau$ which is depicted in Fig. 5.3(b). $\diamond$

5.3.3 Algorithm and complexity

Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ be a locally uniform CTMDP, G a set of goal states and z a time bound. For some error bound $\varepsilon > 0$, let k be the number of steps needed to satisfy $\varepsilon \geq \frac{(\lambda z)^2}{2k}$. Then $\tau = \frac{z}{k}$ induces the discretized MDP $\mathcal{C}_\tau$ of $\mathcal{C}$ with step duration τ . By Thm. 5.5, the maximum probability to reach G within z time units in $\mathcal{C}$ can be approximated (up to ε) by maximizing the step-bounded reachability $p_{\max}^{\mathcal{C}_\tau}$ for G in $\mathcal{C}_\tau$ within k steps. The latter can be computed efficiently by the well-known *value iteration* approach [Ber95]. Briefly, it starts with a probability vector $\vec{v}_0$ with $\vec{v}_0(s) = 1$ if $s \in G$ and 0, otherwise. In each iteration, $\vec{v}_i$ is obtained from $\vec{v}_{i-1}$ according to Eq. (5.18). In each round, i corresponds to the number of steps in the MDP $\mathcal{C}_\tau$; hence, $\vec{v}_i(s)$ equals $p_{\max}^{\mathcal{C}_\tau}(s, i)$.

The value iteration approach on the discretized MDP $\mathcal{C}_\tau$ has the following complexity. For $s \in \mathcal{S}$ and $\alpha \in \text{Act}(s)$, let $\text{post}(s, \alpha) = \{s' \in \mathcal{S} \mid \mathbf{R}(s, \alpha, s') > 0\}$ be the set of α -successors of state s . The size of $\mathcal{C}$ is denoted by $m = \sum_{s \in \mathcal{S}} \sum_{\alpha \in \text{Act}} |\text{post}(s, \alpha)|$. In the worst case, $\mathcal{C}_\tau$ is obtained by adding a self-loop for each state $s \in \mathcal{S}$ and action $\alpha \in \text{Act}(s)$. Thus, the size of $\mathcal{C}_\tau$ is bounded by $2m$. For a given error bound ε , it is easy to derive the number k of value-iteration steps: By Thm. 5.5, $|p_{\max}^{\mathcal{C}}(s, z) - p_{\max}^{\mathcal{C}_\tau}(s, k)| \leq \frac{(\lambda z)^2}{2k}$. Letting $\frac{(\lambda z)^2}{2k} \leq \varepsilon$, we conclude that the smallest k to guarantee ε is $\frac{(\lambda z)^2}{2\varepsilon}$. In each value iteration step, the update of the vector $\vec{v}_i$ takes time $2m$. Thus, the worst-case time complexity of our approach is $\mathcal{O}(m \cdot (\lambda z)^2 / \varepsilon)$.

5.3.4 Synthesis of ε -optimal schedulers

Let $\mathcal{C}$, G , z , k , $\tau = \frac{z}{k}$ and $\mathcal{C}_\tau$ be as before. A byproduct of the value iteration on the discretized MDP $\mathcal{C}_\tau$ is an ε -optimal scheduler for the set of goal states G and time bound z . More precisely, in any of the i value iteration steps, for each state $s \in \mathcal{S}$, an action $\alpha_{s,i}$

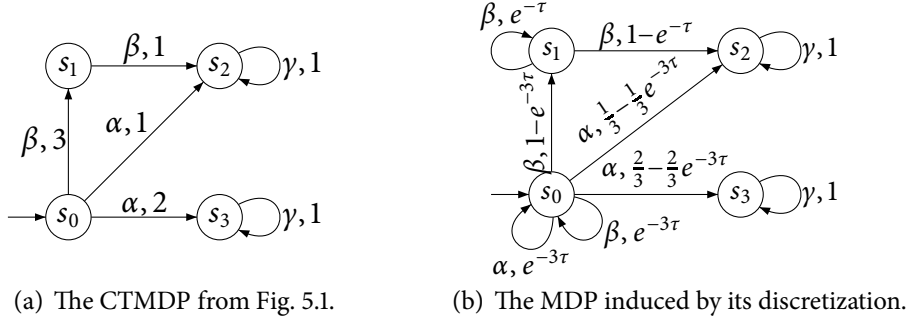


Figure 5.3: The discretization of a locally uniform CTMDP.

is chosen according to Eq. (5.18). In this way, we obtain a history-dependent (or, to be more precise, *step-dependent*) scheduler for the MDP $\mathcal{C}_\tau$. This scheduler induces a τ -scheduler (denoted D_τ) of the original CTMDP $\mathcal{C}$ as follows: $D_\tau(s, t_\pi) = \alpha_{s,i}$ if $t_\pi \in [(k-i)\tau, (k-i+1)\tau)$. The following theorem shows that D_τ is an ε -optimal scheduler in the underlying CTMDP $\mathcal{C}$:

Theorem 5.6 (ε -optimal scheduler). *The scheduler D_τ is an ε -optimal scheduler for $\mathcal{C}$ w.r.t. the maximum time-bounded reachability probability.*

Proof. Let $\mathcal{C} = (\mathcal{S}, Act, \mathbf{R}, \nu)$ be a locally uniform CTMDP, G a set of goal states and z a time bound. For some error bound $\varepsilon > 0$, let k be the number of steps needed to satisfy $\varepsilon \geq \frac{(\lambda z)^2}{2k}$. Let $\mathcal{C}_\tau$ be the induced MDP with $\tau = \frac{z}{k}$, and D_τ be the τ -scheduler as described. To show that D_τ is an ε -optimal scheduler for $\mathcal{C}$ w.r.t. the maximum time-bounded reachability probability, we prove that for all states $s \in \mathcal{S}$ it holds that

$$|Pr_{\nu_s, D_\tau}^\omega(\Diamond^{[0,z]}G) - p_{max}^{\mathcal{C}_\tau}(s, k)| \leq \varepsilon.$$

It is sufficient to show the following equality:

$$p_{max}^{\mathcal{C}_\tau}(s, k) \leq Pr_{\nu_s, D_\tau}^\omega(\Diamond^{[0,z]}G) \leq p_{max}^{\mathcal{C}_\tau}(s, k) + \varepsilon. \quad (5.23)$$

By Theorem 5.5, the upper bound can be shown directly:

$$Pr_{\nu_s, D_\tau}^\omega(\Diamond^{[0,z]}G) \leq p_{max}^{\mathcal{C}}(s, z) \leq p_{max}^{\mathcal{C}_\tau}(s, k) + \frac{(\lambda z)^2}{2k} \leq p_{max}^{\mathcal{C}_\tau}(s, k) + \varepsilon.$$

Now we discuss how to show the lower bound of Eq. (5.23). First, note that under any *TTPDL* scheduler D , the CTMDP $\mathcal{C}$ is totally stochastic and for $s \notin G$, the probability $Pr_{\nu_s, D}^\omega(\Diamond^{[0,z]}G)$ can be computed by:

$$Pr_{\nu_s, D}^\omega(\Diamond^{[0,z]}G) = \int_0^z E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, D(s, t), s') \cdot Pr_{\nu_{s'}, D}^\omega(\Diamond^{[0, z-t]}G) dt.$$

Note D_τ is a *TTPDL* scheduler, thus it holds that

$$Pr_{v_s, D_\tau}^\omega (\diamond^{[0, z]} G) = \int_0^z E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, D_\tau(s, t), s') \cdot Pr_{v_{s'}, D_\tau}^\omega (\diamond^{[0, z-t]} G) dt.$$

This integral can then be split into two parts $A(s, z)$ and $B(s, z)$ at time $t = \tau$: it follows in a similar way as Eq. (5.6) with the difference of taking the action $D_\tau(s, t)$ instead of the maximum over all $\alpha \in \text{Act}$. The lower bound can then be established by induction on k , by adapting the lower bound proof of Eq. (5.19) of Thm. 5.5 appropriately. $\square$

5.4 A case study: The stochastic job scheduling problem

We illustrate the applicability of our approach by considering the stochastic job scheduling problem (sJSP) from [BDF81]. In their paper, the authors analyze the *expected* time to complete a set of stochastic jobs on a number of identical processors under a preemptive scheduling policy. An instance of the sJSP is a tuple (m, n, μ) , where $m \geq 2$ is the number of processors, $J = \{1, \dots, n\}$ is the set of stochastic jobs and $\mu : J \rightarrow \mathbb{R}_{>0}$ specifies the jobs' exponential service times, i.e. $\mu(i)$ is the rate of job i . Each time a job finishes, the preemptive scheduling allows us to assign each processor one of the k remaining jobs, giving rise to $\binom{k}{m}$ nondeterministic choices.

The sJSP can be considered as a locally uniform CTMDP: A state of the sJSP is a tuple (R, W) , where $R, W \subseteq J$ are the sets of running and waiting jobs, respectively. When a job $j \in R$ completes, the decision which jobs to schedule next is nondeterministic.

An action $\alpha \in \text{Act}((R, W))$ is a preemptive schedule: If state (R, W) is left because a job $j \in R$ finishes and if $\alpha : R \rightarrow 2^{R \cup W}$ is chosen, the set $\alpha(j)$ defines the jobs that are executed next. In each state (R, W) , let $\text{Act}((R, W)) = \{\alpha : R \rightarrow 2^{R \cup W} \mid \forall j \in R. j \notin \alpha(j) \wedge |\alpha(j)| \leq m \wedge |\alpha(j)| \text{ maximal}\}$. For $\alpha \in \text{Act}((R, W))$, we define the $\alpha(j)$ -successor (R', W') of (R, W) , denoted $(R, W) \xrightarrow{\alpha(j)} (R', W')$, such that $R' = \alpha(j)$ and $W' = (R \cup W) \setminus (\{j\} \cup \alpha(j))$:

Definition 5.10 (Modelling the sJSP as a CTMDP). Let $P = (m, n, \mu)$ be a sJSP and (R, W) a state. The induced CTMDP $(\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ is defined such that $\mathcal{S} = 2^J \times 2^J$, $\nu = \{(R, W) \mapsto 1\}$, $\text{Act} = \bigcup_{(R', W') \in \mathcal{S}} \text{Act}((R', W'))$ and

$$\mathbf{R}((R, W), \alpha, (R', W')) = \begin{cases} \mu(j) & \text{if } (R, W) \xrightarrow{\alpha(j)} (R', W') \text{ and} \\ 0 & \text{otherwise.} \end{cases}$$

Thus, given state (R, W) , for every job $j \in R$ and action α , there exists an α -transition with the rate $\mu(j)$ of job j that leads to the $\alpha(j)$ -successor (R', W') .

Figure 5.4(a) depicts a fragment of the CTMDP induced by the $(2, 4, \mu)$ sJSP with initial state (R, W) where R is given by the underlined process identifiers (i.e. $R = \{1, 3\}$) and $W = \{2, 4\}$. Action α_1 represents a replacement strategy where jobs $\{3, 4\}$ are executed next if job $1 \in R$ finishes first and otherwise, the next jobs are $\{2, 4\}$. Similarly, for action α_2 , the jobs $\{2, 4\}$ (or $\{1, 4\}$) are scheduled next if job 1 (job 3, resp.) completes first.

The stochastic job scheduling problem is a classical example of a queueing system. At the beginning of this chapter, we claimed that local uniformity is commonly found in this setting. In fact, for our model of the sJSP we can prove local uniformity:

Lemma 5.6 (The sJSP is locally uniform). *For any sJSP $P = (m, n, \mu)$ and all initial states, the CTMDP model induced by Def. 5.10 is locally uniform.*

Proof. From Def. 5.10 it directly follows that for all states (R, W) it holds

$$\begin{aligned} E((R, W), \alpha) &= \sum_{(R', W') \in \mathcal{S}} \mathbf{R}((R, W), \alpha, (R', W')) \\ &= \sum_{(R, W) \xrightarrow{\alpha(j)} (R', W')} \mathbf{R}((R, W), \alpha, (R', W')) = \sum_{j \in R} \mu(j). \end{aligned}$$

Hence, $E((R, W), \alpha) = E((R, W), \beta)$ for all $\alpha, \beta \in \text{Act}((R, W))$. □

Applying the results from Sec. 5.3, we are now able to algorithmically compute the maximum and minimum probabilities to finish all jobs within some time bound z . In Fig. 5.4(b), we plot the maximum and minimum probabilities to finish jobs $\{1, \dots, 4\}$ over a time bound $z \in [0, 15]$ for different values of μ . The probabilities that are shown in Fig. 5.4(b) were obtained by implementing the discretization approach of Sec. 5.3 for maximum and minimum time-bounded reachability. Clearly, for equally distributed job durations, i.e. if $\mu(i) = \mu(k)$ for all i, k , the maximum and minimum probabilities coincide. However, if $\mu(i) \neq \mu(k)$, the probabilities depend on the scheduling policy: In [BDF81], the authors prove that a shortest expected processing time first (SEPT) strategy minimizes the *expected* completion time of the sJSP; reversely, the longest expected processing time strategy (LEPT) is proved to maximize the *expected* completion time.

Although we consider a different quantitative measure (i.e. maximum time-bounded *reachability* instead of *expected* completion time), we observe in our examples, that the ε -optimal τ -scheduler that maximizes the reachability probabilities adheres to the SEPT strategy; moreover, the optimal τ -scheduler for the minimum probabilities obeys the LEPT strategy.

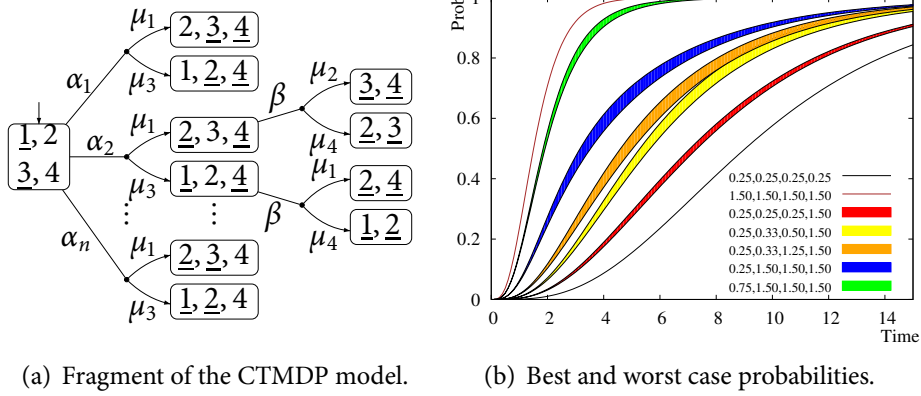


Figure 5.4: Modeling and analysis of the stochastic job scheduling problem.

5.5 Conclusion and related work

In this chapter, we have introduced an efficient discretization algorithm in PTIME that solves the problem of computing time-bounded reachability probabilities in locally uniform CTMDPs with respect to time- and history-dependent late schedulers.

To the best of our knowledge, this is the first time that an automatic analysis of time-bounded reachability objectives becomes feasible for time-dependent schedulers. Moreover, the main advantage of our approach is that we are able to bound the error that is induced by the approximation algorithm in advance. In particular, the maximal admissible error $\varepsilon > 0$ can be specified a priori.

The computation is done by applying the well-known value iteration algorithm [Ber95] to the CTMDP's discretized MDP. We choose the value iteration approach over other methods like LP-solvers, as it has major advantages in our setting: During the value iteration steps, it is possible to extract the optimal scheduling decisions and to synthesize an ε -optimal τ -scheduler whose decisions maximize the reachability objective. Further, the iterative computation allows us to compute time-bounded reachability probabilities incrementally: As a byproduct of the value iteration for a time bound z , we obtain the reachability probabilities for all smaller time bounds $z' < z$ (where z' is a multiple of τ) with minimal computational overhead.

Related work. In the literature, the analysis of CTMDPs has received scant attention. Most of the existing results focus on optimizing criteria such as the expected total reward [GHLPR06, Mil68a] or the expected long-run average reward [dA97, GHLPR06, Mil68b]. Directly related to the results of this chapter is the work in [BHKH05], which provides an algorithm that computes time-bounded reachability probabilities in globally uniform CTMDPs. However, its applicability is severely restricted, as global uniformity — which requires the sojourn times in all states to be identically distributed — is hard to

achieve. We shortly discuss the reason for this:

The approach for the analysis of time-bounded reachability probabilities that is taken in [BHKH05] refers only to time-abstract schedulers, which are strictly less powerful than time-dependent ones [BHKH05, NSK09]. Moreover, as observed in [BHKH05], the uniformization approach that is known from Markov chain theory does not work for CTMDPs and time-abstract scheduler classes: Intuitively, uniformization introduces self loops (or copy states, in case of local uniformization) in the CTMDP model. Thereby uniformization changes the structure of the model. These structural changes expose significant information to history dependent (but time-abstract) schedulers and can be used to estimate the timed behaviour of the system (although the scheduler class is time-abstract). A formal proof of this is included in [BHKH05]. Due to similar reasons, local uniformization fails for all non-trivial time-abstract scheduler classes as proved in Sec. 4.3 (see page 103).

Recently, maximal reachability probabilities in CTMDPs have been studied in stochastic timed games [BF09, BFK⁺09]: However, the authors of [BFK⁺09] also consider the strictly weaker classes of time abstract schedulers, while [BF09] addresses the decidability problem for qualitative reachability probabilities in stochastic timed games, that is, reachability probabilities that are 1 or 0, respectively.

Hence, both approaches differ considerably from our results: The time-dependent scheduler *ML*-schedulers that we use are proved to be strictly more expressive (that is, they generally induce strictly higher probability bounds) than the time-abstract schedulers that are considered in the related work. To the best of our knowledge, no analysis techniques are known for time-dependent scheduler classes.

Therefore, this chapter extends the existing results considerably: We provide an efficient algorithm that computes time-bounded reachability probabilities for the class of time- and history-dependent schedulers up to an a priori given error bound ε . Moreover, we relax the restriction to global uniformity in [BHKH05] and allow different states to have different sojourn time distributions.

6 Model Checking Interactive Markov Chains

It is what I sometimes have called "the separation of concerns", which, even if not perfectly possible, is yet the only available technique for effective ordering of one's thoughts, that I know of.

(Edsger W. Dijkstra)

Interactive Markov chains (IMCs) comprise both nondeterministic choices and exponentially distributed delays. Hence, in the family of stochastic models they are related to CTMDPs. However, subtle differences exist: Whereas CTMDPs closely entangle nondeterminism and stochastic behavior in their transition relation, IMCs strictly separate the two aspects and distinguish between Markovian and interactive transitions.

The different approach taken in IMCs is not surprising, given the fact that IMCs originate in stochastic extensions of classical process algebras. As such, they overcome the absence of hierarchical and compositional facilities in purely stochastic dependability models like CTMCs and SPNs [Mol81, Nat80]. Apart from IMCs, many efforts have been undertaken to vanquish this limitation, including formalism like the stochastic Petri box calculus [MVCR08], Statecharts [BHH⁺09] and in particular, the TIPP [GHR93], PEPA [Hil96] and EMPA [BG98, BG01] process algebras. In this thesis, we focus on IMCs which share most of the other approaches' benefits while preserving a succinct and accurate semantics.

Since IMCs smoothly extend labeled transition systems (LTSs), the model has received attention in academic and in industrial settings [BCH⁺08, CGH⁺08, CHLS09]. In practice however, the theoretical benefits have partly been foiled by the fact that for a long time, the analysis of IMCs was restricted to those instances, where the composed IMC could be transformed into a CTMC.

Beyond these special cases, IMCs also support nondeterminism which arises both implicitly from parallel composition and explicitly by the deliberate use of underspecification in the model [HHK02]. In contrast to CTMC-based models, all of these aspects can neatly be represented in the IMC formalism; therefore, IMCs are strictly more expressive than CTMCs.

The work in [Joh07] is the first approach towards an analysis of nondeterministic IMCs, i.e. of IMCs that cannot be transformed into a CTMC. It relies on a measure preserving transformation from IMCs to CTMDPs and the time-bounded reachability algorithm from [BHKH05]. The latter relies on globally uniform CTMDPs which are obtained by the transformation in [Joh07, BHH⁺09] if the underlying IMC is also globally uniform, that is, if all Markovian states have the same sojourn time distribution.

Apart from these special cases, no analysis techniques exist for the general setting where IMCs are neither globally uniform nor can they be transformed into an equivalent CTMC. In this chapter, we close this gap and provide a model checking algorithm that works for arbitrary IMCs. Our approach extends the discretization technique that is used in Chapter 5: Instead of only considering *time-bounded* reachability objectives, we extend our results to time intervals, that is, we maximize the probability to visit a goal state during a given *time interval*. We then use a fixed-point characterization to discretize an IMC and to obtain an *interactive probabilistic chain* (IPC) [CHLS09]. Our main contribution is the proof that the IPC's maximum step-interval bounded reachability coincides (up to ϵ) with the maximum time-interval bounded reachability probability in the underlying IMC. As a final step, we adapt the value iteration algorithm to IPCs and compute the step-interval bounded reachability probabilities.

On the specification side, the continuous stochastic logic (CSL) [ASSB96, BHHK03] permits to specify a wide variety of performance and dependability measures. It has originally been devised for model checking CTMCs. Therefore, Sec. 6.5 proposes an adaptation of CSL to IMC which enables us to reason about the maximum and minimum achievable probability for CSL path formulas. We then develop an algorithm to automatically model check CSL formulas on arbitrary IMCs.

The crucial point in model checking CSL is the computation of time-interval bounded reachability probabilities. Having achieved the latter, we obtain a model checking algorithm which has a worst-case time complexity of $\mathcal{O}(|\Phi| \cdot (n^{2.376} + (m + n^2) \cdot (\lambda b)^2 / \epsilon))$, where $|\Phi|$ denotes the size of the CSL formula, n, m are the number of states and transitions of the IMC, resp., and b and λ are the maximum upper time interval bound in Φ and the IMC's maximum exit rate, respectively.

As in the previous chapter, we present all results only for maximum time-bounded reachability probabilities. However, all proofs carry over when minimizing the interval-bounded reachability probabilities.

Organization of this chapter. Section 6.1 formally introduces IMCs. In Sec. 6.2 we obtain a fixed-point characterizations for time-interval (and step-interval) bounded reachability in IMCs (respectively in IPCs). A major contribution are the correctness proofs in Sec. 6.3 which provide the theoretical basis for the value iteration algorithm that we present in Sec. 6.4. Section 6.5 introduces the logic CSL and discusses how the interval bounded reachability analysis can be applied to the model checking problem for CSL on IMCs. Finally, we provide some experimental results obtained by our prototypical

implementation in Sec. 6.6.

6.1 Interactive Markov chains

IMCs strictly separate *interactive* from *Markovian* transitions; therefore, they can be seen as a fully orthogonal extension of labeled transition systems with exponentially distributed delays. This enables compositional modeling with intermittent weak bisimulation minimization [Her02] and even allows us to augment existing untimed process algebra specifications with random timing [HK00, BHH⁺09]. Moreover, the IMC formalism is not restricted to exponential delays but permits to encode arbitrary phase-type distributions such as hyper- and hypoexponentials [Pul09]. An excellent and detailed discussion of the advantages of the IMC modeling formalism can be found in the paper [BHK06].

6.1.1 Preliminaries

Opposed to CTMDPs, interactive Markov chains (IMCs) disentangle the relation between Markovian and nondeterministic behaviors: Therefore, IMCs strictly separate *Markovian* from *interactive transitions*. We restate the definition of IMCs from [Her02]:

Definition 6.1 (Interactive Markov chain). An interactive Markov chain is a tuple $\mathcal{M} = (\mathcal{S}, Act, IT, MT, \nu)$ where $\mathcal{S}$ and Act are nonempty sets of states and actions, $IT \subseteq \mathcal{S} \times Act \times \mathcal{S}$ is a set of interactive transitions and $MT \subseteq \mathcal{S} \times \mathbb{R}_{>0} \times \mathcal{S}$ is a set of Markovian transitions. Further, $\nu \in \text{Distr}(\mathcal{S})$ is the initial distribution.

We distinguish *external* actions in Act_e from *internal* actions in Act_i and set $Act = Act_e \cup Act_i$. The reason for this distinction is that IMCs may be composed via synchronization over the set of external actions Act_e , while internal actions in Act_i are not observable from the outside environment. For a detailed discussion of the compositional aspects of IMCs, we refer the reader to [Her02]. For the scope of this thesis, we consider *closed* IMCs [Her02, Joh07], that is, we focus on the IMC $\mathcal{M}$ that is obtained as the final outcome of the composition. Accordingly, $\mathcal{M}$ is not subject to any further synchronization and all remaining external actions can safely be hidden. In our analysis, we therefore assume that $Act_e = \emptyset$ and identify the sets Act and Act_i .

For Markovian transitions, we use λ and μ to denote rates of exponential distributions. Moreover, $IT(s) = \{(s, \alpha, s') \in IT\}$ is the set of interactive transitions that leave state s ; similarly, for Markovian transitions, we set $MT(s) = \{(s, \lambda, s') \in MT\}$. A state $s \in \mathcal{S}$ is *Markovian* iff $MT(s) \neq \emptyset$ and $IT(s) = \emptyset$; it is *interactive* iff $MT(s) = \emptyset$ and $IT(s) \neq \emptyset$. Further, s is a *hybrid state* iff $MT(s) \neq \emptyset$ and $IT(s) \neq \emptyset$; finally, s is a *deadlock state* iff $MT(s) = IT(s) = \emptyset$. We use $MS \subseteq \mathcal{S}$ and $IS \subseteq \mathcal{S}$ to refer to the sets of Markovian and interactive states in $\mathcal{M}$.

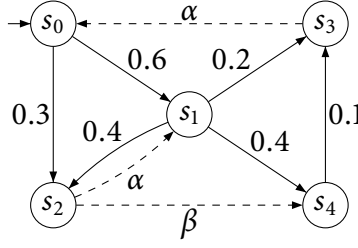


Figure 6.1: Example of an IMC with Markovian and interactive states.

For a Markovian state $s \in MS$, we define $\mathbf{R}(s, s') = \sum \{\lambda \mid (s, \lambda, s') \in MT(s)\}$ as the *rate* to move from state s to state s' and $E(s) = \sum_{s' \in \mathcal{S}} \mathbf{R}(s, s')$ as the *exit rate* of state s ; further, $post^M(s) = \{s' \in \mathcal{S} \mid \mathbf{R}(s, s') > 0\}$ denotes the set of *successor states* of state s . The *discrete branching probability* to move from state s to state s' is $\mathbf{P}(s, s') = \frac{\mathbf{R}(s, s')}{E(s)}$.

Example 6.1. Let $\mathcal{M}$ be the IMC depicted in Fig. 6.1. The semantics of Markovian states equals that of a CTMC state: More precisely, consider the Markovian state s_0 and the transition $(s_0, 0.3, s_2) \in MT(s)$ (depicted by a solid line) that leads from state s_0 to state s_2 with rate $\lambda = 0.3$. The transition's delay is exponentially distributed with rate λ ; hence, it expires in the next $z \in \mathbb{R}_{\geq 0}$ time units with probability $\int_0^z \lambda e^{-\lambda t} dt = (1 - e^{-0.3z})$. As state s_0 has two Markovian transitions, they compete for execution and the IMC moves along the transition whose delay expires first. Clearly, in such a race, the sojourn time in s_0 is determined by the first transition that executes. As the minimum of exponential distributions is exponentially distributed with the sum of their rates, the sojourn time in a state s is determined by the exit rate $E(s)$ of state s . In general, the probability to move from a state $s \in MS$ to a successor state $s' \in \mathcal{S}$ equals the probability that (one of) the Markovian transitions that lead from s to s' wins the race. Accordingly, for state s_0 of our example, we have $\mathbf{R}(s_0, s_2) = 0.3$, $E(s_0) = 0.3 + 0.6 = 0.9$ and $\mathbf{P}(s_0, s_2) = \frac{1}{3}$. $\diamond$

For interactive transitions, we adopt the *maximal progress assumption* [Her02, p. 71] which states that internal transitions (i.e. interactive transitions labeled with internal actions) trigger instantaneously. This implies that they take precedence over all Markovian transitions whose probability to execute immediately is 0. Therefore all Markovian transitions that emanate a hybrid state can be removed without altering the IMC's behavior. This allows us to assume throughout this chapter that $MT(s) \cap IT(s) = \emptyset$ for all $s \in \mathcal{S}$.

To ease the development of the theory, we assume w.l.o.g. that each internal action $\alpha \in Act$ has a unique successor state, denoted $succ(\alpha)$; note that this is no restriction, for if $(s, \alpha, u), (s, \alpha, v) \in IT(s)$ are internal transitions with $u \neq v$, we may replace them by new transitions (s, α_u, u) and (s, α_v, v) with fresh internal actions α_u and α_v .

The *internal successor relation* $\rightsquigarrow_i \subseteq \mathcal{S} \times \mathcal{S}$ is given by $s \rightsquigarrow_i s'$ iff $(s, \alpha, s') \in IT$; furthermore, the *internal reachability relation* $\rightsquigarrow_i^*$ is the reflexive and transitive closure of $\rightsquigarrow_i$. Accordingly, we define $post^i(s) = \{s' \in \mathcal{S} \mid s \rightsquigarrow_i s'\}$ and $Reach^i(s) = \{s' \in \mathcal{S} \mid s \rightsquigarrow_i^* s'\}$.

Finally, entering a deadlock state results in a time lock, as neither internal nor Marko-

vian transitions are available. Therefore, we equip deadlock states $s \in \mathcal{S}$ with interactive self-loops (s, α, s) . Note that the occurrence of time locks breaks compositionality; however, note that our analysis takes place on the closed model which is the monolithic result that is obtained after all compositions.

We justify the modification of deadlock states as follows: Whereas each interactive or Markovian state has an associated sojourn time distribution (which is either 0 or an exponential distribution), the sojourn time in deadlock states remains unquantified. In this case, we encounter a *time lock* situation where the global time does not proceed any further: If a deadlock state is reached at global time t_{dead} , the probability distribution of the associated stochastic process $\{X_t\}_{t \in \mathbb{R}_{\geq 0}}$ is undefined for time-points $t > t_{dead}$. The same phenomenon occurs if a closed IMC eventually remains in a cycle of interactive transitions. In this case, the global time also stops, resulting in a time lock. Hence, the two situations are semantically equivalent which justifies to equip any deadlock state with an interactive self-loop.

Note however, that our approach also allows for a different deadlock state semantics, where the global clock continues; in this case, we would add a Markovian instead of an internal self-loop.

6.1.2 Paths in interactive Markov chains

To unify the notation for interactive and Markovian transitions, we introduce a special action $\perp \notin Act$ and let σ range over $Act_{\perp} = Act \cup \{\perp\}$. In this way, we can denote a finite *path* as a sequence $\pi = s_0 \xrightarrow{\sigma_0, t_0} s_1 \xrightarrow{\sigma_1, t_1} \dots \xrightarrow{\sigma_{n-1}, t_{n-1}} s_n$, where $s_i \in \mathcal{S}$, $\sigma_i \in Act_{\perp}$ and $t_i \in \mathbb{R}_{\geq 0}$ for $i \leq n$. We write $s_i \xrightarrow{\perp, t_i} s_{i+1}$ for Markovian and $s_i \xrightarrow{\alpha_i, 0} s_{i+1}$ for interactive transitions in π . As before, $|\pi|$ denotes the length of path π . Moreover, $\pi[k] = s_k$ and $\delta(\pi, k) = t_k$ refer to the $(k+1)$ -th state on π and its associated sojourn time. Accordingly, $\Delta(\pi, i) = \sum_{k=0}^{i-1} t_k$ is the total time spent on π (where $\Delta(\pi, 0) = 0$) when reaching state $\pi[i]$. If π is finite with $|\pi| = n$, then $\Delta(\pi) = \Delta(\pi, n)$ is the total time spent on π ; similarly, $\pi \downarrow = s_n$ is the last state on π . The *path infix* between the $(i+1)$ -th and the $(j+1)$ -th state of π is denoted $\pi[i..j]$.

Because internal transitions occur immediately in IMCs, an IMC can traverse several states at once. Therefore, we modify the definition of $\pi@t$ such that $\pi@t \in (\mathcal{S}^* \cup \mathcal{S}^\omega)$ denotes the sequence of states that are traversed on π at time point $t \in \mathbb{R}_{\geq 0}$.

The formal derivation of $\pi@t$ is slightly involved: Let i be the smallest index such that $t \leq \Delta(\pi, i)$. Then $\pi[i]$ is the first state on π that is visited at or after time t ; if no such state exists, we set $\pi@t = \langle \rangle$. Otherwise we distinguish two cases: If $t < \Delta(\pi, i)$, we define $\pi@t = \langle s_{i-1} \rangle$; if $t = \Delta(\pi, i)$, let j be the largest index (or $+\infty$, if no such finite index exists) such that $t = \Delta(\pi, j)$ and define $\pi@t = \langle s_i \dots s_j \rangle$.

Example 6.2. Consider the path $\pi = s_0 \xrightarrow{\alpha_0, 0} s_1 \xrightarrow{\alpha_1, 0} s_2 \xrightarrow{\perp, t_2} s_3 \xrightarrow{\alpha_3, 0} s_4 \xrightarrow{\alpha_4, 0} s_5 \xrightarrow{\perp, t_5} s_6$ and let $0 < \varepsilon < \min\{t_2, t_5\}$. The derivations for the sequences $\pi@0$, $\pi@(t_2 - \varepsilon)$, $\pi@t_2$ and $\pi@(t_2 + \varepsilon)$ are sketched in Tab. 6.1:

Intuitively, the $(i+1)$ -th state on path π (i.e. $\pi[i]$) is entered at time $\Delta(\pi, i)$. To find the first state of the sequence $\pi@t$, let i be the first index on π where at least t time units have passed. Formally, we have to choose the minimal i that satisfies $t \leq \Delta(\pi, i)$. For such a minimal i , $t < \Delta(\pi, i)$ implies that time has passed in the previous state $\pi[i-1]$ and that we have been in that state at time point t . Hence, $\pi[i-1]$ must be a Markovian state and we set $\pi@t = \langle \pi[i-1] \rangle$. Otherwise $t = \Delta(\pi, i)$, implying that state $\pi[i]$ is entered at time point t . If it is an interactive state, further transitions can occur immediately. Hence, we look for the maximal index j , for which $\Delta(\pi, j)$ still equals t and define $\pi@t = \langle \pi[i] \dots \pi[j] \rangle$. $\diamond$

We write $s \in \langle s_i \dots s_j \rangle$ if $s \in \{s_i, \dots, s_j\}$; further, for states $s \in \langle s_i \dots s_j \rangle$ we define $\text{Pref}(\langle s_i \dots s_j \rangle, s) = \langle s_i, \dots s_k \rangle$, where $s = s_k$ and k is minimal. If $s \notin \langle s_i \dots s_j \rangle$, we set $\text{Pref}(\langle s_i \dots s_j \rangle, s) = \langle \rangle$. The definitions for *time-abstract* paths are similar.

6.1.3 Events and measurable spaces

A path π (time-abstract path π') as defined in Sec. 6.1.2 is a concatenation of a state and a sequence of *combined transitions* (*time-abstract combined transitions*) from the set $\Omega = \mathbb{R}_{\geq 0} \times \text{Act}_{\perp} \times \mathcal{S}$ ($\Omega_{\text{abs}} = \text{Act}_{\perp} \times \mathcal{S}$); hence, $\pi = s_0 \circ m_0 \circ m_1 \circ \dots \circ m_{n-1}$ with $m_i = (t_i, \sigma_i, s_{i+1}) \in \Omega$ ($m_i = (\sigma_i, s_{i+1}) \in \Omega_{\text{abs}}$). Thus $\text{Paths}^n(\mathcal{M}) = \mathcal{S} \times \Omega^n$ is the set of paths of length n in an IMC $\mathcal{M}$; further, $\text{Paths}^*(\mathcal{M})$, $\text{Paths}^\omega(\mathcal{M})$ and $\text{Paths}(\mathcal{M})$ are the sets of finite, infinite and all paths in $\mathcal{M}$. To refer to time-abstract paths, we add the subscript *abs*; further the reference to $\mathcal{M}$ is omitted wherever possible. The measure-theoretic concepts are mentioned only briefly, as they directly carry over from the definitions for the CTMDP case (cf. Sec. 3.3.2 on page 76): Events in $\mathcal{M}$ are measurable sets of paths; as paths are Cartesian products of combined transitions, we define the σ -field $\mathfrak{F} = \sigma(\mathfrak{B}(\mathbb{R}_{\geq 0}) \times \mathfrak{F}_{\text{Act}_{\perp}} \times \mathfrak{F}_{\mathcal{S}})$ on subsets of Ω where $\mathfrak{F}_{\mathcal{S}} = 2^{\mathcal{S}}$ and $\mathfrak{F}_{\text{Act}_{\perp}} = 2^{\text{Act}_{\perp}}$.

The product σ -field $\mathfrak{F}_{\text{Paths}^n}$ of measurable subsets of Paths^n is defined as usual, that is, $\mathfrak{F}_{\text{Paths}^n} = \sigma(\{S_0 \times M_1 \times \dots \times M_n \mid S_0 \in \mathfrak{F}_{\mathcal{S}}, M_i \in \mathfrak{F}\})$. As for CTMDPs, the cylinder-set construction [ADD00] extends this to infinite paths: A set $B \in \mathfrak{F}_{\text{Paths}^n}$ is called a *base* of an infinite *cylinder* C where $C = \text{Cyl}(B) = \{\pi \in \text{Paths}^\omega \mid \pi[0..n] \in B\}$. Finally, the cylinders generate the σ -field $\mathfrak{F}_{\text{Paths}^\omega} = \sigma(\bigcup_{n=0}^{\infty} \{\text{Cyl}(B) \mid B \in \mathfrak{F}_{\text{Paths}^n}\})$.

$t \leq \Delta(\pi, i)$	0	1	2	3	4	5	6	min i	max j	$\pi@t$
0	✓	✓	✓	✓	✓	✓	✓	0	2	$\langle s_0 s_1 s_2 \rangle$
$t_2 - \varepsilon$	×	×	×	✓	✓	✓	✓	3	NA	$\langle s_2 \rangle$
t_2	×	×	×	✓	✓	✓	✓	3	5	$\langle s_3 s_4 s_5 \rangle$
$t_2 + \varepsilon$	×	×	×	×	×	×	✓	6	NA	$\langle s_5 \rangle$

Table 6.1: An example for the derivation of $\pi@t$ for interactive Markov chains.

6.1.4 Resolving nondeterminism by schedulers

An IMC $\mathcal{M}$ is *nondeterministic* iff for some $s \in IS$, there exist interactive transitions $(s, \alpha, u), (s, \beta, v) \in IT(s)$ with $u \neq v$: For example, nondeterminism arises in the IMC in Fig. 6.1: In state s_2 , two internal transitions (with actions α and β) lead to states s_1 and s_4 , respectively. By the maximal progress assumption, they both execute instantaneously at time point 0. Hence, no order of execution can be fixed, which leads to the situation that the successor state of state s_2 (either s_1 or s_4) is not uniquely determined. To resolve this nondeterministic choice, we use *schedulers*: If $\mathcal{M}$ reaches state s_2 along a *history* $\pi \in Paths^*$, a scheduler yields a probability distribution over the set $Act(\pi \downarrow) = \{\alpha, \beta\}$. Formally, we define the set of *enabled actions* in an interactive state $s \in IS$ of an IMC as follows:

$$Act(s) = \{\alpha \in Act \mid \exists s' \in S. (s, \alpha, s') \in IT\}.$$

IMC schedulers are closely related to CTMDP schedulers and most of the concepts from Sec. 3.3.2 and Chapters 4 and 5 apply analogously. The only notable difference is the distinction between interactive and Markovian states: Nondeterminism does not occur in the latter, as the successor states are probabilistically quantified. Hence, the only source of nondeterminism are competing internal transitions in interactive states.

Definition 6.2 (Generic measurable scheduler). A generic scheduler on an IMC $\mathcal{M} = (S, Act, IT, MT, \nu)$ is a partial mapping $D : Paths^* \times \mathfrak{F}_{Act} \rightarrow [0, 1]$ such that $D(\pi, \cdot) \in \text{Distr}(Act(\pi \downarrow))$ for all $\pi \in Paths^*$ with $\pi \downarrow \in IS$. A generic scheduler D is *measurable* (that is, a *GM scheduler*) iff for all $A \in \mathfrak{F}_{Act}$, $D^{-1}(A) : Paths^* \rightarrow [0, 1]$ is measurable.

Measurability states that $\{\pi \mid D(\pi, A) \in B\} \in \mathfrak{F}_{Paths^*}$ holds for all $A \in \mathfrak{F}_{Act}$ and $B \in \mathfrak{B}([0, 1])$; intuitively, it excludes schedulers which resolve the nondeterminism in a way that induces non-measurable sets. Recall that no nondeterminism occurs if $\pi \downarrow \in MS$. However, we slightly abuse notation and assume that $D(\pi, \cdot) = \{\perp \mapsto 1\}$ if $\pi \downarrow \in MS$ so that D yields a distribution over $Act_\perp$. In this way, we can treat a GM-scheduler D as a total function $D : Paths^* \times \mathfrak{F}_{Act_\perp} \rightarrow [0, 1]$.

A GM scheduler D is *deterministic* iff $D(\pi, \cdot)$ is degenerate for all $\pi \in Paths^*$. We use *GM* (and *GMD*) to denote the class of generic measurable (deterministic) schedulers. Further, a GM scheduler D_{abs} is *time-abstract* (GM_{abs}) iff $abs(\pi) = abs(\pi')$ implies $D_{abs}(\pi, \cdot) = D_{abs}(\pi', \cdot)$.

Example 6.3. If state s_2 in Fig. 6.1 is reached along path $\pi = s_0 \xrightarrow{0.4, \perp} s_2$, then $D(\pi)$ might yield the distribution $\{\alpha \mapsto \frac{1}{2}, \beta \mapsto \frac{1}{2}\}$, whereas for history $\pi' = s_0 \xrightarrow{1.5, \perp} s_2$, it might return a different distribution, say $D(\pi) = \{\alpha \mapsto 1\}$. $\diamond$

6.1.5 Probability measures for IMCs

In this section, we define the probability measure [Joh07] induced by D on the measurable space $(Paths^\omega, \mathfrak{F}_{Paths^\omega})$. We first derive the probability of measurable sets of combined transitions, i.e. of subsets of Ω :

Definition 6.3 (Probability of combined transitions). Let $\mathcal{M} = (\mathcal{S}, Act, IT, MT, \nu)$ be an IMC and $D \in GM$. For $\pi \in Paths^*$, we define the probability measure $\mu_D(\pi, \cdot) : \mathfrak{F} \rightarrow [0, 1]$:

$$\mu_D(\pi, M) = \begin{cases} \sum_{\alpha \in Act(\pi \downarrow)} \mathbf{I}_M(\alpha, 0, succ(\alpha)) \cdot D(\pi, \{\alpha\}) & \text{if } \pi \downarrow \in IS \\ \int_{\mathbb{R}_{\geq 0}} E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{I}_M(\perp, t, s') \cdot \mathbf{P}(s, s') \, dt & \text{if } \pi \downarrow \in MS. \end{cases} \quad (6.1)$$

As usual, $\mathbf{I}_M$ denotes the indicator function for the set M . Intuitively, $\mu_D(\pi, M)$ is the probability to continue along one of the combined transition in the set M . For an interactive state $s \in IS$, it is the probability of choosing $\alpha \in Act(\pi \downarrow)$ such that $(\alpha, 0, succ(\alpha))$ is a transition in M . Stated differently, we sum up the probabilities of all combined transitions in M that lead immediately with an interactive transition to a successor state of $\pi \downarrow$. If $s \in MS$, $\mu_D(\pi, M)$ is given by the density for the Markovian transition to trigger at time t and the probability that the IMC moves to a successor state s' according to a combined transition in M . As paths are inductively defined using combined transitions, we can lift the probability measure $\mu_D(\pi, \cdot)$ to $\mathfrak{F}_{Paths^n}$ as usual:

Definition 6.4 (Probability measure). Let $\mathcal{M} = (\mathcal{S}, Act, IT, MT, \nu)$ be an IMC and $D \in GM$. For $n \geq 0$, we define the probability measures $Pr_{\nu, D}^n$ inductively on the measurable space $(Paths^n, \mathfrak{F}_{Paths^n})$:

$$\begin{aligned} Pr_{\nu, D}^0 : \mathfrak{F}_{Paths^0} &\rightarrow [0, 1] : \Pi \mapsto \sum_{s \in \Pi} \nu(s) \quad \text{and} \\ Pr_{\nu, D}^{n+1} : \mathfrak{F}_{Paths^{n+1}} &\rightarrow [0, 1] : \Pi \mapsto \int_{Paths^n} Pr_{\nu, D}^n(d\pi) \int_{\Omega} \mathbf{I}_{\Pi}(\pi \circ m) \mu_D(\pi, dm). \end{aligned}$$

6.1.6 Interactive probabilistic chains

In this section, we introduce *interactive probabilistic chains* (IPCs) [CHLS09] which serve as the discrete-time analogon of IMCs. In an IPC, Markovian transitions are replaced by *probabilistic transitions*. As a consequence, no delay time distribution is associated with probabilistic states. Therefore, taking a probabilistic transitions corresponds to a discrete time step in the IPC.

The semantics of interactive transitions remains the same as in the IMC case. Open IMCs can synchronize over the set of *external actions*, whereas *internal actions* are unobservable for the environment.

Definition 6.5 (Interactive probabilistic chain). An interactive probabilistic chain (IPC) is a tuple $\mathcal{P} = (\mathcal{S}, \text{Act}, IT, PT, \nu)$, where $\mathcal{S}, \text{Act}, IT$ and ν are as in Def. 6.1 and $PT : \mathcal{S} \times \mathcal{S} \rightarrow [0, 1]$ is a transition probability function s.t. $\forall s \in \mathcal{S}. PT(s, \mathcal{S}) \in \{0, 1\}$.

A state s in an IPC $\mathcal{P}$ is *probabilistic* iff $\sum_{s' \in \mathcal{S}} PT(s, s') = 1$ and $IT(s) = \emptyset$; PS denotes the set of all probabilistic states. The sets of interactive, hybrid and deadlock states are defined as for IMCs, with the same assumption imposed on deadlock states. Further, we assume any IPC to be closed, that is $(s, \alpha, s') \in IT$ implies $\alpha \in \text{Act}_i$. Hence, $\text{Act}_e = \emptyset$ and we identify the sets Act_i and Act .

As for IMCs, we adopt the *maximal progress assumption* [Her02, p. 71]; hence, internal transitions take precedence over probabilistic transitions and their execution takes 0 discrete time steps. In this way, we obtain a full correspondence between IMCs and IPCs, as in both cases internal transitions consume no time.

Definition 6.6 (IPC scheduler). Let $\mathcal{P} = (\mathcal{S}, \text{Act}, IT, PT, \nu)$ be an IPC. A partial function $D : \text{Paths}_{abs}^* \rightsquigarrow \text{Distr}(\text{Act})$ with $D(\pi) \in \text{Distr}(\text{Act}(\pi \downarrow))$ is a time-abstract history-dependent randomized (GM_{abs}) scheduler.

Note that in the discrete-time setting, measurability issues do not arise. Moreover, we extend $D \in GM_{abs}$ to a complete function $D : \text{Paths}_{abs}^* \rightarrow \text{Distr}(\text{Act}_\perp)$ and assume that $D(\pi) = \{\perp \mapsto 1\}$ iff $\pi \downarrow \in PS$. To define a probability measure on sets of paths in $\mathcal{P}$, we define the probability of a single transition:

Definition 6.7 (Combined transitions in IPCs). Let $\mathcal{P} = (\mathcal{S}, \text{Act}, IT, PT, \nu)$ be an IPC, $s \in \mathcal{S}$, $\sigma \in \text{Act}_\perp$, $\pi \in \text{Paths}_{abs}^*$ and $(\sigma, s) \in \Omega_{abs}$ a time-abstract combined transition. For scheduler $D \in GM_{abs}$, we define

$$\mu_D^{abs}(\pi, \{(\sigma, s)\}) = \begin{cases} PT(\pi \downarrow, s) & \text{if } \pi \downarrow \in PS \wedge \sigma = \perp \\ D(\pi, \{\sigma\}) & \text{if } \pi \downarrow \in IS \wedge \text{succ}(\sigma) = s \\ 0 & \text{otherwise.} \end{cases}$$

For a set of combined transitions $M \subseteq \Omega_{abs}$, we set $\mu_D^{abs}(\pi, M) = \sum_{(\sigma, s) \in M} \mu_D^{abs}(\pi, \{(\sigma, s)\})$.



Figure 6.2: An example for an IMC and its embedded IPC.

The measures μ_D^{abs} extend to a unique measure on sets of paths in $\mathcal{P}$ in the same way as it was shown for the IMC case in Sec. 6.1.5.

Example 6.4. *Each IMC induces an embedded IPC: Consider the IMC $\mathcal{M}$ in Fig. 6.2(a), with initial state s_0 and interactive states s_1 and s_3 . A scheduler D has to resolve the nondeterminism in state s_1 : If $\pi = s_0 \xrightarrow{\perp, t_0} s_0 \xrightarrow{\perp, t_1} s_1$ is the path that led into state s_1 , then $D(\pi)(\alpha)$ is the probability that α is chosen in s_1 . In Fig. 6.2(b), we depict the embedded IPC $emb(\mathcal{M})$ of $\mathcal{M}$: It is obtained by disregarding $\mathcal{M}$'s timed behavior and considering the IMC's discrete branching probabilities $\mathbf{P}(s, s')$ only. Hence $emb(\mathcal{M})$ is the IPC $(\mathcal{S}, Act, PT, IT, \nu)$, where $PT(s, s') = \frac{\mathbf{R}(s, s')}{E(s)}$ if $s \in MS$ and $PT(s, s') = 0$, otherwise. $\diamond$*

6.2 Interval bounded reachability probability

We discuss how to compute the maximum probability to visit a set $G \subseteq \mathcal{S}$ of goal states during a given time interval I . Therefore, let $\mathcal{I}$ be the set of nonempty intervals over the nonnegative reals and let $\mathcal{Q}$ be the set of nonempty intervals with nonnegative rational bounds. For $t \in \mathbb{R}_{\geq 0}$ and $I \in \mathcal{I}$, we define $I \ominus t = \{x - t \mid x \in I \wedge x \geq t\}$ and $I \oplus t = \{x + t \mid x \in I\}$. Obviously, if $I \in \mathcal{Q}$ and $t \in \mathbb{Q}_{\geq 0}$, this implies $I \ominus t \in \mathcal{Q}$ and $I \oplus t \in \mathcal{Q}$.

6.2.1 A fixed point characterization for IMCs

Let $\mathcal{M}$ be an IMC. For a time interval $I \in \mathcal{I}$ and a set $G \subseteq \mathcal{S}$ of goal states, we define the event $\diamond^I G = \{\pi \in Paths^\omega \mid \exists t \in I. \exists s' \in \pi @ t. s' \in G\}$ as the set of all paths that hit a state in G during time interval I . The maximum probability induced by $\diamond^I G$ in $\mathcal{M}$ is denoted $p_{max}^\mathcal{M}(s, I)$. Formally, it is obtained by the supremum under all GM schedulers:

$$p_{max}^\mathcal{M}(s, I) = \sup_{D \in GM} Pr_{\nu_s, D}^\omega(\diamond^I G). \quad (6.2)$$

For a scheduler $D \in GM$, $s \in \mathcal{S}$ and interval $I \in \mathcal{I}$ with $\inf I = a$ and $\sup I = b$, consider the functions $Pr_{\nu_s, D}^\omega(\diamond^{I \ominus [\cdot]} G) : t \mapsto Pr_{\nu_s, D}^\omega(\diamond^{I \ominus t} G)$. Then $Pr_{\nu_s, D}^\omega(\diamond^{I \ominus [\cdot]} G)$ is piecewise continuous in $\mathbb{R}_{\geq 0}$ by definition. As the following lemma proves, continuity (and thereby measurability) extends to $p_{max}^\mathcal{M}(s, I \ominus [\cdot])$:

Lemma 6.1 (Continuity of $p_{\max}^{\mathcal{M}}$). *Let $\mathcal{M} = (\mathcal{S}, \text{Act}, IT, MT, \nu)$ be an IMC, $G \subseteq \mathcal{S}$ a set of goal states and $I \in \mathcal{I}$ an interval. The functions $p_{\max}^{\mathcal{M}}(s, I \ominus [\cdot]) : \mathbb{R}_{\geq 0} \rightarrow [0, 1] : t \mapsto p_{\max}^{\mathcal{M}}(s, I \ominus t)$ are piecewise continuous and measurable for all $s \in \mathcal{S}$.*

Proof. For continuity, we prove that for all $s \in \mathcal{S}$ and $t \in (\mathbb{R}_{>0} \setminus \inf I)$ it holds that

$$\lim_{\delta \rightarrow 0^+} p_{\max}^{\mathcal{M}}(s, I \ominus (t - \delta)) = p_{\max}^{\mathcal{M}}(s, I \ominus t) = \lim_{\delta \rightarrow 0^+} p_{\max}^{\mathcal{M}}(s, I \ominus (t + \delta)). \quad (6.3)$$

Observe that $t = 0$ and $t = \inf I$ are the only discontinuities of $Pr_{\nu_s, D}^{\omega}(\Diamond^{I \ominus t} G)$: To see this, note that $0 \notin I \ominus t$ for $t < \inf I$ and $0 \in I \ominus t$ for $t > \inf I$. Hence, if $t = \inf I$, interactive transitions may reach a goal state directly without requiring integration over the time domain.

Further, observe that $Pr_{\nu_s, D}^{\omega}(s, I \ominus t') \leq p_{\max}^{\mathcal{M}}(s, I \ominus t')$ for all $t' \in \mathbb{R}_{\geq 0}$ by definition of $p_{\max}^{\mathcal{M}}$. To prove that $p_{\max}^{\mathcal{M}}(s, I \ominus [\cdot])$ is piecewise continuous, we proceed by contraposition and assume there exists $t \in (\mathbb{R}_{>0} \setminus \inf I)$ such that Eq. (6.3) is violated: Here we consider left-continuity and distinguish two cases: Assume that $p_{\max}^{\mathcal{M}}(s, I \ominus [\cdot])$ is not continuous from the left at point $t \in \mathbb{R}_{\geq 0}$ and that there exists $\varepsilon > 0$ such that

$$\lim_{\delta \rightarrow 0^+} p_{\max}^{\mathcal{M}}(s, I \ominus (t - \delta)) = p_{\max}^{\mathcal{M}}(s, I \ominus t) - \varepsilon. \quad (6.4)$$

Now, choose $D \in GM$ such that $p_{\max}^{\mathcal{M}}(s, I \ominus t) - Pr_{\nu_s, D}^{\omega}(\Diamond^{I \ominus t} G) = \xi$ for some $\xi \leq \frac{\varepsilon}{2}$. Then

$$\begin{aligned} p_{\max}^{\mathcal{M}}(s, I \ominus t) - \xi &= Pr_{\nu_s, D}^{\omega}(\Diamond^{I \ominus t} G) = \lim_{\delta \rightarrow 0^+} Pr_{\nu_s, D}^{\omega}(\Diamond^{I \ominus (t - \delta)} G) \\ &\leq \lim_{\delta \rightarrow 0^+} p_{\max}^{\mathcal{M}}(s, I \ominus (t - \delta)). \end{aligned}$$

But then, $\lim_{\delta \rightarrow 0^+} p_{\max}^{\mathcal{M}}(s, I \ominus (t - \delta)) \geq p_{\max}^{\mathcal{M}}(s, I \ominus t) - \xi > p_{\max}^{\mathcal{M}}(s, I \ominus t) - \varepsilon$, contradicting Eq. (6.4). For the second case, assume that left-continuity at t is violated because there exists $\varepsilon > 0$ such that

$$\lim_{\delta \rightarrow 0^+} p_{\max}^{\mathcal{M}}(s, I \ominus (t - \delta)) = p_{\max}^{\mathcal{M}}(s, I \ominus t) + \varepsilon. \quad (6.5)$$

Choose $D \in GM$ such that $\lim_{\delta \rightarrow 0^+} Pr_{\nu_s, D}^{\omega}(\Diamond^{I \ominus (t - \delta)} G) = \lim_{\delta \rightarrow 0^+} p_{\max}^{\mathcal{M}}(s, I \ominus (t - \delta)) - \xi$ for some $\xi \leq \frac{\varepsilon}{2}$. Then

$$\begin{aligned} p_{\max}^{\mathcal{M}}(s, I \ominus t) &\geq Pr_{\nu_s, D}^{\omega}(\Diamond^{I \ominus t} G) = \lim_{\delta \rightarrow 0^+} Pr_{\nu_s, D}^{\omega}(\Diamond^{I \ominus (t - \delta)} G) \\ &= \lim_{\delta \rightarrow 0^+} p_{\max}^{\mathcal{M}}(s, I \ominus (t - \delta)) - \xi. \end{aligned}$$

But then, $\lim_{\delta \rightarrow 0^+} p_{\max}^{\mathcal{M}}(s, I \ominus (t - \delta)) \leq p_{\max}^{\mathcal{M}}(s, I \ominus t) + \xi < p_{\max}^{\mathcal{M}}(s, I \ominus t) + \varepsilon$, contradicting Eq. (6.5). Thus, $p_{\max}^{\mathcal{M}}(s, I \ominus [\cdot])$ is piecewise left-continuous. The fact that it is piecewise right-continuous follows along the same lines. Hence, $p_{\max}^{\mathcal{M}}(s, I \ominus [\cdot])$ is piecewise continuous. As piecewise continuous functions are Borel measurable [Ros00, Prop. 3.1.8], we are done. $\square$

Based on the measurability of $p_{\max}^{\mathcal{M}}(s, I \ominus [\cdot])$, we are now ready to derive a fixed point characterization of the maximum probability $p_{\max}^{\mathcal{M}}(s, I)$. More specifically, we prove that $p_{\max}^{\mathcal{M}}$ is the least fixed-point of a higher-order operator Ω :

Theorem 6.1 (Fixed point characterization for IMCs). *Let $\mathcal{M} = (S, \text{Act}, IT, PT, \nu)$ be an IMC, $G \subseteq S$ a set of goal states and $I \in \mathcal{I}$ a time interval with $\inf I = a$ and $\sup I = b$ for some $a, b \in \mathbb{R}_{\geq 0}$. The function $p_{\max}^{\mathcal{M}} : S \times \mathcal{I} \rightarrow [0, 1]$ is the least fixed point of the higher-order operator $\Omega : (S \times \mathcal{I} \rightarrow [0, 1]) \rightarrow (S \times \mathcal{I} \rightarrow [0, 1])$, which is defined as follows:*

1. For Markovian states $s \in MS$:

$$\Omega(F)(s, I) = \begin{cases} \int_0^b E(s) e^{-E(s)t} \cdot \sum_{s' \in S} \mathbf{P}(s, s') \cdot F(s', I \ominus t) dt & \text{if } s \notin G \\ e^{-E(s)a} + \int_0^a E(s) e^{-E(s)t} \cdot \sum_{s' \in S} \mathbf{P}(s, s') \cdot F(s', I \ominus t) dt & \text{if } s \in G. \end{cases}$$

2. For interactive states $s \in IS$:

$$\Omega(F)(s, I) = \begin{cases} 1 & \text{if } s \in G \text{ and } 0 \in I, \\ \max\{F(s', I) \mid s' \in \text{post}^i(s)\} & \text{otherwise.} \end{cases}$$

Proof. The proof is split in two parts: First, we prove that $p_{\max}^{\mathcal{M}}$ is a fixed point of Ω and second, we show that it is the least fixed point.

Recall that in Eq. (6.2) we defined $p_{\max}^{\mathcal{M}}(s, I) = \sup_{D \in GM} Pr_{\nu_s, D}^{\omega}(\diamond^I G)$. To prove that $p_{\max}^{\mathcal{M}}$ is a fixed point of Ω , we first provide a disjoint decomposition of the event $\diamond^I G$: Let $\gamma(\pi, n)$ be the time interval which is spent in the n -th state of path π , measured in absolute time. Formally, $\gamma(\pi, n) = [\Delta(\pi, n), \Delta(\pi, n+1))$ if $\Delta(\pi, n) < \Delta(\pi, n+1)$ and $\gamma(\pi, n) = \{\Delta(\pi, n)\}$, otherwise. Now define the set $\Gamma(I, n)$ of all paths whose $(n+1)$ -th state is in G and lies within time interval I , that is, $\Gamma(I, n) = \{\pi \in \text{Paths}^{\omega} \mid \pi[n] \in G \wedge \gamma(\pi, n) \cap I \neq \emptyset\}$. To achieve a disjoint decomposition of $\diamond^I G$, set $\Pi(I, n) = \Gamma(I, n) \setminus \bigcup_{k=0}^{n-1} \Gamma(I, k)$. Then $\diamond^I G = \bigcup_{n=0}^{\infty} \Pi(I, n)$. For $D \in GM$ it holds:

$$Pr_{\nu, D}^{\omega}(\diamond^I G) = Pr_{\nu, D}^{\omega}\left(\bigcup_{n=0}^{\infty} \Pi(I, n)\right) = \sum_{n=0}^{\infty} Pr_{\nu, D}^{\omega}(\Pi(I, n)).$$

Further, let $p_{\max}^{\mathcal{M}, n}(s, I) = \sup_{D \in GM} Pr_{\nu_s, D}^{\omega}(\bigcup_{i=0}^n \Pi(I, i))$ be the upper bound on the probability to visit G during time interval I and within at most n transitions. First, we show that $p_{\max}^{\mathcal{M}, n+1}(s, I) = \Omega(p_{\max}^{\mathcal{M}, n})(s, I)$. It suffices to consider two cases:

1. Let $s \in MS$ and assume that $s \notin G$ (the case $s \in G$ follows similarly). Then:

$$\Omega(p_{\max}^{\mathcal{M}, n})(s, I) = \int_0^b E(s) e^{-E(s)t} \cdot \sum_{s' \in S} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}, n}(s', I \ominus t) dt$$

$$= \int_0^b E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot \sup_{D \in GM} Pr_{v_{s'}, D}^\omega \left(\bigcup_{i=0}^n \Pi(I \ominus t, i) \right) dt. \quad (6.6)$$

Let $D \in GM$, $s \in \mathcal{S}$, $\sigma \in Act_\perp$ and $t \in \mathbb{R}_{\geq 0}$. We define the GM scheduler $D_{s, \sigma, t}$ such that $D_{s, \sigma, t}(\pi) = D(s \xrightarrow{\sigma, t} \pi)$ for all $\pi \in Paths^*$. Hence, $D_{s, \sigma, t}$ yields the same decisions for history π as the original scheduler D does for the history $s \xrightarrow{\sigma, t} \pi$, where we define $s \xrightarrow{\sigma, t} \pi = s \xrightarrow{\sigma, t} s_0 \xrightarrow{\sigma_0, t_0} s_1 \xrightarrow{\sigma_1, t_1} \dots$ if $\pi = s_0 \xrightarrow{\sigma_0, t_0} s_1 \xrightarrow{\sigma_1, t_1} \dots$. This shift allows us to rewrite $\Omega(p_{max}^{\mathcal{M}, n})(s, I)$ further:

$$\begin{aligned} \Omega(p_{max}^{\mathcal{M}, n})(s, I) &= \sup_{D \in GM} \int_0^b E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot Pr_{v_{s'}, D_{s, \perp, t}}^\omega \left(\bigcup_{i=0}^n \Pi(I \ominus t, i) \right) dt \\ &= \sup_{D \in GM} Pr_{v_s, D}^\omega \left(\bigcup_{i=0}^{n+1} \Pi(I, i) \right) = p_{max}^{\mathcal{M}, n+1}(s, I). \end{aligned}$$

2. Now let $s \in IS$. If $s \in G$ and $0 \in I$, it holds that $\Omega(p_{max}^{\mathcal{M}, n})(s, I) = 1 = p_{max}^{\mathcal{M}, n+1}(s, I)$ and we are finished. Otherwise

$$\begin{aligned} \Omega(p_{max}^{\mathcal{M}, n})(s, I) &= \max_{s' \in post^i(s)} p_{max}^{\mathcal{M}, n}(s', I) = \max_{s' \in post^i(s)} \left(\sup_{D \in GM} Pr_{v_{s'}, D}^\omega \left(\bigcup_{i=0}^n \Pi(I, i) \right) \right) \\ &= \max_{\alpha \in Act(s)} \left(\sup_{D \in GM} Pr_{v_{succ(\alpha)}, D}^\omega \left(\bigcup_{i=0}^n \Pi(I, i) \right) \right) \\ &= \sup_{D \in GM} \max_{\alpha \in Act(s)} Pr_{v_{succ(\alpha)}, D_{s, \alpha, 0}}^\omega \left(\bigcup_{i=0}^n \Pi(I, i) \right) \\ &= \sup_{D \in GM} Pr_{v_s, D}^\omega \left(\bigcup_{i=0}^{n+1} \Pi(I, i) \right) = p_{max}^{n+1}(s, I). \end{aligned}$$

It is easy to see that $p_{max}^{\mathcal{M}, n}(s, I)$ converges to $p_{max}^{\mathcal{M}}(s, I)$: By definition, $\bigcup_{i=0}^n \Pi(I, i) \rightarrow \diamond^I G$ for $n \rightarrow +\infty$. Further, Lemma 2.2(a) implies that for each $D \in GM$ we have that $\lim_{n \rightarrow \infty} Pr_{v_s, D}^\omega \left(\bigcup_{i=0}^n \Pi(I, i) \right) = Pr_{v_s, D}^\omega (\diamond^I G)$. As this applies to all $D \in GM$, it holds $\sup \{ Pr_{v_s, D}^\omega \left(\bigcup_{i=0}^n \Pi(I, i) \right) \mid D \in GM \} \rightarrow \sup \{ Pr_{v_s, D}^\omega (\diamond^I G) \mid D \in GM \}$ for $n \rightarrow +\infty$. Taking the limit on both sides of the equation $\Omega(p_{max}^{\mathcal{M}, n})(s, I) = p_{max}^{\mathcal{M}, n+1}(s, I)$ yields that $\Omega(p_{max}^{\mathcal{M}})(s, I) = p_{max}^{\mathcal{M}}(s, I)$. Hence $p_{max}^{\mathcal{M}}$ is a fixed point of Ω .

It remains to show that $p_{max}^{\mathcal{M}}$ is the least fixed point of Ω . Therefore, let $F : \mathcal{S} \times \mathcal{I} \rightarrow [0, 1]$ be another fixed point of Ω . By induction on the number of (interactive and Markovian) transitions n , we show that $p_{max}^{\mathcal{M}, n}(s, I) \leq F(s, I)$ for all $n \in \mathbb{N}$.

1. In the induction base, it holds that $p_{max}^{\mathcal{M}, 0}(s, I) = 1 = \Omega(F(s, I)) = F(s, I)$ if $s \in G$ and $a = 0$; otherwise $p_{max}^{\mathcal{M}, 0}(s, I) = 0 \leq F(s, I)$.
2. For the induction step, we distinguish between Markovian and interactive states:

- (a) Let $s \in MS$ and $s \notin G$ (the case $s \in G$ can be shown similarly). Then

$$\begin{aligned}
 p_{\max}^{\mathcal{M},n+1}(s, I) &= \Omega(p_{\max}^{\mathcal{M},n})(s, I) \\
 &= \int_0^b E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M},n}(s', I \ominus t) dt \\
 &\leq \int_0^b E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot F(s', I \ominus t) dt \quad (* \text{ ind. hyp. } *) \\
 &= \Omega(F(s, I)) = F(s, I). \quad (* F \text{ is fixed point } *)
 \end{aligned}$$

- (b) Now let $s \in IS$. If $s \in G$ and $0 \in I$, we have $\Omega(F)(s, I) = F(s, I) = 1 \geq p_{\max}^{\mathcal{M},n+1}(s, I)$. Otherwise, the induction hypothesis yields

$$p_{\max}^{\mathcal{M},n+1}(s, I) = \Omega(p_{\max}^{\mathcal{M},n})(s, I) = \max_{s' \in \text{post}^i(s)} p_{\max}^{\mathcal{M},n}(s', I) \leq \max_{s' \in \text{post}^i(s)} F(s', I).$$

By definition of Ω , we have $\max_{s' \in \text{post}^i(s)} F(s', I) = \Omega(F)(s, I) = F(s, I)$, proving that $p_{\max}^{\mathcal{M},n+1}(s, I) \leq F(s, I)$.

Hence, $F(s, I) \geq \lim_{n \rightarrow \infty} p_{\max}^{\mathcal{M},n}(s, I) = p_{\max}^{\mathcal{M}}(s, I)$ and the claim follows. $\square$

Example 6.5. The fixed point characterization suggests to compute $p_{\max}^{\mathcal{M}}(s, I)$ analytically: Consider the IMC $\mathcal{M}$ depicted in Fig. 6.1 and assume that $G = \{s_3\}$. For $I = [0, b]$, $b > 0$ we have $p_{\max}^{\mathcal{M}}(s_3, I) = 1$ and $p_{\max}^{\mathcal{M}}(s_4, I) = 1 - e^{-0.1b}$. For state s_1 , we derive that $p_{\max}^{\mathcal{M}}(s_1, I) = \int_0^b e^{-t} \left(\frac{2}{5} \cdot p_{\max}^{\mathcal{M}}(s_2, I \ominus t) + \frac{1}{5} \cdot p_{\max}^{\mathcal{M}}(s_3, I \ominus t) + \frac{2}{5} \cdot p_{\max}^{\mathcal{M}}(s_4, I \ominus t) \right) dt$. In interactive state s_2 , we obtain that $p_{\max}^{\mathcal{M}}(s_2, I) = \max \{ p_{\max}^{\mathcal{M}}(s_4, I), p_{\max}^{\mathcal{M}}(s_1, I) \}$, which yields that $p_{\max}^{\mathcal{M}}(s_0, I) = \int_0^b 0.9 e^{-0.9t} \cdot \left(\frac{2}{3} \cdot p_{\max}^{\mathcal{M}}(s_1, I \ominus t) + \frac{1}{3} \cdot p_{\max}^{\mathcal{M}}(s_2, I \ominus t) \right) dt$. $\diamond$

From this example, it is easy to see that an IMC generally induces an integral equation system over the maximum over functions, which is not tractable. Moreover, the iterated integrations that occur are known to be numerically unstable [BHHK03].

Therefore, we resort to a discretization approach: Informally, we divide the time horizon into small time slices. Then we consider IPCs as a discrete-time model which we define such that its steps correspond to the IMC's behavior during a single time slice.

First, we develop a fixed-point characterization for step bounded reachability in IPCs. Then we reduce the maximum time interval bounded reachability problem in IMCs to the step interval bounded reachability problem in the discretized IPC. Finally, we show how to solve the latter by a modified value iteration algorithm.

6.2.2 A fixed point characterization for IPCs

Similar to the timed paths in IMCs, we define $\pi @ n \in \mathcal{S}^* \cup \mathcal{S}^\omega$ for the time abstract paths in IPCs: Let $\#^{PS}(\pi, k) = |\{i < k \mid \pi[i] \in PS\}|$; then $\#^{PS}(\pi, k)$ is the number of probabilistic transitions that complete up to the $(k+1)$ -th state on π . For fixed $n \in \mathbb{N}$, let i be the

smallest index such that $n = \#^{PS}(\pi, i)$. If no such i exists, we set $\pi@n = \langle \rangle$; otherwise i is the index of the state that is reached on path π directly after the n -th probabilistic transition executed (or the first state on π , if $n = 0$). Similarly, let $j \in \mathbb{N}$ be the largest index (or $+\infty$, if no such finite index exists) such that $n = \#^{PS}(\pi, j)$. Then j denotes the position of the $(n+1)$ -th probabilistic state on π . With these preliminaries, we define $\pi@n = \langle s_i, s_{i+1}, \dots, s_{j-1}, s_j \rangle$ to denote the state sequence after the n -th and up to the $(n+1)$ -th probabilistic state of π . Intuitively, $\pi@n$ is the sequence of states which are traversed during the $(n+1)$ -th discrete time unit.

To define step-interval bounded reachability in an IPC $\mathcal{P}$, let $[k_a, k_b] \subseteq \mathbb{N}$ be a step interval. Then

$$\Diamond^{[k_a, k_b]} G = \{ \pi \in Paths_{abs}^\omega \mid \exists n \in \{k_a, k_a + 1, \dots, k_b\}. \exists s' \in \pi@n. s' \in G \}$$

is the set of paths that visit G between discrete time steps k_a and k_b in $\mathcal{P}$. Accordingly, we define the maximum probability for the event $\Diamond^{[k_a, k_b]} G$:

$$p_{max}^{\mathcal{P}}(s, [k_a, k_b]) = \sup_{D \in GM_{abs}} Pr_{v_s, D}^\omega(\Diamond^{[k_a, k_b]} G).$$

Now, we are ready to provide a fixed point characterization for $p_{max}^{\mathcal{P}}$:

Theorem 6.2 (Fixed point characterisation for IPCs). *Let $\mathcal{P} = (\mathcal{S}, Act, IT, PT, v)$ be an IPC, $G \subseteq \mathcal{S}$ a set of goal states and $I = [k_a, k_b]$ a step interval. The function $p_{max}^{\mathcal{P}}$ is the least fixed point of the higher-order operator $\Omega : (\mathcal{S} \times \mathbb{N} \times \mathbb{N} \rightarrow [0, 1]) \rightarrow (\mathcal{S} \times \mathbb{N} \times \mathbb{N} \rightarrow [0, 1])$ which is stated as follows:*

1. For probabilistic states $s \in PS$:

$$\Omega(F)(s, [k_a, k_b]) = \begin{cases} 1 & \text{if } s \in G \wedge k_a = 0 \\ 0 & \text{if } s \notin G \wedge k_a = k_b = 0 \\ \sum_{s' \in \mathcal{S}} PT(s, s') \cdot F(s', [k_a, k_b]) \ominus 1 & \text{otherwise.} \end{cases}$$

2. For interactive states $s \in IS$:

$$\Omega(F)(s, [k_a, k_b]) = \begin{cases} 1 & \text{if } s \in G \text{ and } k_a = 0 \\ \max_{s' \in post^i(s)} F(s', [k_a, k_b]) & \text{otherwise.} \end{cases}$$

Proof. The proof goes along the same lines as the proof of Thm. 6.1. First, we decompose the event $\Diamond^{[k_a, k_b]}$ into disjoint subsets. Therefore, define

$$\Gamma([k_a, k_b], n) = \{ \pi \in Paths_{abs}^\omega \mid \pi[n] \in G \wedge k_a \leq \#^{PS}(\pi, n) \leq k_b \}.$$

To achieve a disjoint decomposition of $\diamond^{[k_a, k_b]} G$, we set $\Pi([k_a, k_b], n) = \Gamma([k_a, k_b], n) \setminus \bigcup_{i=0}^{n-1} \Gamma([k_a, k_b], i)$. Then $\Pi([k_a, k_b], n)$ is the set of paths that visit G in the probabilistic step interval $[k_a, k_b]$ for the first time after exactly n (probabilistic or interactive) transitions. Then $\diamond^{[k_a, k_b]} G = \bigcup_{n=0}^{\infty} \Pi([k_a, k_b], n)$. Thus, it holds for all $D \in GM_{abs}$:

$$Pr_{v,D}^{\omega}(\diamond^{[k_a, k_b]} G) = Pr_{v,D}^{\omega}\left(\bigcup_{n=0}^{\infty} \Pi([k_a, k_b], n)\right) = \sum_{n=0}^{\infty} Pr_{v,D}^{\omega}(\Pi([k_a, k_b], n)).$$

We maximize the probability of the sets $\bigcup_{i=0}^n \Pi([k_a, k_b], i)$ separately: Therefore, let

$$p_{max}^{\mathcal{P},n}(s, [k_a, k_b]) = \sup_{D \in GM_{abs}} Pr_{v_s,D}^{\omega}\left(\bigcup_{i=0}^n \Pi([k_a, k_b], i)\right)$$

be the upper bound on the probability to reach G during the probabilistic step interval $[k_a, k_b]$ with at most n (interactive or probabilistic) transitions. Now we show that $p_{max}^{\mathcal{P},n+1}(s, [k_a, k_b]) = \Omega(p_{max}^{\mathcal{P},n})(s, [k_a, k_b])$:

1. Let $s \in PS$: If $s \in G$ and $k_a = 0$, we have $p_{max}^{\mathcal{P},n+1}(s, [0, k_b]) = 1$. Further, by definition of Ω , it also holds that $\Omega(p_{max}^{\mathcal{P},n})(s, [0, k_b]) = 1$. Hence $\Omega(p_{max}^{\mathcal{P},n})(s, [0, k_b]) = p_{max}^{\mathcal{P},n+1}(s, [0, k_b])$ and we are done.

The case $s \notin G$ and $k_a = k_b = 0$ is similar: We have $p_{max}^{\mathcal{P},n+1}(s, [0, 0]) = 0$, as no probabilistic step may occur in step interval $[0, 0]$. Further $\Omega(p_{max}^{\mathcal{P},n})(s, [0, 0]) = 0$ by definition of Ω . Hence $\Omega(p_{max}^{\mathcal{P},n})(s, [0, 0]) = p_{max}^{\mathcal{P},n+1}(s, [0, 0])$ and we are done.

In the remaining cases, we proceed as follows:

$$\begin{aligned} \Omega(p_{max}^{\mathcal{P},n})(s, [k_a, k_b]) &= \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{max}^{\mathcal{P},n}(s', [k_a, k_b] \ominus 1) \\ &= \sum_{s' \in \mathcal{S}} PT(s, s') \cdot \sup_{D \in GM_{abs}} Pr_{v_{s'},D}^{\omega}\left(\bigcup_{i=0}^n \Pi([k_a, k_b] \ominus 1, i)\right). \end{aligned} \quad (6.7)$$

For $D \in GM_{abs}$, $s \in \mathcal{S}$ and $\sigma \in Act_{\perp}$, we define the scheduler $D_{s,\sigma} \in GM_{abs}$ such that $D_{s,\sigma}(\pi) = D(s \xrightarrow{\sigma} \pi)$ for all $\pi \in Paths_{abs}^*$. This allows us to derive from Eq. (6.7) that

$$\begin{aligned} \Omega(p_{max}^{\mathcal{P},n})(s, [k_a, k_b]) &= \sup_{D \in GM_{abs}} \sum_{s' \in \mathcal{S}} PT(s, s') \cdot Pr_{v_{s'},D_{s,\perp}}^{\omega}\left(\bigcup_{i=0}^n \Pi([k_a, k_b] \ominus 1, i)\right) \\ &= \sup_{D \in GM_{abs}} Pr_{v_s,D}^{\omega}\left(\bigcup_{i=0}^{n+1} \Pi([k_a, k_b], i)\right) = p_{max}^{\mathcal{P},n+1}(s, [k_a, k_b]). \end{aligned}$$

2. Second, we prove that $\Omega(p_{max}^{\mathcal{P},n})(s, [k_a, k_b]) = p_{max}^{\mathcal{P},n+1}(s, [k_a, k_b])$ for interactive states $s \in IS$: If $s \in G$ and $k_a = 0$, it holds that $p_{max}^{\mathcal{P},n}(s, [0, k_b]) = 1$. Further, the definition of Ω implies that $\Omega(p_{max}^{\mathcal{P},n})(s, [0, k_b]) = 1$. Hence $\Omega(p_{max}^{\mathcal{P},n})(s, [0, k_b]) = p_{max}^{\mathcal{P},n+1}(s, [0, k_b])$.

For the other cases, it holds that

$$\begin{aligned}
\Omega(p_{\max}^{\mathcal{P},n})(s, [k_a, k_b]) &= \max_{s' \in \text{post}^i(s)} p_{\max}^{\mathcal{P},n}(s', [k_a, k_b]) \\
&= \max_{s' \in \text{post}^i(s)} \sup_{D \in GM_{\text{abs}}} Pr_{v_{s'}, D}^{\omega} \left(\bigcup_{i=0}^n \Pi([k_a, k_b], i) \right) \\
&= \sup_{D \in GM_{\text{abs}}} \max_{\alpha \in \text{Act}(s)} Pr_{v_{\text{succ}(\alpha), D, \alpha}}^{\omega} \left(\bigcup_{i=0}^n \Pi([k_a, k_b], i) \right) \\
&= \sup_{D \in GM_{\text{abs}}} Pr_{v_{s, D}}^{\omega} \left(\bigcup_{i=0}^{n+1} \Pi([k_a, k_b], i) \right) = p_{\max}^{\mathcal{P},n+1}(s, [k_a, k_b]).
\end{aligned}$$

Hence, $\Omega(p_{\max}^{\mathcal{P},n})(s, [k_a, k_b]) = p_{\max}^{\mathcal{P},n+1}([k_a, k_b], s)$. Further, $p_{\max}^{\mathcal{P},n}(s, [k_a, k_b])$ converges to $p_{\max}^{\mathcal{P}}(s, [k_a, k_b])$ for $n \rightarrow +\infty$: To see this, note that $\bigcup_{i=0}^n \Pi([k_a, k_b], i) \uparrow \Diamond^{[k_a, k_b]}G$ for $n \rightarrow +\infty$. But then Lemma 2.2 implies for all $D \in GM_{\text{abs}}$ that

$$\lim_{n \rightarrow \infty} Pr_{v_{s, D}}^{\omega} \left(\bigcup_{i=0}^n \Pi([k_a, k_b], i) \right) = Pr_{v_{s, D}}^{\omega} (\Diamond^{[k_a, k_b]}G). \quad (6.8)$$

Now, let $\Pi(n) = \bigcup_{i=0}^n \Pi([k_a, k_b], i)$. As Eq. (6.8) applies to all $D \in GM_{\text{abs}}$, it implies that $\sup\{Pr_{v_{s, D}}^{\omega}(\Pi(n)) \mid D \in GM_{\text{abs}}\} \rightarrow \sup\{Pr_{v_{s, D}}^{\omega}(\Diamond^{[k_a, k_b]}G) \mid D \in GM_{\text{abs}}\}$ for $n \rightarrow +\infty$. Taking the limit on both sides of the equation $\Omega(p_{\max}^{\mathcal{P},n})(s, [k_a, k_b]) = p_{\max}^{\mathcal{P},n+1}(s, [k_a, k_b])$ yields that $\Omega(p_{\max}^{\mathcal{P}})(s, [k_a, k_b]) = p_{\max}^{\mathcal{P}}(s, [k_a, k_b])$. Hence $p_{\max}^{\mathcal{P}}$ is a fixed point of Ω .

It remains to show that $p_{\max}^{\mathcal{P}}$ is the least fixed point of Ω : Thus, let F be another fixed point of Ω . By induction on n , we show that $p_{\max}^{\mathcal{P},n}(s, [k_a, k_b]) \leq F(s, [k_a, k_b])$:

1. For the base case, $p_{\max}^{\mathcal{P},0}(s, [k_a, k_b]) = 1 = \Omega(F(s, [k_a, k_b])) = F(s, [k_a, k_b])$ if $s \in G$ and $k_a = 0$ and $p_{\max}^{\mathcal{P},0}(s, [k_a, k_b]) = 0 \leq F(s, [k_a, k_b])$, otherwise. To see this, note that in the event $\Pi([k_a, k_b], 0)$ a G -state must be visited before any (probabilistic or interactive) transition executes.
2. For the induction step, we distinguish two cases:

- (a) Let $s \in PS$: If $s \notin G$ (the case $s \in G$ is similar), then

$$\begin{aligned}
p_{\max}^{\mathcal{P},n+1}(s, [k_a, k_b]) &= \Omega(p_{\max}^{\mathcal{P},n})(s, [k_a, k_b]) \\
&= \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{\max}^{\mathcal{P},n}(s', [k_a, k_b] \ominus 1) \\
&\leq \sum_{s' \in \mathcal{S}} PT(s, s') \cdot F(s', [k_a, k_b] \ominus 1) \quad (* \text{ ind. hyp. } *) \\
&= \Omega(F(s, [k_a, k_b])) = F(s, [k_a, k_b]). \quad (* F \text{ is a fixed point } *)
\end{aligned}$$

- (b) The case $s \in IS$: If $s \in G$ and $k_a = 0$, we have $p_{\max}^{\mathcal{P},n+1}(s, [0, k_b]) = 1$; further, $F(s, [0, k_b]) = \Omega(F)(s, [0, k_b]) = 1$. Hence $p_{\max}^{\mathcal{P},n+1}(s, [0, k_b]) \leq F(s, [0, k_b])$. Otherwise, applying the induction hypothesis yields

$$p_{\max}^{\mathcal{P},n+1}(s, [k_a, k_b]) = \Omega(p_{\max}^{\mathcal{P},n})(s, [k_a, k_b])$$

$$\begin{aligned}
&= \max_{s' \in \text{post}^i(s)} p_{\max}^{\mathcal{P},n}(s', [k_a, k_b]) \\
&\leq \max_{s' \in \text{post}^i(s)} F(s', [k_a, k_b]).
\end{aligned}$$

The definition of Ω implies $\max_{s' \in \text{post}^i(s)} F(s', [k_a, k_b]) = \Omega(F)(s, [k_a, k_b]) = F(s, [k_a, k_b])$, proving that $p_{\max}^{\mathcal{P},n+1}(s, [k_a, k_b]) \leq F(s, [k_a, k_b])$.

Hence $F(s, [k_a, k_b]) \geq \lim_{n \rightarrow \infty} p_{\max}^{\mathcal{P},n}(s, [k_a, k_b]) = p_{\max}^{\mathcal{P}}(s, [k_a, k_b])$, proving the claim. $\square$

Observe the similarity in the treatment of interactive states in the fixed point characterizations for IMCs and IPCs: In an interactive state, the recursive expression of the time-interval bounded reachability in an IMC does not decrease the time interval I for interactive states, whereas for IPCs, the recursive expression does not decrease the step interval $[k_a, k_b]$.

In this way, we have established a close relationship between IMCs and IPCs which allows us to discretize an IMC into an IPC. The details are the topic of the next section.

6.3 A discretization that reduces IMCs to IPCs

For an IMC $\mathcal{M}$ and a *step duration* $\tau > 0$, we define the *discretized IPC* $\mathcal{M}_\tau$ of $\mathcal{M}$ as follows:

Definition 6.8 (Discretization). An IMC $\mathcal{M} = (S, \text{Act}, IT, MT, \nu)$ and a step duration $\tau > 0$ induce the discretized IPC $\mathcal{M}_\tau = (S, \text{Act}, IT, PT, \nu)$, where

$$PT(s, s') = \begin{cases} (1 - e^{-E(s)\tau}) \cdot \mathbf{P}(s, s') & \text{if } s \neq s' \\ (1 - e^{-E(s)\tau}) \cdot \mathbf{P}(s, s') + e^{-E(s)\tau} & \text{if } s = s'. \end{cases} \quad (6.9)$$

Recall, that $\mathbf{P}(s, s') = \frac{\mathbf{R}(s, s')}{E(s)}$ is the discrete branching probability in the IMC $\mathcal{M}$. Moreover, the term $(1 - e^{-E(s)\tau})$ is the probability to leave state s within τ time units; accordingly, $e^{-E(s)\tau}$ denotes the probability to stay in state s for at least τ time units.

Therefore, we can see that in $\mathcal{M}_\tau$, each probabilistic transition $PT(s, s') > 0$ corresponds to one *time step* of length τ in the underlying IMC $\mathcal{M}$: More precisely, $PT(s, s')$ is the probability that a transition to state s' occurs within τ time units. In case that $s' = s$, the first summand in $PT(s, s')$ is the probability to take a self-loop back to s , i.e. a transition that leads from s back to s executes; the second summand denotes the probability that no transition occurs for the next τ time units and the system stays in state $s = s'$.

6.3.1 Approximating time-bounded reachability probabilities

In the next two sections, we prove the correctness of the discretization given in Def. 6.8. To compute the probability $p_{\max}^{\mathcal{M}}(s, [a, b])$, we analyze step-interval bounded reachabil-

ity in the discretized IPC $\mathcal{M}_\tau$, where each step *approximately* corresponds to τ time units. The goal of this section (cf. Thm. 6.3 below on page 171) is the proof that $p_{\max}^{\mathcal{M}_\tau}(s, [0, \lceil \frac{b}{\tau} \rceil])$ converges from below to $p_{\max}^{\mathcal{M}}(s, [0, b])$ if $\tau \rightarrow 0$. Note the restriction in the type of intervals that we allow here: We only consider intervals with closed lower bound 0. Therefore, in this section we only deal with time-bounded reachability probabilities. This is similar to the discretization that we have devised for locally uniform CTMDPs in Sec. 5.3. We address the more complex issue of computing interval-bounded reachability probabilities (where we also allow for lower bounds greater than 0) in Sec. 6.3.2.

Let $\mathcal{M} = (\mathcal{S}, \text{Act}, IT, MT, \nu)$ be an IMC, $G \subseteq \mathcal{S}$ a set of goal states, $I = [0, b] \in \mathcal{Q}$ a time interval with $b > 0$ and $\lambda = \max_{s \in MS} E(s)$. Further, let $\tau > 0$ be such that $b = k_b \tau$ for some $k_b \in \mathbb{N}_{>0}$. Formally, we aim to prove the inequality

$$p_{\max}^{\mathcal{M}_\tau}(s, [0, k_b]) \leq p_{\max}^{\mathcal{M}}(s, I) \leq p_{\max}^{\mathcal{M}_\tau}(s, [0, k_b]) + k_b \cdot \frac{(\lambda\tau)^2}{2}.$$

Both the upper and lower bounds will be proved by induction on k_b . Because of the constraint $k_b \in \mathbb{N}_{>0}$, the induction base is $k_b = 1$. For the induction step ($k_b \rightsquigarrow k_b+1$), we must establish the connection of the probability $p_{\max}^{\mathcal{M}}(s, I)$ in the IMC $\mathcal{M}$ and $p_{\max}^{\mathcal{M}_\tau}(s, [0, k_b])$ in its discretized IPC $\mathcal{M}_\tau$.

This is the main task of the next section, where we first approximate $p_{\max}^{\mathcal{M}}(s, I)$ recursively in terms of $p_{\max}^{\mathcal{M}}(s, I \ominus \tau)$ by exploiting the fixed point characterization of $p_{\max}^{\mathcal{M}}$ which we have established in Thm. 6.1. Intuitively, we express the probability $p_{\max}^{\mathcal{M}}(s, I)$ as the sum of the integration from 0 to τ and the integration from τ to b . Based on this idea, Lemma 6.3 establishes the one-step approximation of $p_{\max}^{\mathcal{M}}(s, I)$.

One-step approximation

We approximate the probability $p_{\max}^{\mathcal{M}}(s, I)$ for all Markovian states $s \in MS \setminus G$ by reducing it to an expression that depends on $p_{\max}^{\mathcal{M}}(s, I \ominus \tau)$. Since $s \notin G$, we obtain a recursive definition of $p_{\max}^{\mathcal{M}}(s, I)$ which is based on the fixed point characterization which is given by Thm. 6.1. Noting that $b \geq \tau$, we obtain:

$$p_{\max}^{\mathcal{M}}(s, I) = \Omega(p_{\max}^{\mathcal{M}})(s, I) = \int_0^b E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus t) dt. \quad (6.10)$$

We let $A(s, I)$ denote the probability that at least one Markovian transition executes at some time point $t \in [0, \tau]$. Accordingly

$$A(s, I) = \int_0^\tau E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus t) dt. \quad (6.11)$$

Splitting the integral on the right-hand side of Eq. (6.10) then yields

$$p_{\max}^{\mathcal{M}}(s, I) = A(s, I) + \int_\tau^b E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus t) dt$$

$$\begin{aligned}
&= A(s, I) + \int_0^{b-\tau} E(s) e^{-E(s)(t+\tau)} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus (t + \tau)) dt \\
&= A(s, I) + e^{-E(s)\tau} \cdot \int_0^{b-\tau} E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus (t + \tau)) dt \\
&= A(s, I) + \underbrace{e^{-E(s)\tau} \cdot p_{\max}^{\mathcal{M}}(s, I \ominus \tau)}_{B(s, I)}.
\end{aligned}$$

Then $B(s, I)$ is the probability that no Markovian transition occurs before time τ (given by the term $e^{-E(s)\tau}$) multiplied with the probability to reach a G -state within the remaining time interval $I \ominus \tau$ (given by the term $p_{\max}^{\mathcal{M}}(s, I \ominus \tau)$).

From the above derivations, we obtain the result that if $s \in MS$ and $p_{\max}^{\mathcal{M}}(s, I)$ is not determined directly (which is the case if $b = 0$ and $s \notin G$ or if $s \in G$), we may express $p_{\max}^{\mathcal{M}}(s, I)$ recursively:

$$\begin{aligned}
p_{\max}^{\mathcal{M}}(s, I) &= \underbrace{\int_0^{\tau} E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus t) dt}_{A(s, I)} \\
&\quad + \underbrace{e^{-E(s)\tau} \cdot p_{\max}^{\mathcal{M}}(s, I \ominus \tau)}_{B(s, I)}.
\end{aligned} \tag{6.12}$$

This recursive characterization of the IMC's behavior in Markovian states permits to derive our discretization: If we define the random variable $\#_{[0, \tau]}$ such that

$$\#_{[0, \tau]} : Paths^{\omega} \rightarrow \mathbb{N} : \pi \mapsto |\{i \in \mathbb{N} \mid \pi[i] \in MS \wedge \Delta(\pi, i+1) \leq \tau\}|.$$

Informally, $\#_{[0, \tau]}(\pi)$ is the number of Markovian transitions that have completed on path π within the first τ time units. For a given $\tau > 0$, we use $\#_{[0, \tau]}$ to decompose the event $\Diamond^I G$ into disjoint sets of paths and obtain

$$\Diamond^I G = \bigcup_{n=0}^{\infty} (\Diamond^I G \cap \#_{[0, \tau]} = n).$$

The term $B(s, I)$ is already suitable for our discretization: Its first factor represents the probability that no transition occurs during the first discretization step, and $p_{\max}^{\mathcal{M}}(s, I \ominus \tau)$ corresponds to the achievable probability in the following discretization steps.

Similarly, $A(s, I)$ is the probability that starting in state s , at least one transition (or equivalently, *one or more* transitions) occurs in time interval $[0, \tau]$. However, its analytic expression given in Eq. (6.11) must be refined before it can be used for a discretization.

Therefore, let us investigate $A(s, I)$ in more detail. Using the random variable $\#_{[0, \tau]}$, we can characterize the event that is associated with the probability $A(s, I)$. This yields

$$A(s, I) = \sup_{D \in GM} Pr_{v_s, D}^{\omega} (\Diamond^I G \cap \#_{[0, \tau]} \geq 1).$$

Further, we can consider each event $(\Diamond^I G \cap \#_{[0,\tau]} = n)$ separately and maximize its probability. Accordingly, define

$$A_n(s, I) = \sup_{D \in GM} Pr_{v_s, D}^\omega(\Diamond^I G \cap \#_{[0,\tau]} = n) \quad (6.13)$$

for all $n \geq 1$. To relate $A(s, I)$ and $A_n(s, I)$, observe that

$$\begin{aligned} A(s, I) &= \sup_{D \in GM} Pr_{v_s, D}^\omega\left(\bigcup_{n=1}^{\infty} (\Diamond^I G \cap \#_{[0,\tau]} = n)\right) \\ &= \sup_{D \in GM} \sum_{n=1}^{\infty} Pr_{v_s, D}^\omega(\Diamond^I G \cap \#_{[0,\tau]} = n) \\ &\leq \sum_{n=1}^{\infty} \sup_{D \in GM} Pr_{v_s, D}^\omega(\Diamond^I G \cap \#_{[0,\tau]} = n) = \sum_{n=1}^{\infty} A_n(s, I). \end{aligned}$$

With these preliminaries, we can approximate the probability $A(s, I)$ by another term $X(s, I)$ which is closely linked to our discretization. The difference between $A(s, I)$ and $X(s, I)$ that makes $X(s, I)$ suitable for our approximation and $A(s, I)$ not, is the fact that $X(s, I)$ does not require an integration over the time interval $[0, \tau]$:

Lemma 6.2 (An approximation for $A(s, I)$). *Let $\mathcal{M} = (\mathcal{S}, Act, IT, MT, v)$ be an IMC, $G \subseteq \mathcal{S}$ a set of goal states, $\tau > 0$ a step duration, $I = [0, b]$ a time-interval with $b \geq \tau$ such that $b = k_b \tau$ for some $k_b \in \mathbb{N}_{>0}$. Further, let $s \in MS \setminus G$ and $\lambda = \max_{s \in MS} E(s)$ be the maximum exit rate in $\mathcal{M}$. We define*

$$X(s, I) = (1 - e^{-E(s)\tau}) \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus \tau). \quad (6.14)$$

Then $X(s, I)$ approximates $A(s, I)$ in the following sense:

$$X(s, I) \leq A(s, I) \leq X(s, I) + \frac{(\lambda\tau)^2}{2}. \quad (6.15)$$

Proof. First we show the lower bound. Obviously, the function $p_{\max}^{\mathcal{M}}(s, [0, b] \ominus t)$ is monotone decreasing for increasing t . Thus:

$$\begin{aligned} A(s, I) &= \int_0^\tau E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus t) dt \\ &\geq \int_0^\tau E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus \tau) dt \\ &= \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus \tau) \int_0^\tau E(s) e^{-E(s)t} dt = X(s, I). \end{aligned}$$

Hence the lower bound follows. To establish the upper bound, first observe that

$$A_1(s, I) \leq X(s, I). \quad (6.16)$$

To see this, note that

$$\begin{aligned} X(s, I) &= \int_0^\tau E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^M(s', I \ominus \tau) dt \quad \text{and} \\ A_1(s, I) &= \int_0^\tau E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot \kappa(s', t, \tau) \cdot p_{\max}^M(s', I \ominus \tau) dt, \end{aligned}$$

where $\kappa(s', t, \tau)$ is the probability that no further Markovian transition occurs in time interval $(t, \tau]$. As $0 \leq \kappa(s', t, \tau) \leq 1$ for all $s' \in \mathcal{S}$ and $t \in [0, \tau]$, Eq. (6.16) follows.

In the following, we first consider the relation between $A_1(s, I)$ and $A(s, I)$. Recall that by definition,

$$A_n(s, I) = \sup_{D \in GM} Pr_{v_s, D}^\omega \left(\diamond^I G \cap \#_{[0, \tau]} = n \right) \leq \sup_{D \in GM} Pr_{v_s, D}^\omega \left(\#_{[0, \tau]} = n \right).$$

Moreover, $\#_{[0, \tau]} = n$ is defined as the event that n Markovian transitions complete within τ time units. Further, $\lambda = \max_{s' \in \mathcal{S}} E(s')$ is the maximum exit rate over all Markovian states in $\mathcal{M}$. Thus $A_n(s, I)$ is bounded by the Poisson distribution $\rho(n, \lambda\tau)$, which gives the probability that exactly n transitions occur within τ time units with rate λ . As $\rho(n, \lambda\tau) = e^{-\lambda\tau} \cdot \frac{(\lambda\tau)^n}{n!}$, we have that $A_n(s, I) \leq \rho(n, \lambda\tau) = e^{-\lambda\tau} \cdot \frac{(\lambda\tau)^n}{n!}$.

If we approximate $A(s, I)$ by considering the term $A_1(s, I)$ only, the probability that we neglect (i.e. the error that we make) is given by the expression $A(s, I) - A_1(s, I)$. This error can be bounded as follows: We have $A(s, I) \leq \sum_{n=1}^\infty A_n(s, I)$; hence $A(s, I) - A_1(s, I) \leq \sum_{n=2}^\infty A_n(s, I)$. Further, the Poisson distribution provides an upper bound for each $A_n(s, I)$. This yields

$$\begin{aligned} A(s, I) &\leq \sum_{n=1}^\infty A_n(s, I) = A_1(s, I) + \sum_{n=2}^\infty A_n(s, I) \\ &\leq A_1(s, I) + \sum_{n=2}^\infty \rho(n, \lambda\tau) = A_1(s, I) + \sum_{n=2}^\infty e^{-\lambda\tau} \cdot \frac{(\lambda\tau)^n}{n!} \\ &= A_1(s, I) + e^{-\lambda\tau} \cdot \sum_{n=2}^\infty \frac{(\lambda\tau)^n}{n!} = A_1(s, I) + e^{-\lambda\tau} \cdot R_1(\lambda\tau), \end{aligned}$$

where $R_1(x) = \sum_{n=2}^\infty \frac{x^n}{n!}$ is the remainder term of the Taylor expansion of $f(x) = e^x$ at point $a = 0$. By Taylor's theorem, there exists $\xi \in [0, \lambda\tau]$ such that

$$R_1(\lambda\tau) = \frac{f^{(2)}(\xi)}{(2)!} \cdot (\lambda\tau)^2 = \frac{e^\xi}{2} \cdot (\lambda\tau)^2. \quad (6.17)$$

To derive an upper bound, choose ξ maximal in $[0, \lambda\tau]$. Then

$$A(s, I) \leq A_1(s, I) + \sum_{n=2}^\infty A_n(s, I) \leq A_1(s, I) + e^{-\lambda\tau} \cdot R_1(\lambda\tau)$$

$$\begin{aligned}
& \stackrel{(6.17)}{=} A_1(s, I) + e^{-\lambda\tau} \cdot \frac{e^\xi}{2} \cdot (\lambda\tau)^2 \leq A_1(s, I) + e^{-\lambda\tau} \cdot \frac{e^{\lambda\tau}}{2} \cdot (\lambda\tau)^2 \\
& = A_1(s, I) + \frac{(\lambda\tau)^2}{2} \stackrel{(6.16)}{\leq} X(s, I) + \frac{(\lambda\tau)^2}{2}. \quad \square
\end{aligned}$$

Lemma 6.2 justifies to use $X(s, I)$ to approximate the probability $A(s, I)$. Now we can establish the relation between $X(s, I)$ and the one-step transition probabilities in the discretized IPC $\mathcal{M}_\tau$ that belongs to $\mathcal{M}$ (cf. Def. 6.8):

Lemma 6.3 (One-step approximation). *Let $\mathcal{M} = (\mathcal{S}, \text{Act}, IT, MT, \nu)$ be an IMC, $\tau > 0$ a step duration and let $\mathcal{M}_\tau = (\mathcal{S}, \text{Act}, IT, PT, \nu)$ be the discretized IPC of $\mathcal{M}$. Further, let $I = [0, b]$ a time-interval with $b \geq \tau$ such that $b = k_b\tau$ for some $k_b \in \mathbb{N}_{>0}$. For all $s \in MS \setminus G$ it holds*

$$p_{\max}^{\mathcal{M}}(s, I) \geq \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus \tau) \quad \text{and} \quad (6.18)$$

$$p_{\max}^{\mathcal{M}}(s, I) \leq \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus \tau) + \frac{(\lambda\tau)^2}{2}. \quad (6.19)$$

Proof. Let $X(s, I)$ be defined as in Lemma 6.2. First, we observe:

$$\begin{aligned}
X(s, I) + B(s, I) &= \underbrace{\left(1 - e^{-E(s)\tau}\right) \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus \tau)}_{X(s, I)} \\
&\quad + \underbrace{e^{-E(s)\tau} \cdot p_{\max}^{\mathcal{M}}(s, I \ominus \tau)}_{B(s, I)} \\
&= \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus \tau).
\end{aligned}$$

Since $p_{\max}^{\mathcal{M}}(s, I) = A(s, I) + B(s, I)$, the statement follows directly by applying Eq. (6.29) of Lemma 6.2. $\square$

Correctness of the reduction to IPC

In this section, we use Lemma 6.3 to prove the correctness of our discretization for computing time-bounded reachability probabilities. However, up to the present point, we only considered states in the set $MS \setminus G$. As a preparation for dealing with interactive states, the following lemma first handles a few special cases:

Lemma 6.4. Let $\mathcal{M} = (\mathcal{S}, \text{Act}, IT, MT, \nu)$ be an IMC, $\tau > 0$ be a step duration and $\mathcal{M}_\tau = (\mathcal{S}, \text{Act}, IT, PT, \nu)$ its discretized IPC. Let $G \subseteq \mathcal{S}$ be a set of goal states, $[0, b] \in \mathcal{Q}$ a time interval such that $b = k_b \tau$ for some $k_b \in \mathbb{N}$. For all $s \in \mathcal{S}$ it holds:

$$(a) \ p_{\max}^{\mathcal{M}}(s, [0, 0]) = p_{\max}^{\mathcal{M}_\tau}(s, [0, 0]).$$

(b) If $\text{Reach}^i(s) \cap MS = \emptyset$ or if $\text{Reach}^i(s) \cap G \neq \emptyset$, then

$$p_{\max}^{\mathcal{M}}(s, [0, b]) = p_{\max}^{\mathcal{M}_\tau}(s, [0, k_b]). \quad (6.20)$$

Proof. We prove each claim separately:

(a) This case is trivial, as both probabilities are 0 if $\text{Reach}^i(s) \cap G = \emptyset$ and 1, otherwise.

(b) For this part we consider the two conditions separately:

- If $\text{Reach}^i(s) \cap MS = \emptyset$, then state s cannot reach a Markovian state. Hence, no more time can pass (time lock).
 - If $\text{Reach}^i(s) \cap G \neq \emptyset$, then a goal state can be reached by taking interactive transitions only. Hence $p_{\max}^{\mathcal{M}}(s, [0, b]) = 1 = p_{\max}^{\mathcal{M}_\tau}(s, [0, k_b])$.
 - If $\text{Reach}^i(s) \cap G = \emptyset$, we cannot reach G along interactive transitions only. Thus $p_{\max}^{\mathcal{M}}(s, [0, b]) = 0 = p_{\max}^{\mathcal{M}_\tau}(s, [0, k_b])$. $\square$
- If $\text{Reach}^i(s) \cap G \neq \emptyset$, then $p_{\max}^{\mathcal{M}}(s, [0, b]) = 1 = p_{\max}^{\mathcal{M}_\tau}(s, [0, k_b])$.

With Lemma 6.4 we have covered three special cases which do not require a discretization to determine the reachability probabilities: No time may pass (no probabilistic transitions may be taken) in the point interval $[0, 0]$ before reaching a G -state. Hence, if $s \notin G$ directly, the set G must be reachable via internal transitions (which consume no time) only. Similarly, if $s \in IS$ is an interactive state such that no Markovian (probabilistic) state is reachable from s , a *time lock* occurs. In this case, the probabilities $p_{\max}(s, I)$ and $p_{\max}^{\mathcal{M}_\tau}(s, I)$ are both 1 if a G -state is reachable via internal transitions and 0, otherwise.

In the remaining cases, we need the discretization technique to compute the time-bounded reachability probabilities. In the following Lemma, we therefore establish the upper error bound of the approximation:

Lemma 6.5 (Upper error bound). Let $\mathcal{M} = (\mathcal{S}, \text{Act}, IT, MT, \nu)$ be an IMC, $G \subseteq \mathcal{S}$ a set of goal states, $\tau > 0$ a step duration, $[0, b]$ a time interval with $b > 0$ such that $b = k_b \tau$ for some $k_b \in \mathbb{N}_{>0}$. Further, let $\lambda = \max_{s \in MS} E(s)$. For all $s \in \mathcal{S}$ it holds:

$$p_{\max}^{\mathcal{M}}(s, [0, b]) \leq p_{\max}^{\mathcal{M}_\tau}(s, [0, k_b]) + k_b \cdot \frac{(\lambda \tau)^2}{2}. \quad (6.21)$$

Proof. We prove Eq. (6.21) by induction on k_b :

1. In the induction base, let $k_b = 1$. We distinguish two cases:

- (a) The case $s \in MS$: If $s \in G$, we have $p_{max}^{\mathcal{M}}(s, [0, \tau]) = 1 = p_{max}^{\mathcal{M}_\tau}(s, [0, 1])$ directly. For $s \notin G$ we can apply Lemma 6.3 and proceed as follows:

$$\begin{aligned} p_{max}^{\mathcal{M}}(s, [0, \tau]) &\stackrel{(6.19)}{\leq} \frac{(\lambda\tau)^2}{2} + \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{max}^{\mathcal{M}}(s', [0, \tau] \ominus \tau) \\ &= \frac{(\lambda\tau)^2}{2} + \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{max}^{\mathcal{M}}(s', [0, 0]) \\ &= \frac{(\lambda\tau)^2}{2} + \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{max}^{\mathcal{M}_\tau}(s', [0, 0]) \quad (* \text{ Lemma 6.4 } *) \\ &= p_{max}^{\mathcal{M}_\tau}(s, [0, 1]) + \frac{(\lambda\tau)^2}{2}. \end{aligned}$$

- (b) The case $s \in IS$: If $Reach^i(s) \cap MS = \emptyset$ or if $Reach^i(s) \cap G \neq \emptyset$, the claim follows by Lemma 6.4 directly. Otherwise, $Reach^i(s) \cap G = \emptyset$ and $Y = Reach^i(s) \cap MS$, where $Y = \{s_1, s_2, \dots, s_n\}$ for some $n \geq 1$. For the induction base, let $I_d = [0, 1]$ be the step-interval that corresponds to the time interval $I = [0, \tau]$. By the fixed-point characterizations of $p_{max}^{\mathcal{M}}(s, I)$ and $p_{max}^{\mathcal{M}_\tau}(s, I_d)$, it holds that

$$\begin{aligned} p_{max}^{\mathcal{M}}(s, I) &= \max \{p_{max}^{\mathcal{M}}(s_1, I), p_{max}^{\mathcal{M}}(s_2, I), \dots, p_{max}^{\mathcal{M}}(s_n, I)\} \\ p_{max}^{\mathcal{M}_\tau}(s, I_d) &= \max \{p_{max}^{\mathcal{M}_\tau}(s_1, I_d), p_{max}^{\mathcal{M}_\tau}(s_2, I_d), \dots, p_{max}^{\mathcal{M}_\tau}(s_n, I_d)\}. \end{aligned}$$

Case (1a) implies for all $s_i \in Y$ that

$$p_{max}^{\mathcal{M}}(s_i, I) \leq p_{max}^{\mathcal{M}_\tau}(s_i, I_d) + \frac{(\lambda\tau)^2}{2}. \quad (6.22)$$

Now pick the state s_k with the maximum probability in $\mathcal{M}$: Formally, choose $s_k \in Y$ such that $p_{max}^{\mathcal{M}}(s, I) = p_{max}^{\mathcal{M}}(s_k, I)$. Then

$$p_{max}^{\mathcal{M}}(s, I) = p_{max}^{\mathcal{M}}(s_k, I) \stackrel{(6.22)}{\leq} p_{max}^{\mathcal{M}_\tau}(s_k, I_d) + \frac{(\lambda\tau)^2}{2} \leq p_{max}^{\mathcal{M}_\tau}(s, I_d) + \frac{(\lambda\tau)^2}{2}.$$

2. In the induction step ($k_b \rightsquigarrow k_b + 1$), we distinguish two cases:

- (a) The case $s \in MS$: If $s \in G$, this case is trivial. Otherwise $s \notin G$ and we apply Lemma 6.3 to derive

$$\begin{aligned} p_{max}^{\mathcal{M}}(s, [0, b + \tau]) &\stackrel{(6.19)}{\leq} \frac{(\lambda\tau)^2}{2} + \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{max}^{\mathcal{M}}(s', [0, b + \tau] \ominus \tau) \\ &= \frac{(\lambda\tau)^2}{2} + \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{max}^{\mathcal{M}}(s', [0, b]) \end{aligned}$$

$$\begin{aligned}
&\stackrel{i.h.}{\leq} \frac{(\lambda\tau)^2}{2} + \sum_{s' \in \mathcal{S}} PT(s, s') \cdot \left(p_{\max}^{\mathcal{M}_\tau}(s', [0, k_b]) + k_b \cdot \frac{(\lambda\tau)^2}{2} \right) \\
&= \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{\max}^{\mathcal{M}_\tau}(s', [0, k_b]) + k_b \cdot \frac{(\lambda\tau)^2}{2} + \frac{(\lambda\tau)^2}{2} \\
&= p_{\max}^{\mathcal{M}_\tau}(s, [0, k_b + 1]) + (k_b + 1) \cdot \frac{(\lambda\tau)^2}{2}.
\end{aligned}$$

- (b) The case $s \in IS$: The same proof as in case (1b) in the induction base applies verbatim, if I and I_d are defined such that $I = [0, b + \tau]$ and $I_d = [0, k_b + 1]$ and if instead of case (1a), the case (2a) of the induction step is used. $\square$

After having established the upper bound, we now complete the proof for the discretization of time-bounded reachability probabilities and establish the lower error bound. Again, we only consider those cases which are not already covered by Lemma 6.4:

Lemma 6.6 (Lower error bound). *Let $\mathcal{M} = (\mathcal{S}, Act, IT, MT, \nu)$ be an IMC, $G \subseteq \mathcal{S}$ a set of goal states, $\tau > 0$ a step duration, $I = [0, b] \in \mathcal{Q}$ a time interval with $b > 0$ such that $b = k_b \tau$ for some $k_b \in \mathbb{N}_{>0}$. Further, let $\lambda = \max_{s \in MS} E(s)$. Then it holds for all $s \in \mathcal{S}$:*

$$p_{\max}^{\mathcal{M}_\tau}(s, [0, k_b]) \leq p_{\max}^{\mathcal{M}}(s, [0, b]). \quad (6.23)$$

Proof. The proof of Eq. (6.23) is by induction on k_b :

1. In the induction base, let $k_b = 1$ (and hence, $b = \tau$). We distinguish two cases:

- (a) The case $s \in MS$: We prove Eq. (6.23) as follows:

$$\begin{aligned}
p_{\max}^{\mathcal{M}}(s, [0, \tau]) &\stackrel{(6.18)}{\geq} \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{\max}^{\mathcal{M}}(s', [0, \tau] \ominus \tau) \\
&= \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{\max}^{\mathcal{M}}(s', [0, 0]) \\
&= \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{\max}^{\mathcal{M}_\tau}(s', [0, 0]) \quad (* \text{ Lemma 6.4 } *) \\
&= p_{\max}^{\mathcal{M}_\tau}(s, [0, 1]).
\end{aligned}$$

- (b) The case $s \in IS$: If $Reach^i(s) \cap MS = \emptyset$ or if $Reach^i(s) \cap G \neq \emptyset$, the claim follows by Lemma 6.4 directly. Otherwise, $Reach^i(s) \cap G = \emptyset$ and $Y = Reach^i(s) \cap MS$, where $Y = \{s_1, s_2, \dots, s_n\}$ for some $n \geq 1$. For the induction base, let $I_d = [0, 1] \subseteq \mathbb{N}$ be the step-interval that corresponds to the time interval $I = [0, \tau]$. By the fixed-point characterizations of $p_{\max}^{\mathcal{M}}(s, I)$ and $p_{\max}^{\mathcal{M}_\tau}(s, I_d)$, it holds that

$$p_{\max}^{\mathcal{M}}(s, I) = \max \{ p_{\max}^{\mathcal{M}}(s_1, I), p_{\max}^{\mathcal{M}}(s_2, I), \dots, p_{\max}^{\mathcal{M}}(s_n, I) \}$$

$$p_{\max}^{\mathcal{M}_\tau}(s, I_d) = \max \{p_{\max}^{\mathcal{M}_\tau}(s_1, I_d), p_{\max}^{\mathcal{M}_\tau}(s_2, I_d), \dots, p_{\max}^{\mathcal{M}_\tau}(s_n, I_d)\}.$$

Case (1a) implies for all $s_i \in Y$ that

$$p_{\max}^{\mathcal{M}_\tau}(s_i, I_d) \leq p_{\max}^{\mathcal{M}}(s_i, I). \quad (6.24)$$

Now pick the state s_k with the maximum probability in $\mathcal{M}_\tau$: Formally, choose $s_k \in Y$ such that $p_{\max}^{\mathcal{M}_\tau}(s, I_d) = p_{\max}^{\mathcal{M}_\tau}(s_k, I_d)$. Then

$$p_{\max}^{\mathcal{M}_\tau}(s, I_d) = p_{\max}^{\mathcal{M}_\tau}(s_k, I_d) \stackrel{(6.24)}{\leq} p_{\max}^{\mathcal{M}}(s_k, I) \leq p_{\max}^{\mathcal{M}}(s, I).$$

2. For the induction step ($k_b \rightsquigarrow k_b + 1$), we distinguish two cases:

(a) The case $s \in MS$: If $s \in G$, this case is trivial. For $s \notin G$ we can apply Lemma 6.3 to derive

$$\begin{aligned} p_{\max}^{\mathcal{M}}(s, [0, b + \tau]) &\stackrel{(6.18)}{\geq} \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{\max}^{\mathcal{M}}(s', [0, b + \tau] \ominus \tau) \\ &= \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{\max}^{\mathcal{M}}(s', [0, b]) \\ &\stackrel{i.h.}{\geq} \sum_{s' \in \mathcal{S}} PT(s, s') \cdot (p_{\max}^{\mathcal{M}_\tau}(s', [0, k_b])) \\ &= p_{\max}^{\mathcal{M}_\tau}(s, [0, k_b + 1]). \end{aligned}$$

(b) The case $s \in IS$: The same proof as in case (1b) in the induction base applies verbatim, if I and I_d are defined such that $I = [0, b + \tau]$ and $I_d = [0, k_b + 1]$ and if instead of case (1a), the case (2a) of the induction step is used. $\square$

Theorem 6.3. Let $\mathcal{M} = (\mathcal{S}, Act, IT, MT, \nu)$ be an IMC, $G \subseteq \mathcal{S}$ a set of goal states, $I = [0, b] \in \mathcal{Q}$ a time interval with $b > 0$ and $\lambda = \max_{s \in MS} E(s)$. Further, let $\tau > 0$ be such that $b = k_b \tau$ for some $k_b \in \mathbb{N}_{>0}$. For all $s \in \mathcal{S}$ it holds:

$$p_{\max}^{\mathcal{M}_\tau}(s, [0, k_b]) \leq p_{\max}^{\mathcal{M}}(s, I) \leq p_{\max}^{\mathcal{M}_\tau}(s, [0, k_b]) + k_b \cdot \frac{(\lambda \tau)^2}{2}.$$

Proof. The upper bound follows by Lemma 6.5 and the lower bound by Lemma 6.6. $\square$

We conclude the discussion for time-bounded reachability with a small example, which also allows us to bridge the gap towards interval bounded reachability in the next section:

Example 6.6. Consider the IMC $\mathcal{M}$ and its discretized IPC $\mathcal{M}_\tau$ in Fig. 6.3(a) and Fig. 6.3(b), resp. Assume that $G = \{s_2\}$ and fix some $\tau > 0$ and $k \in \mathbb{N}_{>0}$. We consider the time interval $I = [0, k\tau]$: In the IMC $\mathcal{M}$, we have $p_{\max}^{\mathcal{M}}(s_0, I) = \int_0^{k\tau} \lambda e^{-\lambda t} \cdot p_{\max}^{\mathcal{M}}(s_1, I \ominus t) dt = 1 - e^{-\lambda k\tau}$. In the IPC $\mathcal{M}_\tau$, we derive $p_{\max}^{\mathcal{M}_\tau}(s_0, [0, k]) = \sum_{i=1}^k (e^{-\lambda\tau})^{i-1} (1 - e^{-\lambda\tau}) = 1 - e^{-\lambda k\tau}$, which is the geometric distribution function for parameter $p = 1 - e^{-\lambda\tau}$. $\diamond$



Figure 6.3: Interval and time-bounded reachability in IMCs.

6.3.2 Approximating interval-bounded reachability probabilities

So far, we only considered intervals of the form $I = [0, b]$, $b > 0$. In what follows, we extend our results to arbitrary intervals. However, this is slightly involved and several aspects have to be considered:

- (a) If $s \in MS$ is a Markovian state and $b > 0$, then $p_{\max}^{\mathcal{M}}(s, (0, b]) = p_{\max}^{\mathcal{M}}(s, [0, b])$. However, this is not true for interactive states: If s_1 (instead of s_0) is made the only initial state in $\mathcal{M}$ and $\mathcal{M}_\tau$ of Fig. 6.3, the probability to reach s_2 in $\mathcal{M}$ within interval $[0, b]$ is 1 whereas it is 0 for the right-semiclosed interval $(0, b]$.
- (b) Further, the discretization does not work for point intervals: To see this, consider Fig. 6.3 again: If $I = [\tau, \tau]$, then $p_{\max}^{\mathcal{M}}(s_0, I) = 0$, as the probability for the Markovian transition that leads from state s_0 to state s_1 to execute exactly at time τ is 0. On the other hand, the corresponding probability in $\mathcal{M}_\tau$ is $p_{\max}^{\mathcal{M}_\tau}(s_0, [1, 1]) = 1 - e^{-\lambda\tau}$.
- (c) Now, let $I = [k_a\tau, k_b\tau]$ be a *closed* interval with $k_a, k_b \in \mathbb{N}$ and $0 < k_a < k_b$. That is, we consider an interval with a lower bound that is larger than 0. Then, in the IMC $\mathcal{M}$ in Fig. 6.3(a), we obtain $p_{\max}^{\mathcal{M}}(s_0, I) = \int_{k_a\tau}^{k_b\tau} \lambda e^{-\lambda t} \cdot p_{\max}^{\mathcal{M}}(s_1, I \ominus t) dt = e^{-\lambda k_a\tau} - e^{-\lambda k_b\tau}$, whereas for its discretized IPC $\mathcal{M}_\tau$ (see Fig. 6.3(b)), we derive

$$p_{\max}^{\mathcal{M}_\tau}(s_0, [k_a, k_b]) = \sum_{i=k_a}^{k_b} (e^{-\lambda\tau})^{i-1} \cdot (1 - e^{-\lambda\tau}) = e^{-\lambda(k_a-1)\tau} - e^{-\lambda k_b\tau}.$$

Clearly, the two probabilities differ in the first term by a factor of $e^{\lambda\tau}$. To see the reason, let $k_a = 2$ and $k_b = 3$: We have $p_{\max}^{\mathcal{M}}(s, [2\tau, 3\tau]) = e^{-2\lambda\tau} - e^{-3\lambda\tau}$; however, in $\mathcal{M}_\tau$ it holds $p_{\max}^{\mathcal{M}_\tau}(s, [2, 3]) = e^{-\lambda\tau} \cdot (1 - e^{-\lambda\tau}) + e^{-2\lambda\tau} \cdot (1 - e^{-\lambda\tau}) = e^{-\lambda\tau} - e^{-3\lambda\tau}$.

This can be explained as follows: As each step in $\mathcal{M}_\tau$ corresponds to a time interval of length τ (cf. Fig. 6.4), the interval bounds 2τ and 3τ fall in different discretization steps. Hence in the discretization, we add two step (instead of only one) which leads to an error.

It is important to note that if we had computed $p_{\max}^{\mathcal{M}}(s, (2\tau, 3\tau])$ instead, we would have obtained the desired result $p_{\max}^{\mathcal{M}}(s, (2, 3]) = p_{\max}^{\mathcal{M}_\tau}(s, [3, 3]) = e^{-2\lambda\tau} - e^{-3\lambda\tau}$.

In the remainder of this section, we prove that our discretization approach also works for approximating time interval-bounded reachability probabilities. Similar to Thm. 6.3

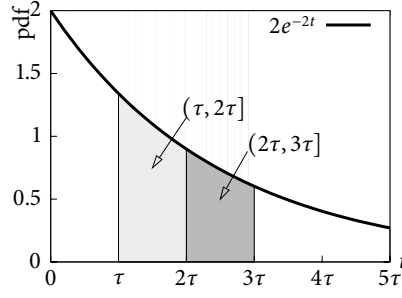


Figure 6.4: Discretization steps.

we obtain a “sandwich” theorem (cf. Thm. 6.4) which provides upper and lower bounds for the discretization error.

We proceed roughly in the same way as in the time-bounded case. However, the technical details are different. In particular, the lower bound proof is completely new, as an important continuity property is violated which holds for time-bounded reachability but not for intervals with lower bounds that are greater than 0.

Let $\mathcal{M} = (\mathcal{S}, \text{Act}, IT, MT, \nu)$ be an IMC, $G \subseteq \mathcal{S}$ a set of goal states, $I = (a, b] \in \mathcal{Q}$ a time interval with $0 \leq a < b$ and $\lambda = \max_{s \in MS} E(s)$. Further, let $\tau > 0$ be such that $a = k_a \tau$ and $b = k_b \tau$ for some $k_a, k_b \in \mathbb{N}$. Formally, we aim to prove that for all $s \in \mathcal{S}$ it holds

$$p_{\max}^{\mathcal{M}_\tau}(s, (k_a, k_b]) - k_a \cdot \frac{(\lambda\tau)^2}{2} \leq p_{\max}^{\mathcal{M}}(s, I) \leq p_{\max}^{\mathcal{M}_\tau}(s, (k_a, k_b]) + k_b \cdot \frac{(\lambda\tau)^2}{2} + \lambda\tau.$$

Similar to the time-bounded case, we begin the discussion in the next section with a one-step approximation. Then we prove in Sec. 6.3.2 that we can reduce the problem of computing (time-)interval bounded reachability probabilities in an IMC $\mathcal{M}$ to computing step-interval bounded reachability probabilities in $\mathcal{M}$'s discretized IPC $\mathcal{M}_\tau$.

One-step approximation

As for the case of time-bounded reachability, we approximate the interval-bounded reachability probability $p_{\max}^{\mathcal{M}}(s, I)$ for intervals $I = (a, b]$ with $0 \leq a < b$ via a discretization technique. For a given step duration $\tau > 0$, we aim to compute the probability that $\mathcal{M}$ moves to a successor state within the next τ time units. Based on the fixed point characterization for $p_{\max}^{\mathcal{M}}$, we distinguish two cases:

1. Let $s \in (MS \setminus G)$: The fact that $a < b$ and $b = k_b \tau$ implies that $b \geq \tau$. We obtain a recursive definition of $p_{\max}^{\mathcal{M}}(s, I)$ by the fixed point theorem (Thm. 6.1 on page 156) as follows:

$$\begin{aligned} p_{\max}^{\mathcal{M}}(s, I) &= \Omega(p_{\max}^{\mathcal{M}})(s, I) \\ &= \int_0^b E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus t) dt. \end{aligned}$$

Similar to Sec. 6.3.1, we can derive that $p_{\max}^{\mathcal{M}}(s, I)$ is the sum $A(s, I) + B(s, I)$ for intervals of the form $(a, b]$: The term

$$A(s, I) = \int_0^\tau E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus t) dt$$

is the probability that a first Markovian transition executes at some time point $t \in [0, \tau]$, and $B(s, I) = e^{-E(s)\tau} \cdot p_{\max}^{\mathcal{M}}(s, I \ominus \tau)$ is the probability that no Markovian transition occurs before time τ and that G is visited in time interval I .

2. If $s \in (MS \cap G)$ and $a = 0$, then $p_{\max}^{\mathcal{M}}(s, I) = 1$ and we can stop; hence, no further recursion is necessary. Otherwise, we have $a \geq \tau$. This case needs further attention: Note that by the fixed point theorem we obtain

$$\begin{aligned} p_{\max}^{\mathcal{M}}(s, I) &= \Omega(p_{\max}^{\mathcal{M}}) \\ &= e^{-E(s)a} + \int_0^a E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus t) dt \\ &= \underbrace{\int_0^\tau E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus t) dt}_{A(s, I)} \\ &\quad + \underbrace{e^{-E(s)a} + \int_\tau^a E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus t) dt}_{B'(s, I)}. \end{aligned} \tag{6.25}$$

Here, as for the previous case, $A(s, I)$ is the probability that a first Markovian transition executes at some time point $t \in [0, \tau]$ and that a G -state is hit afterwards in the remaining time interval $I \ominus t$. It is important to note that the term $B'(s, I)$ in Eq. (6.25) actually corresponds to the term $B(s, I)$ (see Eq. (6.12) on page 164) used for the derivation of the time-bounded case in Sec. 6.3.1. This can be seen by the following derivations:

$$\begin{aligned} B'(s, I) &= e^{-E(s)a} + \int_\tau^a E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus t) dt \\ &= e^{-E(s)a} + \int_0^{a-\tau} E(s) e^{-E(s) \cdot (t+\tau)} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus (t+\tau)) dt \\ &= e^{-E(s)\tau} \left[e^{-E(s) \cdot (a-\tau)} + \int_0^{a-\tau} E(s) e^{-E(s)t} \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus (t+\tau)) dt \right] \\ &= e^{-E(s)\tau} \cdot p_{\max}^{\mathcal{M}}(s, I \ominus \tau) \\ &= B(s, I). \end{aligned}$$

Therefore, $B(s, I)$ can be interpreted as the probability that no Markovian transition occurs before time τ and that G is visited in time interval $I \ominus \tau$.

From the above derivations we conclude that if $s \in MS$ and $p_{max}^{\mathcal{M}}(s, I)$ is not determined directly¹, we may express $p_{max}^{\mathcal{M}}(s, I)$ recursively:

$$\begin{aligned}
 p_{max}^{\mathcal{M}}(s, I) = & \underbrace{\int_0^\tau E(s) e^{-E(s)t} \cdot \sum_{s' \in S} \mathbf{P}(s, s') \cdot p_{max}^{\mathcal{M}}(s', I \ominus t) dt}_{A(s, I)} \\
 & + \underbrace{e^{-E(s)\tau} \cdot p_{max}^{\mathcal{M}}(s, I \ominus \tau)}_{B(s, I)}.
 \end{aligned} \tag{6.26}$$

Note that even though we consider intervals with strict lower bounds, we obtain the same decomposition of $p_{max}^{\mathcal{M}}(s, I)$ as obtained in Eq. (6.12) in the setting of time-bounded reachability objectives. For the remaining derivations in this section, let the random variable $\#_{[0, \tau]}$ and $A_n(s, I)$ (see Eq. (6.13) on page 165) be defined as in Sec. 6.3.1.

We now derive a lower bound for $A(s, I)$: In fact, this is the crucial part for the correctness of our approximation for intervals with lower bounds $a > 0$: Opposed to Sec. 6.3.1, where we make use of the fact that the functions $p_{max}^{\mathcal{M}}(s, [0, b] \ominus t)$ are monotone decreasing for increasing t , this is generally not the case if the lower interval bound a is larger than 0. Thus the way we prove the lower bound in Lemma 6.2 for intervals of the form $[0, b]$ cannot be adapted to the current setting.

For intervals $(a, b]$, the analogue of Lemma 6.2 is Lemma 6.8, where the lower bound is established differently. In its proof, we make use of the following Lemma which considers the case of interval bounds $I = (a, b]$ with $\tau \leq a < b$:

Lemma 6.7 (A lower bound for $A(s, I)$). *Let $\mathcal{M} = (S, Act, IT, MT, \nu)$ be an IMC, $\lambda = \max_{s \in S} E(s)$ be the maximum exit rate in $\mathcal{M}$, $\tau > 0$ a step duration, $s \in MS$ and $I = (a, b]$ a time interval such that $\tau \leq a < b$ and $a = k_a \tau$ and $b = k_b \tau$ for some $k_a, k_b \in \mathbb{N}_{>0}$. Then*

$$A(s, I) \geq \sum_{s' \in S} \mathbf{P}(s, s') \int_0^\tau E(s) e^{-E(s)t} \cdot e^{-\lambda(\tau-t)} \cdot p_{max}^{\mathcal{M}}(s', I \ominus \tau) dt. \tag{6.27}$$

Proof. We have

$$\begin{aligned}
 A(s, I) &= \int_0^\tau E(s) e^{-E(s)t} \cdot \sum_{s' \in S} \mathbf{P}(s, s') \cdot p_{max}^{\mathcal{M}}(s', I \ominus t) dt \\
 &= \sum_{s' \in S} \mathbf{P}(s, s') \int_0^\tau E(s) e^{-E(s)t} \cdot p_{max}^{\mathcal{M}}(s', I \ominus t) dt.
 \end{aligned}$$

¹Examples where the value of $p_{max}^{\mathcal{M}}(s, I)$ is determined directly include the case where $0 \in I$ and $s \in G$ or the case where a time lock occurs.

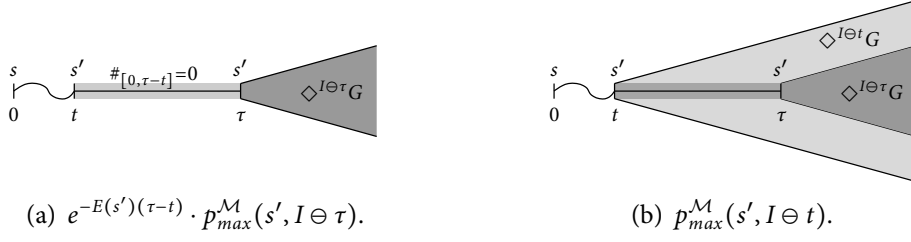


Figure 6.5: Derivation of a lower bound for $A(s, I)$ as used in Lemma 6.7.

Hence, to prove Eq. (6.27) it suffices to show that for all $s' \in \mathcal{S}$ and $t \in [0, \tau]$ it holds that

$$e^{-\lambda(\tau-t)} \cdot p_{\max}^{\mathcal{M}}(s', I \ominus \tau) \leq p_{\max}^{\mathcal{M}}(s', I \ominus t). \quad (6.28)$$

We consider two cases:

1. The case $s' \in MS$: At time t , we took a transition from state s to state $s' \in MS$. Observe that $e^{-E(s')(\tau-t)} \cdot p_{\max}^{\mathcal{M}}(s', I \ominus \tau)$ is the maximum probability for the event that no transition occurs in state s' within the next $(\tau - t)$ time units and that the set G is visited thereafter during the time interval $I \ominus \tau$. Formally, it corresponds to the maximum probability of the event $E_{\text{left}} = (\#_{[0, \tau-t]} = 0 \cap \Diamond^{I \ominus t} G)$ (see Fig. 6.5(a)). On the right hand side, $p_{\max}^{\mathcal{M}}(s', I \ominus t)$ is the maximum probability of the event that G is visited during interval $I \ominus t$, no matter how many transitions occur in the next $(\tau - t)$ time units. Formally, the corresponding event is $E_{\text{right}} = \Diamond^{I \ominus t} G$ (depicted in Fig. 6.5(b)). Hence $E_{\text{left}} \subseteq E_{\text{right}}$. Therefore $e^{-E(s')(\tau-t)} \cdot p_{\max}^{\mathcal{M}}(s', I \ominus \tau) \leq p_{\max}^{\mathcal{M}}(s', I \ominus t)$. Furthermore, $\lambda = \max_{s' \in \mathcal{S}} E(s')$ implies $e^{-\lambda(\tau-t)} \leq e^{-E(s')(\tau-t)}$. Hence Eq. (6.28) follows.
2. The case $s' \in IS$: We consider two sub cases, depending on whether a time lock occurs (the case (2a)) or not (the case (2b)):
 - (a) $\text{Reach}^i(s') \cap MS = \emptyset$: Note that $\text{Reach}^i(s') \cap MS = \emptyset$ implies that only interactive states are reachable from s' , thus the step interval cannot decrease. Further, $I = (a, b]$ and $a \geq \tau$ imply that $0 \notin (I \ominus \tau)$ and $0 \notin (I \ominus t)$. Hence, $p_{\max}^{\mathcal{M}}(s', I \ominus \tau) = 0$ and Eq. (6.28) follows.
 - (b) $\text{Reach}^i(s') \cap MS \neq \emptyset$: Then $\text{Reach}^i(s') \cap MS = Y$, where $Y = \{s_1, s_2, \dots, s_n\}$ for some $n \geq 1$. Then there exist states $s_j, s_k \in Y$ such that $p_{\max}^{\mathcal{M}}(s', I \ominus t) = p_{\max}^{\mathcal{M}}(s_k, I \ominus t)$ and $p_{\max}^{\mathcal{M}}(s', I \ominus \tau) = p_{\max}^{\mathcal{M}}(s_j, I \ominus \tau)$. Therefore we obtain

$$\begin{aligned} e^{-\lambda(\tau-t)} \cdot p_{\max}^{\mathcal{M}}(s', I \ominus \tau) &= e^{-\lambda(\tau-t)} \cdot p_{\max}^{\mathcal{M}}(s_j, I \ominus \tau) \\ &\stackrel{(*)}{\leq} p_{\max}^{\mathcal{M}}(s_j, I \ominus t) \\ &\leq p_{\max}^{\mathcal{M}}(s_k, I \ominus t) = p_{\max}^{\mathcal{M}}(s', I \ominus t), \end{aligned}$$

where $(*)$ follows from case (1). □

With Lemma 6.7 and its new lower bound for $A(s, I)$, we are ready to prove a sandwich lemma that shows that the probabilities $X(s, I)$ approximate $A(s, I)$. It can be regarded as the extension of Lemma 6.2 to the case of intervals with strict lower bounds:

Lemma 6.8 (One-step approximation of $A(s, I)$). *Let $\mathcal{M} = (\mathcal{S}, \text{Act}, IT, MT, \nu)$ be an IMC, $G \subseteq \mathcal{S}$ a set of goal states, $\tau > 0$ a step duration, $I = (a, b]$ a time-interval with $\tau \leq a < b$ such that $a = k_a \tau$ and $b = k_b \tau$ for some $k_a, k_b \in \mathbb{N}_{>0}$. Further, let $s \in MS$ be a Markovian state and $\lambda = \max_{s \in MS} E(s)$ be the maximum exit rate in $\mathcal{M}$. If we define $X(s, I)$ as in Lemma 6.2 then $X(s, I)$ approximates $A(s, I)$ in the following sense:*

$$X(s, I) - \frac{(\lambda \tau)^2}{2} \leq A(s, I) \leq X(s, I) + \frac{(\lambda \tau)^2}{2}. \quad (6.29)$$

Proof. First, let us restate the definition of $X(s, I)$ as given in Lemma. 6.2:

$$X(s, I) = \left(1 - e^{-E(s)\tau}\right) \cdot \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus \tau). \quad (6.30)$$

For the prove, we make use of an approximation of the exponential function e^{-x} . First, note that by the Taylor expansion, $e^{-x} = \sum_{n=0}^{\infty} \frac{(-x)^n}{n!}$. Further, by Taylor's theorem it holds for all $x \geq 0$:

$$e^{-x} \geq 1 - x \quad \text{and} \quad (6.31)$$

$$e^{-x} \leq 1 - x + \frac{x^2}{2}. \quad (6.32)$$

Combining Eq. (6.31) with Lemma 6.7, we have:

$$\begin{aligned} A(s, I) &\geq \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot \int_0^{\tau} E(s) e^{-E(s)t} \cdot e^{-\lambda(\tau-t)} \cdot p_{\max}^{\mathcal{M}}(s', I \ominus \tau) dt \\ &\stackrel{(6.31)}{\geq} \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot \int_0^{\tau} E(s) e^{-E(s)t} \cdot (1 - \lambda(\tau-t)) \cdot p_{\max}^{\mathcal{M}}(s', I \ominus \tau) dt \\ &= \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus \tau) \cdot \int_0^{\tau} E(s) e^{-E(s)t} \cdot (1 - \lambda(\tau-t)) dt. \end{aligned}$$

The integral in the above equation can be simplified as follows:

$$\begin{aligned} &\int_0^{\tau} E(s) e^{-E(s)t} dt - \int_0^{\tau} E(s) e^{-E(s)t} \cdot \lambda(\tau-t) dt \\ &= (1 - e^{-E(s)\tau}) + E(s) \cdot \lambda \cdot \frac{1 - e^{-E(s)\tau} - E(s)\tau}{E(s)^2} \end{aligned}$$

$$\begin{aligned}
&= (1 - e^{-E(s)\tau}) - \frac{\lambda}{E(s)} \cdot \left(E(s)\tau - 1 + \underbrace{e^{-E(s)\tau}}_{\text{Taylor's theorem}} \right) \\
&\stackrel{(6.32)}{\geq} (1 - e^{-E(s)\tau}) - \frac{\lambda}{E(s)} \cdot \left(E(s)\tau - 1 + \left(1 - E(s)\tau + \frac{(E(s)\tau)^2}{2} \right) \right) \\
&= (1 - e^{-E(s)\tau}) - \frac{\lambda}{E(s)} \cdot \left(\frac{(E(s)\tau)^2}{2} \right) \\
&= (1 - e^{-E(s)\tau}) - \frac{\lambda E(s)\tau^2}{2} \\
&\geq (1 - e^{-E(s)\tau}) - \frac{(\lambda\tau)^2}{2}. \quad (* \text{ as } \lambda \geq E(s) *)
\end{aligned}$$

Therefore, we obtain the lower bound for $A(s, I)$:

$$\begin{aligned}
A(s, I) &\geq \sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus \tau) \cdot \left[(1 - e^{-E(s)\tau}) - \frac{(\lambda\tau)^2}{2} \right] \\
&\stackrel{(6.30)}{=} X(s, I) - \left[\sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus \tau) \cdot \frac{(\lambda\tau)^2}{2} \right] \\
&= X(s, I) - \frac{(\lambda\tau)^2}{2} \cdot \underbrace{\sum_{s' \in \mathcal{S}} \mathbf{P}(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus \tau)}_{\leq 1} \\
&\geq X(s, I) - \frac{(\lambda\tau)^2}{2}.
\end{aligned}$$

For the derivation of the upper bound, the respective proof in Lemma 6.2 applies verbatim, with $A(s, I)$ defined for right-semiclosed intervals. $\square$

Now that we have established lower and upper bounds for the approximation of the probability $A(s, I)$, we are ready to extend this result to our discretization. Therefore, in the next lemma we establish the relationship between the approximation for $A(s, I)$ and our discretization. It serves the same purpose as Lemma 6.3 in Sec. 6.3.1, but also accounts for the error that is induced by lower interval bounds that are larger than 0:

Lemma 6.9 (One-step approximation). *Let $\mathcal{M} = (S, \text{Act}, IT, MT, \nu)$ be an IMC, $\tau > 0$ a step duration and let $\mathcal{M}_\tau = (S, \text{Act}, IT, PT, \nu)$ be the discretized IPC of $\mathcal{M}$. Further, let $I = (a, b]$ be a time interval with $\tau \leq a < b$ such that $a = k_a\tau$ and $b = k_b\tau$ for some $k_a, k_b \in \mathbb{N}_{>0}$. For $s \in MS$ it holds:*

$$p_{\max}^{\mathcal{M}}(s, I) \geq \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{\max}^{\mathcal{M}_\tau}(s', I \ominus \tau) - \frac{(\lambda\tau)^2}{2} \quad \text{and} \quad (6.33)$$

$$p_{\max}^{\mathcal{M}}(s, I) \leq \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{\max}^{\mathcal{M}}(s', I \ominus \tau) + \frac{(\lambda\tau)^2}{2}. \quad (6.34)$$

Proof. The proof goes along the same lines as the proof of Lemma 6.3 if the approximation result obtained in Eq. (6.29) of Lemma 6.8 is used. $\square$

Correctness of the reduction to IPC

We first establish the upper bound for Thm. 6.4. Note that in contrast to the Lemmas before, we now allow for intervals of the form $(0, b]$, that is, we allow the lower bound a of the right-semiclosed intervals I to be 0.

Lemma 6.10 (Upper error bound). *Let $\mathcal{M} = (\mathcal{S}, Act, IT, MT, v)$ be an IMC, $G \subseteq \mathcal{S}$ a set of goal states, $\tau > 0$ a step duration, $(a, b]$ a time interval with $0 \leq a < b$ such that $a = k_a\tau$ and $b = k_b\tau$ for some $k_a \in \mathbb{N}$ and $k_b \in \mathbb{N}_{>0}$. Further, let $\lambda = \max_{s \in MS} E(s)$. For all $s \in \mathcal{S}$ it holds:*

$$p_{\max}^{\mathcal{M}}(s, (a, b]) \leq p_{\max}^{\mathcal{M}_\tau}(s, (k_a, k_b]) + k_b \cdot \frac{(\lambda\tau)^2}{2} + \lambda\tau. \quad (6.35)$$

Proof. We prove Eq. (6.35) by induction on k_a :

1. In the induction base, let $k_a = 0$ (implying $a = 0$). We consider three cases:

- (a) The case $s \in MS \setminus G$:

$$\begin{aligned} p_{\max}^{\mathcal{M}}(s, (0, b]) &= p_{\max}^{\mathcal{M}}(s, [0, b]) \\ &\leq p_{\max}^{\mathcal{M}_\tau}(s, [0, k_b]) + k_b \cdot \frac{(\lambda\tau)^2}{2} \quad (* \text{ by Thm. 6.3 } *) \\ &\stackrel{(*)}{=} p_{\max}^{\mathcal{M}_\tau}(s, [1, k_b]) + k_b \cdot \frac{(\lambda\tau)^2}{2} \\ &= p_{\max}^{\mathcal{M}_\tau}(s, (0, k_b]) + k_b \cdot \frac{(\lambda\tau)^2}{2}, \end{aligned}$$

where $(*)$ follows from the fact that $s \in MS \setminus G$ implies $p_{\max}^{\mathcal{M}_\tau}(s, [1, b]) = p_{\max}^{\mathcal{M}_\tau}(s, [0, b])$. Hence, if $a = 0$, Eq. (6.35) even holds for a tighter upper bound.

- (b) The case $s \in MS \cap G$: In this case, the discretization induces an additional error which can be bound from above by the term $\lambda\tau$: In contrast to case (1a), in the case that $s \in MS \cap G$ we have that $p_{\max}^{\mathcal{M}}(s, (0, b]) = 1$, whereas $p_{\max}^{\mathcal{M}_\tau}(s, (0, k_b]) = p_{\max}^{\mathcal{M}_\tau}(s, [1, k_b]) \geq e^{-\lambda\tau}$. Intuitively, the discretization requires one discretization step to pass, in which the goal state s could be left. The probability for this

to happen is $(1 - e^{-E(s)\tau})$ which can be bounded by the Taylor expansion as follows: $(1 - e^{-E(s)\tau}) \leq (1 - e^{-\lambda\tau}) = (1 - (1 - \lambda\tau + R_1(\lambda\tau)))$, where $R_1(\lambda\tau) > 0$. Hence $(1 - e^{-E(s)\tau}) \leq \lambda\tau$. With these remarks we can derive

$$\begin{aligned} p_{\max}^{\mathcal{M}}(s, (0, b]) &= p_{\max}^{\mathcal{M}}(s, [0, b]) \leq p_{\max}^{\mathcal{M}_\tau}(s, [0, k_b]) + k_b \cdot \frac{(\lambda\tau)^2}{2} \\ &\leq p_{\max}^{\mathcal{M}_\tau}(s, [1, k_b]) + k_b \cdot \frac{(\lambda\tau)^2}{2} + \lambda\tau \\ &= p_{\max}^{\mathcal{M}_\tau}(s, (0, k_b]) + k_b \cdot \frac{(\lambda\tau)^2}{2} + \lambda\tau. \end{aligned}$$

(c) If $s \in IS$, we distinguish two cases:

- i. If $\text{Reach}^i(s) \cap MS = \emptyset$, then $p_{\max}^{\mathcal{M}}(s, (0, b]) = 0 = p_{\max}^{\mathcal{M}_\tau}(s, (0, k_b])$.
- ii. Otherwise, $\text{Reach}^i(s) \cap MS \neq \emptyset$ and $\text{Reach}^i(s) \cap MS = Y$ for some $Y = \{s_1, s_2, \dots, s_n\}$ and $n \geq 1$. Let $I = (0, b]$ and $I_d = (0, k_b]$. Then

$$\begin{aligned} p_{\max}^{\mathcal{M}}(s, I) &= \max \{p_{\max}^{\mathcal{M}}(s_1, I), p_{\max}^{\mathcal{M}}(s_2, I), \dots, p_{\max}^{\mathcal{M}}(s_n, I)\} \text{ and} \\ p_{\max}^{\mathcal{M}_\tau}(s, I_d) &= \max \{p_{\max}^{\mathcal{M}_\tau}(s_1, I_d), p_{\max}^{\mathcal{M}_\tau}(s_2, I_d), \dots, p_{\max}^{\mathcal{M}_\tau}(s_n, I_d)\}. \end{aligned}$$

Now choose $s_k \in Y$ such that $p_{\max}^{\mathcal{M}}(s, I) = p_{\max}^{\mathcal{M}}(s_k, I)$. Depending on whether $s_k \notin G$ or $s_k \in G$, cases (1a) or (1b) apply, respectively. Hence

$$\begin{aligned} p_{\max}^{\mathcal{M}}(s, I) &= p_{\max}^{\mathcal{M}}(s_k, I) \leq p_{\max}^{\mathcal{M}_\tau}(s_k, I_d) + k_b \cdot \frac{(\lambda\tau)^2}{2} + \lambda\tau \\ &\leq p_{\max}^{\mathcal{M}_\tau}(s, I_d) + k_b \cdot \frac{(\lambda\tau)^2}{2} + \lambda\tau. \end{aligned}$$

- 2. For the induction step ($k_a \rightsquigarrow k_a + 1$), assume Eq. (6.35) holds for k_a . We show that it holds for $k_a + 1$. Therefore, we distinguish two cases:

(a) The case $s \in MS$: Since $a + \tau \geq \tau$, we can apply Lemma 6.9 and obtain:

$$\begin{aligned} p_{\max}^{\mathcal{M}}(s, (a + \tau, b]) &\stackrel{(6.34)}{\leq} \frac{(\lambda\tau)^2}{2} + \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{\max}^{\mathcal{M}}(s', (a + \tau, b] \ominus \tau) \\ &= \frac{(\lambda\tau)^2}{2} + \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{\max}^{\mathcal{M}}(s', (a, b - \tau]) \\ &\stackrel{i.h.}{\leq} \frac{(\lambda\tau)^2}{2} + \sum_{s' \in \mathcal{S}} PT(s, s') \cdot \left(p_{\max}^{\mathcal{M}_\tau}(s', (k_a, k_b - 1]) + (k_b - 1) \cdot \frac{(\lambda\tau)^2}{2} + \lambda\tau \right) \\ &= \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{\max}^{\mathcal{M}_\tau}(s', (k_a, k_b - 1]) + (k_b - 1) \cdot \frac{(\lambda\tau)^2}{2} + \frac{(\lambda\tau)^2}{2} + \lambda\tau \\ &= p_{\max}^{\mathcal{M}_\tau}(s, (k_a + 1, k_b]) + k_b \cdot \frac{(\lambda\tau)^2}{2} + \lambda\tau. \end{aligned}$$

- (b) The case $s \in IS$: We consider two cases: If $Reach^i(s) \cap MS = \emptyset$, the claim follows directly, as $p_{max}^{\mathcal{M}}(s, (a, b]) = p_{max}^{\mathcal{M}_\tau}(s, (k_a, k_b]) = 0$. Otherwise $Reach^i(s) \cap MS \neq \emptyset$ and $Reach^i(s) \cap MS = Y$ for some $Y = \{s_1, s_2, \dots, s_n\}$ and $n \geq 1$. Now let $I_d = (k_a + 1, k_b] \subseteq \mathbb{N}$ be the step-interval that corresponds to the time interval $I = (a + \tau, b]$. By the fixed-point characterizations of $p_{max}^{\mathcal{M}}(s, I)$ and $p_{max}^{\mathcal{M}_\tau}(s, I_d)$ it holds that

$$\begin{aligned} p_{max}^{\mathcal{M}}(s, I) &= \max \{p_{max}^{\mathcal{M}}(s_1, I), p_{max}^{\mathcal{M}}(s_2, I), \dots, p_{max}^{\mathcal{M}}(s_n, I)\} \\ p_{max}^{\mathcal{M}_\tau}(s, I_d) &= \max \{p_{max}^{\mathcal{M}_\tau}(s_1, I_d), p_{max}^{\mathcal{M}_\tau}(s_2, I_d), \dots, p_{max}^{\mathcal{M}_\tau}(s_n, I_d)\}. \end{aligned}$$

Case (2a) implies for all $s_i \in Y$ that

$$p_{max}^{\mathcal{M}}(s_i, I) \leq p_{max}^{\mathcal{M}_\tau}(s_i, I_d) + k_b \cdot \frac{(\lambda\tau)^2}{2} + \lambda\tau. \quad (6.36)$$

Now pick the state s_k with the maximum probability in $\mathcal{M}$: Formally, choose $s_k \in Y$ such that $p_{max}^{\mathcal{M}}(s_k, I) = p_{max}^{\mathcal{M}}(s, I)$. Then

$$\begin{aligned} p_{max}^{\mathcal{M}}(s, I) &= p_{max}^{\mathcal{M}}(s_k, I) \\ &\stackrel{(6.36)}{\leq} p_{max}^{\mathcal{M}_\tau}(s_k, I_d) + k_b \cdot \frac{(\lambda\tau)^2}{2} + \lambda\tau \\ &\leq p_{max}^{\mathcal{M}_\tau}(s, I_d) + k_b \cdot \frac{(\lambda\tau)^2}{2} + \lambda\tau. \quad \square \end{aligned}$$

We continue and prove the lower bound of Thm. 6.4. Again, we consider right-semi-closed intervals $(a, b]$ and also allow for the case $a = 0$:

Lemma 6.11 (Lower error bound). *Let $\mathcal{M} = (S, Act, IT, MT, v)$ be an IMC, $G \subseteq S$ a set of goal states, $\tau > 0$ a step duration, $I = (a, b] \in \mathcal{Q}$ a time interval with $0 \leq a < b$ such that $a = k_a\tau$ and $b = k_b\tau$ for some $k_a \in \mathbb{N}$ and $k_b \in \mathbb{N}_{>0}$, $k_a < k_b$. Further, let $\lambda = \max_{s \in MS} E(s)$. For for all $s \in S$ it holds:*

$$p_{max}^{\mathcal{M}_\tau}(s, (k_a, k_b]) - k_a \cdot \frac{(\lambda\tau)^2}{2} \leq p_{max}^{\mathcal{M}}(s, (a, b]). \quad (6.37)$$

Proof. The proof is by induction on k_a :

1. For the induction base, let $k_a = 0$ (implying $a = 0$). We consider two cases:

- (a) The case $s \in MS$:

$$\begin{aligned} p_{max}^{\mathcal{M}}(s, (0, b]) &= p_{max}^{\mathcal{M}}(s, [0, b]) \\ &\geq p_{max}^{\mathcal{M}_\tau}(s, [0, k_b]) \quad (* \text{ by Thm. 6.3*}) \\ &\geq p_{max}^{\mathcal{M}_\tau}(s, [1, k_b]) \\ &= p_{max}^{\mathcal{M}_\tau}(s, (0, k_b]). \end{aligned}$$

- (b) The case $s \in IS$: We distinguish two sub cases, depending on whether a time lock occurs or not:

- i. If $Reach^i(s) \cap MS = \emptyset$, then $p_{max}^{\mathcal{M}}(s, (0, b]) = 0 = p_{max}^{\mathcal{M}_\tau}(s, (0, k_b])$.
- ii. Otherwise, $Reach^i(s) \cap MS \neq \emptyset$ and $Reach^i(s) \cap MS = Y$ for some $Y = \{s_1, s_2, \dots, s_n\}$ and $n \geq 1$. Let $I = (0, b]$ and $I_d = (0, k_b]$. Then

$$p_{max}^{\mathcal{M}}(s, I) = \max \{p_{max}^{\mathcal{M}}(s_1, I), p_{max}^{\mathcal{M}}(s_2, I), \dots, p_{max}^{\mathcal{M}}(s_n, I)\} \text{ and} \\ p_{max}^{\mathcal{M}_\tau}(s, I_d) = \max \{p_{max}^{\mathcal{M}_\tau}(s_1, I_d), p_{max}^{\mathcal{M}_\tau}(s_2, I_d), \dots, p_{max}^{\mathcal{M}_\tau}(s_n, I_d)\}.$$

Now, choose $s_k \in Y$ such that $p_{max}^{\mathcal{M}_\tau}(s_k, I_d) = p_{max}^{\mathcal{M}_\tau}(s, I_d)$. Then case (1a) applies and we obtain

$$p_{max}^{\mathcal{M}_\tau}(s, I_d) = p_{max}^{\mathcal{M}_\tau}(s_k, I_d) \leq p_{max}^{\mathcal{M}}(s_k, I) \leq p_{max}^{\mathcal{M}}(s, I).$$

2. For the induction step ($k_a \rightsquigarrow k_a + 1$ and $a \rightsquigarrow a + \tau$), assume that Eq. (6.37) holds for k_a . We show that it also holds for $k_a + 1$. Therefore, consider two cases:

- (a) The case $s \in MS$: Since $a + \tau \geq \tau$, we can apply Lemma 6.9 and obtain:

$$\begin{aligned} p_{max}^{\mathcal{M}}(s, (a + \tau, b]) &\stackrel{(6.33)}{\geq} \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{max}^{\mathcal{M}}(s', (a + \tau, b] \ominus \tau) - \frac{(\lambda\tau)^2}{2} \\ &= \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{max}^{\mathcal{M}}(s', (a, b - \tau]) - \frac{(\lambda\tau)^2}{2} \\ &\stackrel{i.h.}{\geq} \sum_{s' \in \mathcal{S}} PT(s, s') \cdot \left(p_{max}^{\mathcal{M}_\tau}(s', (k_a, k_b - 1]) - k_a \cdot \frac{(\lambda\tau)^2}{2} \right) - \frac{(\lambda\tau)^2}{2} \\ &= p_{max}^{\mathcal{M}_\tau}(s, (k_a + 1, k_b]) - (k_a + 1) \cdot \frac{(\lambda\tau)^2}{2}. \end{aligned}$$

- (b) The case $s \in IS$: We consider two cases: If $Reach^i(s) \cap MS = \emptyset$, the claim follows directly, as $p_{max}^{\mathcal{M}}(s, (a, b]) = p_{max}^{\mathcal{M}_\tau}(s, (k_a, k_b]) = 0$. Otherwise, $Reach^i(s) \cap MS \neq \emptyset$. Hence, $Reach^i(s) \cap MS = Y$ for some $Y = \{s_1, s_2, \dots, s_n\}$ and $n \geq 1$. Now let $I_d = (k_a + 1, k_b] \subseteq \mathbb{N}$ be the step-interval that corresponds to the time interval $I = (a + \tau, b]$. By the fixed-point characterizations of $p_{max}^{\mathcal{M}}(s, I)$ and $p_{max}^{\mathcal{M}_\tau}(s, I_d)$ it holds that

$$p_{max}^{\mathcal{M}}(s, I) = \max \{p_{max}^{\mathcal{M}}(s_1, I), p_{max}^{\mathcal{M}}(s_2, I), \dots, p_{max}^{\mathcal{M}}(s_n, I)\} \\ p_{max}^{\mathcal{M}_\tau}(s, I_d) = \max \{p_{max}^{\mathcal{M}_\tau}(s_1, I_d), p_{max}^{\mathcal{M}_\tau}(s_2, I_d), \dots, p_{max}^{\mathcal{M}_\tau}(s_n, I_d)\}.$$

Case (2a) implies for all $s_i \in Y$ that

$$p_{max}^{\mathcal{M}_\tau}(s_i, I_d) - (k_a + 1) \cdot \frac{(\lambda\tau)^2}{2} \leq p_{max}^{\mathcal{M}}(s_i, I). \quad (6.38)$$

Now pick the state s_k with the maximum probability in $\mathcal{M}_\tau$: Formally, choose $s_k \in Y$ such that $p_{\max}^{\mathcal{M}_\tau}(s, I) = p_{\max}^{\mathcal{M}_\tau}(s_k, I)$. Then

$$\begin{aligned} p_{\max}^{\mathcal{M}_\tau}(s, I_d) - (k_a + 1) \cdot \frac{(\lambda\tau)^2}{2} &= p_{\max}^{\mathcal{M}_\tau}(s_k, I_d) - (k_a + 1) \cdot \frac{(\lambda\tau)^2}{2} \\ &\stackrel{(6.38)}{\leq} p_{\max}^{\mathcal{M}}(s_k, I) \leq p_{\max}^{\mathcal{M}}(s, I). \end{aligned} \quad \square$$

With the technical details in Lemma 6.10 and Lemma 6.11, we have established both a lower and an upper error bound. They are the main result of this section and summarized in the following theorem, which states the correctness of our approximation technique for right-semiclosed intervals:

Theorem 6.4. *Let $\mathcal{M} = (\mathcal{S}, \text{Act}, IT, MT, \nu)$ be an IMC, $G \subseteq \mathcal{S}$ a set of goal states, $I = (a, b] \in \mathcal{Q}$ a time interval with $0 \leq a < b$ and $\lambda = \max_{s \in MS} E(s)$. If $\tau > 0$ is such that $a = k_a\tau$ and $b = k_b\tau$ for some $k_a \in \mathbb{N}$ and $k_b \in \mathbb{N}_{>0}$, then it holds for all $s \in \mathcal{S}$:*

$$p_{\max}^{\mathcal{M}_\tau}(s, (k_a, k_b]) - k_a \cdot \frac{(\lambda\tau)^2}{2} \leq p_{\max}^{\mathcal{M}}(s, I) \leq p_{\max}^{\mathcal{M}_\tau}(s, (k_a, k_b]) + k_b \cdot \frac{(\lambda\tau)^2}{2} + \lambda\tau.$$

Proof. The claim follows directly from Lemma 6.10 and Lemma 6.11. $\square$

With the results of Thm. 6.3 and Thm. 6.4, we have a correct approximation for intervals of the form $[0, b]$ and $(a, b]$, respectively. This suffices to also establish the correctness for open and left-semiclosed intervals and for closed intervals that have a lower bound that is larger than 0:

Theorem 6.5. *Let $\mathcal{M} = (\mathcal{S}, \text{Act}, IT, MT, \nu)$ be an IMC, $G \subseteq \mathcal{S}$ a set of goal states and $\tau > 0$ a step duration. Further, let $I \in \mathcal{Q}$ be a time interval with $\inf I = a$ and $\sup I = b$ such that $a < b$ and $a = k_a\tau$ and $b = k_b\tau$ for some $k_a \in \mathbb{N}$ and $k_b \in \mathbb{N}_{>0}$. If $0 \notin I$ it holds for all $s \in \mathcal{S}$:*

$$p_{\max}^{\mathcal{M}_\tau}(s, (k_a, k_b]) - k_a \cdot \frac{(\lambda\tau)^2}{2} \leq p_{\max}^{\mathcal{M}}(s, I) \leq p_{\max}^{\mathcal{M}_\tau}(s, (k_a, k_b]) + k_b \cdot \frac{(\lambda\tau)^2}{2} + \lambda\tau.$$

Proof. We consider the following cases according to the form of the interval I :

1. The case $I = (a, b]$: Follows directly from Thm. 6.4.
2. The case $I = [a, b]$: By the assumption $0 \notin I$, we have $a > 0$. For $s \in MS$, $[a, b]$ can be replaced by $(a, b]$ without changing the probability. For $s \in IS$, $a > 0$ implies also that $p_{\max}^{\mathcal{M}}(s, I) = p_{\max}^{\mathcal{M}}(s', I)$ for some Markovian state s' . Thus,

$$p_{\max}^{\mathcal{M}}(s, I) = p_{\max}^{\mathcal{M}}(s', I) = p_{\max}^{\mathcal{M}}(s', (a, b]) = p_{\max}^{\mathcal{M}}(s, (a, b])$$

The claims follows then by applying the first case.

3. The case $I = (a, b)$ or $I = [a, b)$: Since $b > 0$, this case can be proved in a similar way as the previous one. $\square$

For the remaining cases, note that for all states $s \in \mathcal{S}$ and time interval $I = \emptyset$ it holds that $p_{\max}^{\mathcal{M}}(s, I) = 0$. As we have shown in the introductory remark, the discretization does not work for general point intervals $[a, a]$. However, if $I = [0, 0]$, an interactive reachability analysis suffices to compute $p_{\max}^{\mathcal{M}}(s, I)$, which is either 1 or 0. Hence, these cases do not require a discretization as the probabilities can be determined directly.

6.4 Solving the problem on the reduced IPC

In Sec. 6.3 we have proved that the interval-bounded reachability probability in an IMC $\mathcal{M}$ can be approximated arbitrarily closely by computing the corresponding step-interval bounded reachability probability in $\mathcal{M}$'s induced (discrete-time) IPC. However, we did not propose an efficient method to compute the latter.

In this section, we will fill this gap. In order to be as general as possible, we consider an arbitrary IPC $\mathcal{P} = (\mathcal{S}, Act, IT, PT, \nu)$ and a set of goal states $G \subseteq \mathcal{S}$ together with a step-interval $[k_a, k_b]$ with $k_a, k_b \in \mathbb{N}$, $k_a < k_b$.

We discuss how to compute $p_{\max}^{\mathcal{P}}(s, [k_a, k_b])$ via a modification of the well known *value iteration* algorithm [Ber95] for MDPs. However, the adaptation is more involved than the one used in Sec. 5.3.1 for locally uniform CTMDPs, as we have to extend the algorithm to correctly handle interactive transitions. More precisely, our adaptation needs to consider step intervals that correspond to the number of *probabilistic steps* that are taken. This is reflected in our algorithm which only decreases the step counter for probabilistic, but not for internal transitions.

As done before, we discuss step bounded reachability first and extend our results to step-intervals later.

6.4.1 Maximum step bounded reachability

We aim at computing $p_{\max}^{\mathcal{P}}(s, [0, k])$ for some $k > 0$. This works as follows: In each step $i = 0, 1, \dots, k$ of the value iteration, we use two vectors $\vec{v}_i \in [0, 1]^{\mathcal{S}}$ and $\vec{u}_i \in [0, 1]^{\mathcal{S}}$, where $\vec{v}_i$ is the probability vector obtained from $\vec{u}_{i-1}$ by one step in the classical value iteration algorithm and $\vec{u}_i$ is obtained by computing the backwards closure along interactive transitions with respect to $\vec{v}_{i-1}$.

Each of the k value iteration steps consists of two phases. We describe the i -th value iteration step:

1. First, $\vec{v}_i$ is computed: For the first value iteration step, we set $\vec{v}_0(s) = 1$ if $s \in G$ and $\vec{v}_0(s) = 0$, otherwise. In the subsequent steps, the vector $\vec{v}_i$ is obtained as follows:

If $s \in PS \cap G$, then $\vec{v}_i(s) = 1$. If $s \in PS \setminus G$, then $\vec{v}_i(s)$ is the weighted sum of the probabilistic successor states s' of s , multiplied by the result $\vec{u}_{i-1}(s')$ of the previous value iteration step. Finally, for interactive states, the result from the previous value iteration step propagates into $\vec{v}_i$. Formally, for all $0 < i \leq k$:

$$\vec{v}_i(s) = \begin{cases} 1 & \text{if } s \in PS \cap G \\ \sum_{s' \in \mathcal{S}} PT(s, s') \cdot \vec{u}_{i-1}(s') & \text{if } s \in PS \setminus G \\ \vec{u}_{i-1}(s) & \text{if } s \in IS. \end{cases} \quad (6.39)$$

2. In the second phase, $\vec{u}_i$ is obtained by the backwards closure of $\vec{v}_i$ along internal transitions. Formally, the vector $\vec{u}_i$ is obtained according to the following equation:

$$\vec{u}_i(s) = \max \{ \vec{v}_i(s') \mid s \rightsquigarrow_i^* s' \}.$$

Note that for efficiency reasons, the set $\{s' \in \mathcal{S} \mid s \rightsquigarrow_i^* s'\}$ can be precomputed by a backwards search in the interactive reachability graph of $\mathcal{P}$.

After k value iteration steps, $p_{\max}^{\mathcal{P}}(s, [0, k])$ equals the probability $\vec{u}_k(s)$.

6.4.2 Maximum step-interval bounded reachability

In this part, we compute $p_{\max}^{\mathcal{P}}(s, [k_a, k_b])$, for interval bounds $0 < k_a < k_b$. As before, the computation proceeds stepwise and produces a sequence of probability vectors $\vec{v}_0, \vec{u}_0, \vec{v}_1, \vec{u}_1, \dots, \vec{v}_{k_b}, \vec{u}_{k_b}$. To allow for lower step bounds $k_a > 0$, we split the value iteration in two parts: In the first $k_b - k_a$ value iteration steps, we proceed as before and compute the probability vectors $\vec{v}_0, \vec{u}_0, \dots, \vec{v}_{k_b - k_a}, \vec{u}_{k_b - k_a}$. Thus, we compute the probabilities $p_{\max}^{\mathcal{P}}(s, [0, k_b - k_a])$ for all $s \in \mathcal{S}$.

The vector $\vec{v}_{k_b - k_a}$ provides the initial probabilities of the second part, which consists of the remaining k_a value iteration steps. For these, we change the way the vectors $\vec{v}_i$ are computed. Instead of Eq. (6.39), we use the defining equation

$$\vec{v}_i(s) = \begin{cases} 0 & \text{if } s \in IS \\ \sum_{s' \in \mathcal{S}} PT(s, s') \cdot \vec{u}_{i-1}(s') & \text{if } s \in PS \end{cases} \quad (6.40)$$

to determine the vectors $\vec{v}_i$. The definition of the vectors $\vec{u}_i$ remains unmodified.

To motivate this definition, note that the value iteration algorithm proceeds in a backwards manner, starting from the goal states. Hence the first $k_b - k_a$ value iteration steps correspond to the specified step interval and we set $\vec{v}_i(s) = 1$ if $s \in G$. However, the remaining k_a steps corresponds to the first k_a transitions that are taken by the IPC. Hence, those steps do not fall into the specified step interval. More specifically, in Eq. (6.40) we do not set $\vec{v}_i(s) = 1$ if $s \in G$, since the fact that a goal state has been hit before k_a steps have occurred does not influence the step-interval bounded reachability probability.

Finally, in order to avoid that the probabilities of interactive states $s \in IS$ erroneously propagate in the vectors $\vec{u}_i(s)$ from the first to the second part, we define $\vec{v}_i(s) = 0$ for all $s \in IS$ (instead of $\vec{v}_i(s) = \vec{u}_{i-1}(s)$ as in the first part). We illustrate this by means of an example.

Example 6.7. We compute $p_{\max}^{\mathcal{P}}(s, [1, 2])$ in the IPC $\mathcal{P}$ in Fig. 6.6 for initial state s_0 and goal state s_3 : In the first part, apply the value iteration to compute $\vec{u}_1$: $\vec{v}_0(s) = 1$ if $s = s_3$ and 0, otherwise. By the backwards closure, $\vec{u}_0 = (1, 0, 0, 1)$. Thus $p_{\max}^{\mathcal{P}}(s_0, [0, 0]) = 1$, as s_0 can reach G by the interactive α -transition. For $\vec{v}_1$, we have $\vec{v}_1(s_0) = \vec{u}_0(s_0) = 1$ and $\vec{v}_1(s_1) = \frac{1}{2}\vec{u}_0(s_3) + \frac{1}{2}\vec{u}_0(s_2) = \frac{1}{2}$. In this way, we obtain $\vec{v}_1 = (1, \frac{1}{2}, \frac{1}{4}, 1)$ and $\vec{u}_1 = (1, \frac{1}{2}, \frac{1}{4}, 1)$. With the probabilities $\vec{u}_1$, the first part ends after $k_b - k_a = 1$ value iteration steps. As $k_a = 1$, one iteration for the lower step bound follows. Here $\vec{v}_2(s_0) = \vec{v}_2(s_3) = 0$ as $s_0, s_3 \in IS$; further $\vec{v}_2(s_1) = \frac{1}{2}\vec{u}_1(s_3) + \frac{1}{2}\vec{u}_1(s_2) = \frac{5}{8}$ and $\vec{v}_2(s_2) = \frac{1}{2}\vec{u}_1(s_2) + \frac{1}{4}\vec{u}_1(s_3) + \frac{1}{4}\vec{u}_1(s_1) = \frac{1}{2}$. Finally, $\vec{u}_2 = (\frac{5}{8}, \frac{5}{8}, \frac{1}{2}, \frac{1}{2})$. Therefore, we obtain that $p_{\max}^{\mathcal{P}}(s_0, [1, 2]) = \vec{u}_2(s_0) = \frac{5}{8}$. $\diamond$

6.4.3 Correctness of the modified value iteration

The following theorem states the correctness of the value iteration algorithm that is informally described in Sec. 6.4.2. More precisely, we prove that the probability $\vec{u}_{k_b}(s)$ is equal to the maximum step-interval bounded reachability probability $p_{\max}^{\mathcal{P}}(s, [k_a, k_b])$.

Although intuitive, the description in Sec. 6.4.2 does not separate the first from the second part of the value iteration algorithm formally. For the correctness proof, we therefore have to extend our notation slightly: Let $[k_a, k_b]$ with $k_a, k_b \in \mathbb{N}$ and $k_a < k_b$ be a step-interval. Then $n = k_b - k_a$ is the number of iteration steps in the first part. Accordingly, the second part consists of the remaining k_a iterations. The idea is to annotate the vectors with the number $n = k_b - k_a$ of value iteration steps that belong to the first part. Therefore, we consider vectors $\vec{v}_0^n, \vec{u}_0^n, \vec{v}_1^n, \vec{u}_1^n, \dots, \vec{v}_{k_b}^n, \vec{u}_{k_b}^n$, where $\vec{v}_0^n, \vec{v}_1^n, \dots, \vec{v}_n^n$ are computed according to Eq. (6.39) and $\vec{v}_{n+1}^n, \vec{v}_{n+2}^n, \dots, \vec{v}_{k_b}^n$ are derived according to Eq. (6.40).

Theorem 6.6 (Maximum value iteration). Let $\mathcal{P} = (\mathcal{S}, \text{Act}, IT, PT, \nu)$ be an IPC, $G \subseteq \mathcal{S}$ a set of goal states, $s \in \mathcal{S}$ a state and $[k_a, k_b]$ with $k_a, k_b \in \mathbb{N}$, $k_a \leq k_b$ a step interval. Further, let $n = k_b - k_a$. For $i = 0, 1, \dots, k_b$, we define the probability vectors $\vec{u}_i^n \in [0, 1]^{\mathcal{S}}$ and $\vec{v}_i^n \in [0, 1]^{\mathcal{S}}$: Initially, $\vec{v}_0^n(s) = 1$ if $s \in G$ and $\vec{v}_0^n(s) = 0$, otherwise. Further, for $i > 0$ we set

$$\vec{v}_i^n(s) = \begin{cases} \sum_{s' \in \mathcal{S}} PT(s, s') \cdot \vec{u}_{i-1}^n(s') & \text{if } s \in PS \wedge (s \notin G \vee i > n) \\ 1 & \text{if } s \in PS \cap G \wedge i \leq n \\ \vec{u}_{i-1}^n(s) & \text{if } s \in IS \wedge i \leq n \\ 0 & \text{if } s \in IS \wedge i > n. \end{cases}$$

For the vectors $\vec{u}_i^n$, we define $\vec{u}_i^n(s) = \max \{ \vec{v}_i^n(s') \mid s \rightsquigarrow_i^* s' \}$ for all $i \leq k_b$. Then it holds

$$p_{\max}^{\mathcal{P}}(s, [k_a, k_b]) = \vec{u}_{k_b}^n(s). \quad (6.41)$$

Observe that if $k_a = 0$, Thm. 6.6 simplifies to the value iteration for the step-bounded reachability computation. Moreover, if $k_a > 0$, the same value iteration is also used in the first $n = k_b - k_a$ steps when maximizing the step-interval bounded reachability for an interval $[k_a, k_b]$. However, in the remaining k_a steps, the vectors $\vec{v}_i^n$ are defined such that visiting a goal state does not imply a probability of 1. We come to the formal proof of Thm. 6.6:

Proof. First, note that by definition of $\rightsquigarrow_i^*$, it holds that $\vec{u}_i^n(s) = \vec{v}_i^n(s)$ for all probabilistic states $s \in PS$. We prove Eq. (6.41) by induction on k_b :

1. For the induction base, assume that $k_b = 0$. As $k_a \leq k_b$, this implies $k_a = 0$. We distinguish between interactive and probabilistic states:
 - (a) The case $s \in PS$: If $s \in G$, then $p_{max}^{\mathcal{P}}(s, [0, 0]) = \Omega(p_{max}^{\mathcal{P}})(s, [0, 0]) = 1 = \vec{v}_0^0(s)$; further, as $s \in PS$ it holds that $\vec{u}_0^0(s) = \vec{v}_0^0(s)$, as desired. With the same reasoning, $p_{max}^{\mathcal{P}}(s, [0, 0]) = \Omega(p_{max}^{\mathcal{P}})(s, [0, 0]) = 0 = \vec{v}_0^0(s) = \vec{u}_0^0(s)$ if $s \notin G$.
 - (b) The case $s \in IS$: As $p_{max}^{\mathcal{P}}$ is the least fixed point of Ω , it holds that $p_{max}^{\mathcal{P}}(s, [0, 0]) = 1$ if $Reach^i(s) \cap G \neq \emptyset$ and $p_{max}^{\mathcal{P}}(s, [0, 0]) = 0$, otherwise. Hence $p_{max}^{\mathcal{P}}(s, [0, 0]) = \max \{ \vec{v}_0^0(s') \mid s \rightsquigarrow_i^* s' \} = \vec{u}_0^0(s)$.
2. In the induction step ($k_b \rightsquigarrow k_b + 1$), we use as induction hypothesis that

$$\forall s \in \mathcal{S}. \forall k_a \leq k_b. p_{max}^{\mathcal{P}}(s, [k_a, k_b]) = \vec{u}_{k_b}^n(s), \text{ where } n = k_b - k_a.$$

The goal is to prove that $p_{max}^{\mathcal{P}}(s, [k_a, k_b + 1]) = \vec{u}_{k_b+1}^{n+1}(s)$ for all $k_a \leq k_b + 1$. We do so by considering two cases, depending on the state s :

- (a) Assume that $s \in PS$. Then $\vec{u}_{k_b+1}^{n+1}(s) = \vec{v}_{k_b+1}^{n+1}(s)$. If $s \in G$ and $k_a = 0$, then $p_{max}^{\mathcal{P}}(s, [0, k_b + 1]) = \Omega(p_{max}^{\mathcal{P}})(s, [0, k_b + 1]) = 1 = \vec{v}_{k_b+1}^{n+1}(s) = \vec{u}_{k_b+1}^{n+1}(s)$. Otherwise $s \notin G$ or $k_a > 0$. If $k_a > 0$ we proceed as follows:

$$\begin{aligned} p_{max}^{\mathcal{P}}(s, [k_a, k_b + 1]) &= \Omega(p_{max}^{\mathcal{P}})(s, [k_a, k_b + 1]) \\ &= \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{max}^{\mathcal{P}}(s', [k_a - 1, k_b]) \\ &\stackrel{i.h.}{=} \sum_{s' \in \mathcal{S}} PT(s, s') \cdot \vec{u}_{k_b}^{n+1}(s') \\ &= \vec{v}_{k_b+1}^{n+1}(s) \quad (* \text{ by def. of } \vec{v}_{k_b+1}^{n+1}(s), \text{ as } n+1 < k_b+1 *) \\ &= \vec{u}_{k_b+1}^{n+1}(s). \quad (* \text{ as } s \in PS *) \end{aligned}$$

If $k_a = 0$ and $s \notin G$, we derive:

$$\begin{aligned} p_{max}^{\mathcal{P}}(s, [0, k_b + 1]) &= \Omega(p_{max}^{\mathcal{P}})(s, [0, k_b + 1]) = \sum_{s' \in \mathcal{S}} PT(s, s') \cdot p_{max}^{\mathcal{P}}(s', [0, k_b]) \\ &\stackrel{i.h.}{=} \sum_{s' \in \mathcal{S}} PT(s, s') \cdot \vec{u}_{k_b}^n(s') = \sum_{s' \in \mathcal{S}} PT(s, s') \cdot \vec{u}_{k_b}^{k_b}(s'). \end{aligned}$$

Observe that by definition, $\vec{v}_i^i = \vec{v}_i^m$ and $\vec{u}_i^i = \vec{u}_i^m$ for all $m \geq i$. Hence:

$$\begin{aligned} p_{\max}^{\mathcal{P}}(s, [0, k_b + 1]) &= \sum_{s' \in \mathcal{S}} PT(s, s') \cdot \vec{u}_{k_b}^{k_b+1}(s') \\ &= \sum_{s' \in \mathcal{S}} PT(s, s') \cdot \vec{u}_{k_b}^{n+1}(s') = \vec{v}_{k_b+1}^{n+1}(s) = \vec{u}_{k_b+1}^{n+1}(s). \end{aligned}$$

(b) The case $s \in IS$: We consider two cases:

- i. The case that $k_a = 0$ and $Reach^i(s) \cap G \neq \emptyset$: If $Reach^i(s) \cap G \neq \emptyset$, then $p_{\max}^{\mathcal{P}}(s, [0, k_b + 1]) = 1$. To see this, choose some state $s' \in Reach^i(s) \cap G$ and apply Ω iteratively until s' is reached.

By definition, we have $\vec{u}_{k_b+1}^{k_b+1}(s) = \max\{\vec{v}_{k_b+1}^{k_b+1}(s'') \mid s \rightsquigarrow_i^* s''\}$. Further, if $s' \in PS$ it holds by definition that $\vec{v}_{k_b+1}^{k_b+1}(s') = 1$. This implies $\vec{u}_{k_b+1}^{k_b+1}(s) = 1$. If $s' \in IS$, we derive $\vec{v}_{k_b+1}^{k_b+1}(s') = \vec{u}_{k_b}^{k_b+1}(s') = \vec{u}_{k_b}^{k_b}(s') = p_{\max}^{\mathcal{P}}(s', [0, k_b]) = 1$ by applying the induction hypothesis to the term $\vec{u}_{k_b}^{k_b}(s')$. Again, $\vec{v}_{k_b+1}^{k_b+1}(s') = 1$ implies that $\vec{u}_{k_b+1}^{k_b+1}(s) = 1$ and we are done.

- ii. The case that $k_a > 0$ or $Reach^i(s) \cap G = \emptyset$: We derive

$$\begin{aligned} p_{\max}^{\mathcal{P}}(s, [k_a, k_b + 1]) &= \Omega(p_{\max}^{\mathcal{P}})(s, [k_a, k_b + 1]) \\ &= \max\{p_{\max}^{\mathcal{P}}(s', [k_a, k_b + 1]) \mid s' \in Reach^i(s)\} \\ &= \max\{p_{\max}^{\mathcal{P}}(s', [k_a, k_b + 1]) \mid s' \in Reach^i(s) \cap PS\} \\ &\quad (* \text{ the case } s \in PS \text{ before } *) \\ &= \max\{\vec{u}_{k_b+1}^{n+1}(s') \mid s' \in Reach^i(s) \cap PS\} \\ &\quad (* \vec{u}_{k_b+1}^{n+1}(s) = \vec{v}_{k_b+1}^{n+1}(s) \text{ for } s \in PS *) \\ &= \max\{\vec{v}_{k_b+1}^{n+1}(s') \mid s' \in Reach^i(s) \cap PS\}. \end{aligned}$$

Now, if $Reach^i(s) \cap G = \emptyset$, it holds that $\max\{\vec{v}_{k_b+1}^{n+1}(s') \mid s' \in Reach^i(s)\} = \vec{v}_{k_b+1}^{n+1}(s'')$ for some $s'' \in Reach^i(s) \cap PS$. Therefore, we obtain $\vec{u}_{k_b+1}^{n+1}(s) = \max\{\vec{v}_{k_b+1}^{n+1}(s') \mid s' \in Reach^i(s) \cap PS\}$, as desired.

Otherwise, $k_a > 0$ and $Reach^i(s) \cap G = \{s_1, s_2, \dots, s_j\}$ for some $j \geq 1$. If $s_i \in G \cap IS$, $k_a > 0$ implies that $k_b + 1 > n + 1$ and hence $\vec{v}_{k_b+1}^{n+1}(s_i) = 0$. Therefore $\max\{\vec{v}_{k_b+1}^{n+1}(s') \mid s' \in Reach^i(s)\} = \vec{v}_{k_b+1}^{n+1}(s'')$ for some $s'' \in Reach^i(s) \cap PS$ and we conclude $\max\{\vec{v}_{k_b+1}^{n+1}(s') \mid s' \in Reach^i(s) \cap PS\} = \vec{u}_{k_b+1}^{n+1}(s)$. $\square$

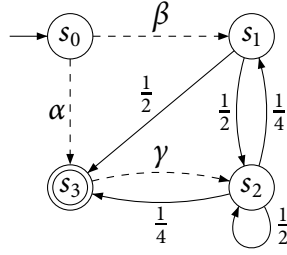


Figure 6.6: Example IPC.

6.4.4 Complexity considerations

Let $\mathcal{M} = (\mathcal{S}, Act, IT, MT, \nu)$ be an IMC, $G \subseteq \mathcal{S}$ a set of goal states and let $I \in \mathcal{Q}$ be a time interval with $b = \sup I$. For the error bound $\varepsilon > 0$, choose k_b such that

$$k_b \cdot \frac{(\lambda\tau)^2}{2} + \lambda\tau \leq \varepsilon.$$

With $\tau = \frac{b}{k_b}$, the smallest such k_b is $k_b = \lceil \frac{\lambda^2 b^2 + 2\lambda b}{2\varepsilon} \rceil$. Then the step duration τ induces the discretized IPC $\mathcal{M}_\tau$. By Thm. 6.5, $p_{max}^{\mathcal{M}}(s_0, I)$ can be approximated (up to ε) by the step-interval bounded reachability $p_{max}^{\mathcal{M}_\tau}(s_0, (k_a, k_b])$ in the discretized IPC $\mathcal{M}_\tau$.

We derive the complexity of our approach: Therefore, let $n = |\mathcal{S}|$ and $m = |IT| + |MT|$ be the number of states and transitions of $\mathcal{M}$, respectively. In the worst case, $\mathcal{M}_\tau$ has n states, and $m + n$ transitions, due to the self-loops which are introduced in the discretization (cf. Def. 6.8 on page 162).

In each value iteration step, the update of the vector $\vec{v}_i$ takes at most $m + n$ time units. When computing $\vec{u}_i$, we assume that the sets $Reach^i(s)$ are precomputed: In the general case, the best theoretical complexity for computing the reflexive transitive closure is in $\mathcal{O}(n^{2.376})$, as given by [CW87]. Let $m^* \subseteq \mathcal{S} \times \mathcal{S}$ denote the reflexive and transitive closure along interactive transitions. As $m^* \subseteq \mathcal{S} \times \mathcal{S}$, the number of transitions in m^* is bounded by n^2 . Hence, with an appropriate precomputation of m^* , updating $\vec{u}_i$ takes time $\mathcal{O}(n^2)$.

Altogether, for $k_b = \lceil \frac{\lambda^2 b^2 + 2\lambda b}{2\varepsilon} \rceil$ value iteration steps, the worst case time complexity of our approach is $n^{2.376} + (m + n + n^2) \cdot (\lambda b) \cdot (\lambda b + 2) / (2\varepsilon) \in \mathcal{O}(n^{2.376} + (m + n^2) \cdot (\lambda b)^2 / \varepsilon)$.

6.5 Model checking the continuous stochastic logic

The crucial point for model checking CSL is to compute the maximum and minimum probability to visit a set of goal states in some time interval I . In this section, we therefore apply the results from Sec. 6.3 and reduce the CSL model checking problem to the time-interval bounded reachability computation. However, this only works for a slightly restricted subset of the logic CSL. We address this restriction in detail in Sec. 6.5.2.

Model checking CSL relies on state labellings; hence, we introduce a finite set $AP = \{a, b, c, \dots\}$ of *atomic propositions* and consider *state labeled* IMCs, where a *state labeling function* $L : \mathcal{S} \rightarrow 2^{AP}$ assigns to each state the set of atomic propositions that hold in that state.

6.5.1 Syntax and semantics of CSL

The continuous stochastic logic (CSL) [BHHK03, CDHS06] is devised for specifying quantitative properties of continuous-time Markov chains. In the first part of this section, we therefore extend its semantics to the nondeterministic setting. However, we omit the steady-state operator from classical CSL [BHHK03], as a steady-state generally does not exist in controlled Markov chains or IMCs.

Definition 6.9 (CSL syntax). For $a \in AP$, $p \in [0,1]$, $I \subseteq \mathcal{Q}$ an interval and $\triangleq \in \{<, \leq, \geq, >\}$, the syntax of CSL state and CSL path formulas is defined by the following grammar rules:

$$\Phi ::= a \mid \neg\Phi \mid \Phi \wedge \Phi \mid \mathcal{P}_{\triangleq p}(\varphi) \quad \text{and} \quad \varphi ::= X^I\Phi \mid \Phi \mathcal{U}^I \Phi.$$

Intuitively, a path $\pi \in \text{Paths}^\omega$ satisfies the next formula $X^I\Phi$ (denoted $\pi \models X^I\Phi$) if the first transition on π occurs in time-interval I and leads to a successor state in $\text{Sat}(\Phi)$. Similarly, π satisfies the until formula $\Phi \mathcal{U}^I \Psi$ if a state in $\text{Sat}(\Psi)$ is visited at some time point $t \in I$ and before that, all states satisfy state formula Φ .

Intuitively, the semantics of the probabilistic state formula $\mathcal{P}_{\triangleq p}(\varphi)$ is defined such that $s \models \mathcal{P}_{\triangleq p}(\varphi)$ holds if the probability of the set of paths that start in state s and that satisfy the CSL path formula φ meets the bound specified by $\triangleq p$.

Definition 6.10 (CSL semantics). Let $\mathcal{M} = (\mathcal{S}, \text{Act}, IT, MT, AP, L, \nu)$ be a state labeled IMC, $s \in \mathcal{S}$ a state, $a \in AP$ an atomic proposition, $I \in \mathcal{Q}$ a time interval, $\triangleq \in \{<, \leq, \geq, >\}$ a comparison operator and $\pi \in \text{Paths}^\omega$ an infinite path.

For CSL state formulas, we define:

$$\begin{aligned} s \models a & \iff a \in L(s) \\ s \models \neg\Phi & \iff s \not\models \Phi \\ s \models \Phi \wedge \Psi & \iff s \models \Phi \wedge s \models \Psi \\ s \models \mathcal{P}_{\triangleq p}(\varphi) & \iff \forall D \in GM. Pr_{\nu, D}^\omega \{ \pi \in \text{Paths}^\omega \mid \pi \models \varphi \} \triangleq p. \end{aligned}$$

The semantics for path formulas is defined as follows:

$$\begin{aligned} \pi \models X^I\Phi & \iff \pi[1] \models \Phi \wedge \delta(\pi, 0) \in I \\ \pi \models \Phi \mathcal{U}^I \Psi & \iff \exists t \in I. \exists s \in \pi@t. s \models \Psi \wedge \forall s' \in \text{Pref}(\pi@t, s). s' \models \Phi \\ & \quad \wedge \forall t' \in [0, t). \forall s'' \in \pi@t'. s'' \models \Phi. \end{aligned}$$

Some remarks are in order: First, the semantics of the until path formula is slightly more involved compared to the original definition in [BHHK03]: Due to interactive transitions that execute instantaneously, an IMC may traverse a (finite or infinite) sequence of states in 0 time units. Therefore $\pi \models \Phi \mathcal{U}^I \Psi$ is defined such that it holds if there exists a state sequence $\pi@t$ that is traversed at some time $t \in I$ and on $\pi@t$, a Ψ -state is visited. Moreover, for $\pi \models \Phi \mathcal{U}^I \Psi$ to be satisfied, all previous states on $\pi@t$ and all states visited at times $t' < t$ must satisfy Φ .

Second, to decide the probabilistic CSL state formula $\mathcal{P}_{\trianglelefteq p}(\varphi)$, we need to distinguish two cases: If $\trianglelefteq = <$ or $\trianglelefteq = \leq$, it suffices to verify that $p_{\max}^{\mathcal{M}}(s, \varphi) \trianglelefteq p$. Reversely, if $\trianglelefteq = >$ or $\trianglelefteq = \geq$, we need to compute the infimum $p_{\min}^{\mathcal{M}}(s, \varphi)$ and to check whether $p_{\min}^{\mathcal{M}}(s, \varphi) \trianglelefteq p$.

6.5.2 Model checking algorithm for CSL

The model checking algorithm that we present in this section works only for a subset of all CSL formulas. More precisely, we restrict to path formulas $\Phi \mathcal{U}^I \Psi$ where $\Psi \Rightarrow \Phi$ if $\inf I > 0$. Note however, that albeit this restriction we preserve most of the expressivity of CSL: For example, the CSL operator $\diamond^I \Phi$ can still be derived, as $\diamond^I \Phi \equiv \text{tt } \mathcal{U}^I \Phi$ for CSL state formula Φ . Moreover, it does not apply to time-bounded reachability objectives, i.e. to the case where $\inf I = 0$. Hence, the restriction does hardly ever hamper the practical applicability of our approach. Intuitively, its consequence can be stated as follows: If we consider interval-bounded until formulas with $\inf I > 0$, we require that on any path π which satisfies the formula $\Phi \mathcal{U}^I \Psi$, the validity of Φ needs to be resolved by a state which satisfies Ψ and Φ .

To model check an IMC with respect to a state formula Φ from this subset of CSL, we successively consider the state subformulas Ψ of Φ and calculate the sets $\text{Sat}(\Psi) = \{s \in \mathcal{S} \mid s \models \Psi\}$. For atomic propositions, conjunction and negation, this is easy, as

$$\begin{aligned} \text{Sat}(a) &= \{s \in \mathcal{S} \mid a \in L(s)\}, \\ \text{Sat}(\neg \Psi) &= \mathcal{S} \setminus \text{Sat}(\Psi) \text{ and} \\ \text{Sat}(\Psi_1 \wedge \Psi_2) &= \text{Sat}(\Psi_1) \cap \text{Sat}(\Psi_2). \end{aligned}$$

In the remainder of this section, we therefore discuss the probabilistic operator $\mathcal{P}_{\trianglelefteq p}(\varphi)$ for next and until formulas. To decide $\text{Sat}(\mathcal{P}_{\trianglelefteq p}(\varphi))$, it suffices to maximize or minimize the probability $Pr_{v_s, D}^{\omega}(\{\pi \in \text{Paths}^{\omega} \mid \pi \models \varphi\})$ with respect to all schedulers $D \in \text{GM}$. Accordingly, we define

$$\begin{aligned} p_{\max}^{\mathcal{M}}(s, \varphi) &= \sup_{D \in \text{GM}} Pr_{v_s, D}^{\omega}(\{\pi \in \text{Paths}^{\omega} \mid \pi \models \varphi\}) \text{ and} \\ p_{\min}^{\mathcal{M}}(s, \varphi) &= \inf_{D \in \text{GM}} Pr_{v_s, D}^{\omega}(\{\pi \in \text{Paths}^{\omega} \mid \pi \models \varphi\}). \end{aligned}$$

As done throughout this chapter, we only consider the details for maximizing the probability $Pr_{v_s, D}^{\omega}(\{\pi \in \text{Paths}^{\omega} \mid \pi \models \varphi\})$ and leave out most of the details for computing the minimum probabilities, which can be done similarly.

The next formula

Computing $p_{\max}^{\mathcal{M}}(s, \mathcal{X}^I \Phi)$ is straightforward: We proceed inductively on the structure of the formula and assume that $Sat(\Phi)$ is already computed. Then we distinguish two cases, depending on whether state s is a Markovian or an interactive state:

- (a) If $s \in MS$ is a Markovian state, no nondeterminism occurs and we derive $p_{\max}^{\mathcal{M}}(s, \mathcal{X}^I \Phi)$ as done for CTMCs in [BHHK03]: Let $a = \inf I$ and $b = \sup I$; then

$$\begin{aligned} p_{\max}^{\mathcal{M}}(s, \mathcal{X}^I \Phi) &= \int_a^b E(s) e^{-E(s)t} \cdot \sum_{s' \in Sat(\Phi)} \mathbf{P}(s, s') dt \\ &= \mathbf{P}(s, Sat(\Phi)) \cdot (e^{-E(s)a} - e^{-E(s)b}), \end{aligned}$$

where $\mathbf{P}(s, Sat(\Phi)) = \sum_{s' \in Sat(\Phi)} \mathbf{P}(s, s')$ is the probability to move to a successor state $s' \in Sat(\Phi)$ when leaving state s .

- (b) If $s \in IS$ is an interactive state, the probability $p_{\max}^{\mathcal{M}}(s, \mathcal{X}^I \Phi)$ depends on the interval I : If $0 \in I$ and $post^i(s) \cap Sat(\Phi) \neq \emptyset$, then $p_{\max}^{\mathcal{M}}(s, \mathcal{X}^I \Phi) = 1$; otherwise it holds that $p_{\max}^{\mathcal{M}}(s, \mathcal{X}^I \Phi) = 0$.

The until formula

Computing $p_{\max}^{\mathcal{M}}(s, \Phi \mathcal{U}^I \Psi)$ is more complex: Let $\varphi = \Phi \mathcal{U}^I \Psi$ be a time-interval bounded until path formula with $I \in \mathcal{Q}$ and the restriction that $\Psi \Rightarrow \Phi$ if $\inf I > 0$. As we will see, this technical restriction is essential for the correctness proof given in Thm. 6.7 below.

As the computation proceeds inductively along the structure of the formula, we may assume that $Sat(\Phi)$ and $Sat(\Psi)$ are already computed. Note that if $\inf I > 0$, the restriction to until formulas $\Phi \mathcal{U}^I \Psi$ where $\Psi \Rightarrow \Phi$ directly implies that $Sat(\Psi) \subseteq Sat(\Phi)$.

We reduce the problem of computing $p_{\max}^{\mathcal{M}}(s, \varphi)$ and $p_{\min}^{\mathcal{M}}(s, \varphi)$ to the maximum and minimum interval-bounded reachability problem, respectively. Therefore, define the set

$$\mathcal{S}_{=0}^{\varphi} = \{s \in \mathcal{S} \mid s \models \neg \Phi \wedge \neg \Psi\}.$$

of absorbing states: A Markovian state $s \in MS$ is called *absorbing* iff $\mathbf{R}(s, \lambda, s) > 0$ and $\mathbf{R}(s, \lambda, s') = 0$ for all $s' \neq s$; hence, absorbing states are states with a single Markovian self loop. Similar to the approach taken for model checking CTMCs and MDPs [BHHK03, BdA95], we make all states $s \in \mathcal{S}_{=0}^{\varphi}$ absorbing by replacing all their outgoing transitions by a single Markovian self loop $(s, 1, s)$.

Intuitively this is justified as follows: Let $Paths^{\omega}(s)$ denote the set of all infinite paths that start in state s . Then the probability of the set $\{\pi \in Paths^{\omega}(s) \mid \pi \models \Phi \mathcal{U}^I \Psi\}$ is 0 for states $s \in \mathcal{S}_{=0}^{\varphi}$: If a state $s \in \mathcal{S}_{=0}^{\varphi}$ is visited, it violates Φ and Ψ . But all paths that start in a $(\neg \Phi \wedge \neg \Psi)$ -state violate the until formula $\Phi \mathcal{U}^I \Psi$. Hence, making those states absorbing does not alter the probabilities $p_{\max}^{\mathcal{M}}(s, \varphi)$ and $p_{\min}^{\mathcal{M}}(s, \varphi)$.

Theorem 6.7 (Time-bounded until). Let $\mathcal{M} = (\mathcal{S}, \text{Act}, IT, MT, AP, L, \nu)$ be a state-labeled IMC, $\varphi = \Phi \mathcal{U}^I \Psi$ a CSL path formula with $I \in \mathcal{Q}$ a time-interval and Φ, Ψ state formulas such that $\Psi \Rightarrow \Phi$ if $\inf I > 0$. Further, let $G = \text{Sat}(\Psi)$ be the set of goal states and assume that all states $s \in \mathcal{S}_{=0}^\varphi$ are made absorbing. Then it holds for all $s \in \mathcal{S}$:

$$p_{\max}^{\mathcal{M}}(s, \Phi \mathcal{U}^I \Psi) = p_{\max}^{\mathcal{M}}(s, I) \quad \text{and} \quad p_{\min}^{\mathcal{M}}(s, \Phi \mathcal{U}^I \Psi) = p_{\min}^{\mathcal{M}}(s, I).$$

Proof. It suffices to prove that for all paths $\pi \in \text{Paths}^\omega(s)$, it holds:

$$\pi \models \Phi \mathcal{U}^I \Psi \iff \pi \models \Diamond^I(\text{Sat}(\Psi)).$$

We show the two directions separately:

“ $\Rightarrow$ ” First, assume that $\pi \models \Phi \mathcal{U}^I \Psi$. Let $\pi \in \text{Paths}^\omega$. By the semantics of the until formula, we have:

$$\begin{aligned} \pi \models \Phi \mathcal{U}^I \Psi \iff \exists t \in I. \exists s \in \pi@t. s \models \Psi \wedge \forall s' \in \text{Pref}(\pi@t, s). s' \models \Phi \\ \wedge \forall t' \in [0, t). \forall s'' \in \pi@t'. s'' \models \Phi. \end{aligned}$$

Thus, for all $t' \in [0, t)$ and $s'' \in \pi@t'$, we have $s'' \models \Phi$ implying $s'' \notin \mathcal{S}_{=0}^\varphi$. Moreover, for all $s' \in \text{Pref}(\pi@t, s)$ it holds that $s' \models \Phi$, implying that $s' \notin \mathcal{S}_{=0}^\varphi$. Hence, none of the states is made absorbing. Let n be the index of π such that $\pi[n] = s$. Then we have that $\pi[n] = s \models \Psi$, implying that $\pi \models \Diamond^I(\text{Sat}(\Psi))$.

“ $\Leftarrow$ ” Now let π be such that $\pi \models \Diamond^I(\text{Sat}(\Psi))$. Thus, there exists $t \in I$ such that

$$\exists s \in \pi@t. s \models \Psi. \tag{6.42}$$

Choose the minimal $t \in I$ such that Eq. (6.42) holds. Moreover, for this t , choose the first occurrence of a state $s \in \text{Sat}(\Psi)$ in $\pi@t$. Now let $n \in \mathbb{N}$ be its position on π and consider all states $\pi[k]$ with $k < n$. Since $\pi[k]$ can reach $\pi[n]$, we have $\pi[k] \notin \mathcal{S}_{=0}^\varphi$. If $\inf I = 0$, the minimality of t implies that $\pi[n]$ is the first occurrence of a Ψ -state on π and therefore, that $\pi[k] \models \Phi$ for all $k < n$. If $\inf I > 0$, we know that $\pi[k] \models \Phi$ or $\pi[k] \models \Psi$ for all $k < n$. In the latter case, the restriction to until formulas where $\Psi \Rightarrow \Phi$ implies that $\pi[k] \models \Phi$. Hence, in both cases it holds that $\pi[k] \models \Phi$ for all $k < n$, proving that $\pi \models \Phi \mathcal{U}^I \Psi$. $\square$

Theorem 6.7 reduces the problem to compute $p_{\max}^{\mathcal{M}}(s, \Phi \mathcal{U}^I \Psi)$ and $p_{\min}^{\mathcal{M}}(s, \Phi \mathcal{U}^I \Psi)$ for interval bounded until formulas to the problem of computing the interval bounded reachability probabilities $p_{\max}^{\mathcal{M}}(s, I)$ and $p_{\min}^{\mathcal{M}}(s, I)$ with respect to the set of goal states $G = \text{Sat}(\Psi)$. The latter can be computed efficiently by the discretization approach introduced in Sec. 6.3.

Remark 6.1 (The restricted until formulas). *Theorem 6.7 relies on the assumption that $\Psi \Rightarrow \Phi$ for intervals I with $\inf I > 0$. Without this restriction, the direction from right to left in the proof of Thm. 6.7 does not hold. To see this, assume that $\Psi \not\Rightarrow \Phi$ and that $\inf I > 0$. If on a path π , a $(\Psi \wedge \neg\Phi)$ -state is visited at time $t < \inf I$, say on position k , then $\pi \not\models \Phi \mathcal{U}^I \Psi$. However, $\pi[k] \notin \mathcal{S}_{=0}^\varphi$, as $\pi[k] \models \Psi$. Hence, state $\pi[k]$ is not made absorbing. Therefore, the path π is erroneously included in the computation of the reachability probability $\Diamond^I G$.*

Complexity of CSL model checking

The complexity of the CSL model checking approach is clearly dominated by the interval-bounded reachability computation: For CSL state-formula Φ , let $|\Phi|$ be the number of state subformulas of Φ . In the worst case, the interval bounded reachability probability is computed $|\Phi|$ times. Using the complexity of the value iteration algorithm (cf. Sec. 6.4.4), the model checking problem has time complexity $\mathcal{O}(|\Phi| \cdot (n^{2.376} + (m + n^2) \cdot (\lambda b)^2 / \epsilon))$.

6.6 Experimental results

We consider the IMC in Fig. 6.7(a), where $\text{Erl}(30, 10)$ denotes a transition with an Erlang (k, λ) distributed delay: This corresponds to $k = 30$ consecutive Markovian transitions each of which has rate λ . The mean time to move from s_2 to the goal s_4 is $\frac{k}{\lambda} = 3$ with a variance of $\frac{k}{\lambda^2} = \frac{3}{10}$. Hence, with very high probability we move from state s_2 to state s_4 after approximately 3 time units. The decision that maximizes the probability to reach s_4 in time interval $[0, b]$ in state s_1 depends on the sojourn in state s_0 . Fig. 6.7(b) depicts the computed maxima for time-dependent schedulers and the upper part of Tab. 6.7(c) lists some performance measurements.

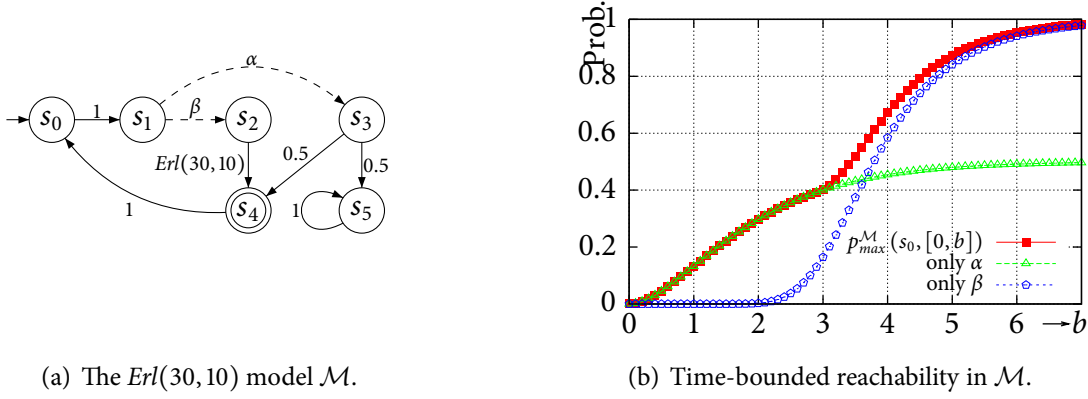
If $AP = \{g\}$ and s_4 is the only state labeled with g , we can verify the CSL formula $\Phi = \mathcal{P}_{\geq 0.5}(\Diamond^{[3,4]} g)$ by computing $p_{\max}^M(s_0, [3, 4])$ with the modified value iteration. The result $p_{\max}^M(s_0, [3, 4]) = 0.6057$ meets the bound ≥ 0.5 in Φ , implying that $s_0 \models \Phi$.

All measurements were carried out on a 2.2GHz Xeon CPU with 16GB RAM.

6.7 Interval bounded reachability in early CTMDPs

In this section, we apply the time-interval bounded reachability analysis that we have developed for closed IMCs to also solve the open problem of computing time-interval bounded reachability probabilities in early CTMDPs. Note the difference compared to Chapter 5, where we considered *locally uniform late CTMDPs*. In this section, we consider *arbitrary early CTMDPs* and transform them into an equivalent *alternating* IMC which is then subject to the analysis techniques developed so far.

As a model that incorporates continuous-time and nondeterminism, IMCs strictly separate interactive from Markovian transitions, whereas CTMDPs combine non-deterministic choices with exponential delays. However, CTMDPs can be considered as the



problem	states	ε	λ	b	prob.	time
$Erl(30, 10)$	35	10^{-3}	10	4	0.672	50s
$Erl(30, 10)$	35	10^{-3}	10	7	0.983	70s
$Erl(30, 10)$	35	10^{-4}	10	4	0.6718	268s

(c) Computation times for different parameters.

Figure 6.7: Experimental results for $Erl(30, 10)$.

subclass of *strictly alternating* IMCs [HJ07]. Briefly, an IMC is strictly alternating if all successor states of interactive states are Markovian states, and all successor states of Markovian states are interactive states. With this definition, an early CTMDP can be considered as a strictly alternating (and closed) IMC in which the Markovian and interactive states are entangled.

In order to reduce the model checking problem for early CTMDPs to the corresponding problem for IMCs, we define the *induced IMC* $\mathcal{M}(\mathcal{C})$ for an early CTMDP $\mathcal{C}$ as follows:

Definition 6.11 (Induced IMC of a CTMDP). Let $\mathcal{C} = (\mathcal{S}, Act, \mathbf{R}, v)$ be a CTMDP. Its induced IMC $\mathcal{M}(\mathcal{C})$ is the tuple $(\mathcal{S}', Act, IT, MT, v')$ such that

$$\begin{aligned}
 \mathcal{S}' &= \mathcal{S} \sqcup \{s^\alpha \mid s \in \mathcal{S} \wedge \alpha \in Act(s)\}, \\
 IT &= \{(s, \alpha, s^\alpha) \mid s \in \mathcal{S} \wedge \alpha \in Act(s)\} \text{ and} \\
 MT &= \{(s^\alpha, \mathbf{R}(s, \alpha, s'), s') \mid s' \in \mathcal{S} \wedge \mathbf{R}(s, \alpha, s') > 0\}.
 \end{aligned}$$

Further, $v'(s) = v(s)$ if $s \in \mathcal{S}$ and $v'(s) = 0$, otherwise.

Example 6.8. Consider the early CTMDP in Fig. 6.8(b) on page 201. Applying Def. 6.11 yields its induced IMC which is depicted in Fig. 6.8. $\diamond$

For model checking purposes, it is useful to extend Def. 6.11 to *state labeled* CTMDPs: A state labeled CTMDP is augmented by a set AP of atomic propositions and a state labeling function $L : \mathcal{S} \rightarrow 2^{AP}$. We define the labeling L' of $\mathcal{C}$'s induced IMC such that the labeling of each interactive state and its corresponding Markovian successor states coincide. Formally: $L'(s) = L(s)$ and $L'(s^\alpha) = L(s)$ for all $s \in \mathcal{S}$ and $\alpha \in Act(s)$.

By definition, the induced IMC of a CTMDP $\mathcal{C}$ is *strictly alternating*: Each state $s \in \mathcal{S}$ in $\mathcal{C}$ becomes an interactive state in the induced IMC which mimics the CTMDP's nondeterministic choices: For each action $\alpha \in Act(s)$, an internal transition leads from interactive state s to a newly introduced Markovian state s^α which represents the race between the exponential delays that lead to the successor states of s in the underlying early CTMDP under action α .

To formally establish the relation between an early CTMDP $\mathcal{C}$ and its induced strictly alternating IMC $\mathcal{M}$, we first observe a correspondence between paths in $\mathcal{M}$ and paths in $\mathcal{C}$: Therefore, let $sep : Paths(\mathcal{C}) \rightarrow Paths(\mathcal{M})$ be such that it *separates* the scheduler choices and the Markovian sojourn times on a path $\pi \in Paths(\mathcal{C})$. Formally:

$$sep\left(s_0 \xrightarrow{\alpha_0, t_0} s_1 \xrightarrow{\alpha_1, t_1} \dots\right) = s_0 \xrightarrow{\alpha_0, 0} s_0^{\alpha_0} \xrightarrow{\perp, t_0} s_1 \xrightarrow{\alpha_1, 0} s_1^{\alpha_1} \xrightarrow{\perp, t_1} \dots.$$

Reversely, we *collapse* paths in $\mathcal{M}$ to obtain the corresponding path in $\mathcal{C}$:

$$col\left(s_0 \xrightarrow{\alpha_0, 0} s_0^{\alpha_0} \xrightarrow{\perp, t_0} s_1 \xrightarrow{\alpha_1, 0} s_1^{\alpha_1} \xrightarrow{\perp, t_1} \dots\right) = s_0 \xrightarrow{\alpha_0, t_0} s_1 \xrightarrow{\alpha_1, t_1} \dots.$$

For infinite paths, we thus have a one-to-one correspondence between infinite paths in $\mathcal{C}$ and infinite paths in $\mathcal{M}$. Moreover, each finite path $\pi \in Paths(\mathcal{C})$ induces a unique path $\bar{\pi} \in Paths(\mathcal{M})$ of length $|\bar{\pi}| = 2|\pi|$; reversely, each path $\bar{\pi} \in Paths(\mathcal{M})$ that starts and ends in an interactive state maps back to a unique path $col(\bar{\pi})$ in the underlying early CTMDP. For the following discussion, we extend the definitions of the functions sep and col to sets of paths in the natural way.

6.7.1 Scheduler correspondence

We aim at establishing a correspondence between sets of paths in the early CTMDP $\mathcal{C}$ and its induced IMC $\mathcal{M}$. Each path $\pi \in Paths(\mathcal{C})$ corresponds to the path $sep(\pi)$ in $\mathcal{M}$, which starts and ends in an interactive state. Further, the initial distribution in $\mathcal{C}$'s induced IMC $\mathcal{M}$ assigns probability 0 to each path in $\mathcal{M}$ that starts in a Markovian state. Hence, such paths can safely be ignored in the remainder of this section.

The above observation allows us to establish a close correspondence between the schedulers in $\mathcal{C}$ and $\mathcal{M}$: Let $D^{\mathcal{C}} \in GM(\mathcal{C})$ be an early scheduler in $\mathcal{C}$ and $\bar{\pi} \in Paths^*(\mathcal{M})$ a path in $\mathcal{M}$. We define the scheduler $D^{\mathcal{M}} \in GM(\mathcal{M})$ such that

$$D^{\mathcal{M}}(\bar{\pi}) = \begin{cases} D^{\mathcal{C}}(col(\bar{\pi})) & \text{if } \bar{\pi} \downarrow \in IS \wedge \bar{\pi}[0] \in IS \\ ? & \text{if } \bar{\pi} \downarrow \in IS \wedge \bar{\pi}[0] \in MS \\ \perp & \text{if } \bar{\pi} \downarrow \in MS, \end{cases} \quad (6.43)$$

where the scheduler decisions taken on paths $\bar{\pi}$ that start in a Markovian state can be chosen arbitrary (as long as $D^{\mathcal{M}}$ remains measurable), as in our setting, the set of such paths has measure 0 anyways. Hence, for our purposes we can identify all schedulers which differ only for the case that $\bar{\pi} \downarrow \in IS$ and $\bar{\pi}[0] \in MS$.

Reversely, if $D^{\mathcal{M}} \in GM(\mathcal{M})$ is a scheduler in the strictly alternating IMC $\mathcal{M}$, it corresponds to a unique early scheduler $D^{\mathcal{C}} \in GM(\mathcal{C})$, which is defined for all $\pi \in Paths^*(\mathcal{C})$ such that $D^{\mathcal{C}}(\pi) = D^{\mathcal{M}}(sep(\pi))$.

Hence, there exists a one-to-one correspondence between schedulers in $\mathcal{C}$ and $\mathcal{M}$.

6.7.2 Measure correspondence

We first prove that the probability measure that is induced for a set of paths $\Pi \in Paths^{\omega}(\mathcal{C})$ by a scheduler $D^{\mathcal{C}} \in GM(\mathcal{C})$ in the early CTMDP $\mathcal{C}$ equals the probability of $sep(\Pi)$ under the corresponding scheduler $D^{\mathcal{M}}$ in the induced IMC $\mathcal{M}$:

Lemma 6.12 (Measure correspondence). *Let $\mathcal{C} = (\mathcal{S}, Act, \mathbf{R}, \nu)$ be a CTMDP and $\mathcal{M} = (\mathcal{S}', Act, IT, MT, \nu')$ be its induced IMC. Further, let $D^{\mathcal{C}} \in GM(\mathcal{C})$ be a scheduler in $\mathcal{C}$ and let $D^{\mathcal{M}} \in GM(\mathcal{M})$ be the corresponding scheduler in $\mathcal{M}$ as defined in Eq. (6.43). For all $s \in \mathcal{S}$ and $\Pi \in \mathfrak{F}_{Paths^{\omega}(\mathcal{C})}$ it holds that*

$$Pr_{\nu_s, D^{\mathcal{C}}}^{\omega}(\Pi) = Pr_{\nu'_s, D^{\mathcal{M}}}^{\omega}(sep(\Pi)).$$

Proof. The proof is along the same lines as in Lemma 4.4 in Sec. 4.2.2: We first prove the claim for measurable rectangles: Let $B = S_0 \times A_0 \times T_0 \times S_1 \times \dots \times S_n \in \mathfrak{F}_{Paths^n(\mathcal{C})}$ be a measurable rectangle in $\mathcal{C}$. Then $\bar{B} = sep(B) = S_0 \times A_0 \times S_0^{A_0} \times T_0 \times S_1 \times A_1 \times S_1^{A_1} \times T_1 \times S_2 \times \dots \times S_n$, where $S_i^{A_i} = \{s^{\alpha} \mid s \in S_i \wedge \alpha \in A_i\}$ for $0 \leq i < n$. We proceed by induction on n and prove for all measurable rectangles $B \in \mathfrak{F}_{Paths^n(\mathcal{C})}$:

$$Pr_{\nu, D^{\mathcal{C}}}^n(B) = Pr_{\nu', D^{\mathcal{M}}}^{2n}(sep(B)). \quad (6.44)$$

In the induction base, $B = S_0$ and $\bar{B} = S_0$. Hence, $Pr_{\nu, D^{\mathcal{C}}}^0(B) = \sum_{s \in S_0} \nu(s) = \sum_{s \in S_0} \nu'(s) = Pr_{\nu', D^{\mathcal{M}}}^0(\bar{B})$. In the induction step, let $I = S_0 \times A_0 \times T_0$ be a set of initial path prefixes (cf. Lemma 3.16) in $\mathcal{C}$ which extend the measurable rectangle $B \in \mathfrak{F}_{Paths^n(\mathcal{C})}$ to a measurable rectangle $I \times B \in \mathfrak{F}_{Paths^{n+1}(\mathcal{C})}$ of length $n+1$. With $i = (s, \alpha, t)$ ranging over I , we derive

$$\begin{aligned} Pr_{\nu, D^{\mathcal{C}}}^{n+1}(I \times B) &= \int_I Pr_{\nu_i, D_i^{\mathcal{C}}}^n(B) \mu_{\nu, D^{\mathcal{C}}}^1(di) \\ &= \int_I Pr_{\nu_i, D_i^{\mathcal{M}}}^{2n}(\bar{B}) \mu_{\nu, D^{\mathcal{C}}}^1(di), \quad (* \text{ by the ind. hyp.} *) \end{aligned}$$

where $\mu_{\nu, D^{\mathcal{C}}}^k$ is the probability measure on initial path prefixes as defined in Sec. 3.3.2 on page 82. Now, if $i = (s, \alpha, t) \in I$ is an initial path prefix in $\mathcal{C}$, let $\bar{i} = (s, \alpha, 0, s^{\alpha}, \perp, t)$ be the

corresponding two-step initial path prefix in $\mathcal{M}$. Then $v_i(s') = \mathbf{P}^C(s, \alpha, s') = \frac{\mathbf{R}(s, \alpha, s')}{E(s, \alpha)} = \mathbf{P}^{\mathcal{M}}(s^\alpha, s') = v_{\bar{i}}(s')$, where $\mathbf{P}^C(s, \alpha, s')$ denotes the branching probability from s to s' under action α in $\mathcal{C}$ and $\mathbf{P}^{\mathcal{M}}(s^\alpha, s')$ denotes the corresponding probability from state s^α to state s' in $\mathcal{M}$.

Moreover it holds that $D_i^{\mathcal{M}}(\text{sep}(\pi)) = D^{\mathcal{M}}(\bar{i} \circ \text{sep}(\pi)) = D^C(i \circ \pi) = D_i^C(\pi) = D_i^{\mathcal{M}}(\text{sep}(\pi))$ for all $\pi \in \text{Paths}^*(C)$. Hence:

$$\begin{aligned}
Pr_{v, D^C}^{n+1}(I \times B) &= \int_I Pr_{v_{\bar{i}}, D_{\bar{i}}^{\mathcal{M}}}^{2n}(\bar{B}) \mu_{v, D^C}^1(di) \\
&= \sum_{s \in S_0} v(s) \sum_{\alpha \in A_0} D^C(s, \alpha) \int_{T_0} Pr_{v_{\bar{i}}, D_{\bar{i}}^{\mathcal{M}}}^{2n}(\bar{B}) \eta_{E(s, \alpha)}(dt) \quad (* \text{ def. of } \mu_{v, D^C}^1 *) \\
&= \sum_{s \in S_0} v(s) \sum_{\alpha \in A_0} D^{\mathcal{M}}(s, \alpha) \int_{T_0} Pr_{v_{\bar{i}}, D_{\bar{i}}^{\mathcal{M}}}^{2n}(\bar{B}) \eta_{E(s^\alpha)}(dt) \quad (* \text{ succ}(\alpha) = s^\alpha *) \\
&= \int_{\bar{I}} Pr_{v_{\bar{i}}, D_{\bar{i}}^{\mathcal{M}}}^{2n}(\bar{B}) \mu_{v, D^{\mathcal{M}}}^2(d\bar{i}) \quad (* \text{ def. of } \mu_{v, D^{\mathcal{M}}}^2 *) \\
&= Pr_{v, D^{\mathcal{M}}}^{2n+2}(\bar{I} \times \bar{B}) = Pr_{v, D^{\mathcal{M}}}^{2(n+1)}(\text{sep}(I \times B)).
\end{aligned}$$

Thus Eq. (6.44) holds for all measurable rectangles. To prove that this result extends to arbitrary measurable sets of paths $\Pi \in \mathfrak{F}_{\text{Paths}^\omega}$, it suffices to prove (6.44) for any measurable base $B \in \mathfrak{F}_{\text{Paths}^n}$. Therefore, let $\mathfrak{G}_{\text{Paths}^n(C)}$ denote the set of all finite disjoint unions of measurable rectangles, which forms a field by Lemma 2.10 (see page 43). Then Eq. (6.44) directly extends to $\mathfrak{G}_{\text{Paths}^n(C)}$: Let $B = \uplus_{i=0}^k B_i$ with all B_i being pairwise disjoint measurable rectangles in $\mathfrak{F}_{\text{Paths}^n(C)}$. Then $Pr_{v, D^C}^n(B) = Pr_{v, D^C}^n(\uplus_{i=0}^k B_i) = \sum_{i=0}^k Pr_{v, D^C}^n(B_i) = \sum_{i=0}^k Pr_{v', D^{\mathcal{M}}}^{2n}(\text{sep}(B_i)) = Pr_{v', D^{\mathcal{M}}}^{2n}(\uplus_{i=0}^k \text{sep}(B_i)) = Pr_{v', D^{\mathcal{M}}}^{2n}(\text{sep}(B))$.

Now, define

$$\mathfrak{C} = \{B \in \mathfrak{F}_{\text{Paths}^n} \mid Pr_{v, D^C}^n(B) = Pr_{v', D^{\mathcal{M}}}^{2n}(\text{sep}(B))\}.$$

Then $\mathfrak{C}$ is a *monotone class*, i.e. for all $B_i \uparrow B$ and $B_i \downarrow B$, it holds $B \in \mathfrak{C}$: Here, we only give the proof for increasing sequences. Let $B_i \uparrow B$. As σ -fields are closed under increasing sequences, we obtain $B \in \mathfrak{F}_{\text{Paths}^n}$. Thus, it remains to prove that $Pr_{v, D^C}^n(B) = Pr_{v', D^{\mathcal{M}}}^{2n}(\text{sep}(B))$. Therefore, note that $\text{sep}(B_i) \uparrow \text{sep}(B)$. From Lemma 2.2 (see page 16), we obtain

$$Pr_{v, D^C}^n(B) = \lim_{i \rightarrow \infty} Pr_{v, D^C}^n(B_i) = \lim_{i \rightarrow \infty} Pr_{v', D^{\mathcal{M}}}^{2n}(\text{sep}(B_i)) = Pr_{v', D^{\mathcal{M}}}^{2n}(\text{sep}(B)).$$

For decreasing sequences, the same argument applies analogously.

Hence, $\mathfrak{C}$ is a monotone class. Further, as all sets in $\mathfrak{G}_{\text{Paths}^n(C)}$ satisfy Eq. (6.44), it holds $\mathfrak{G}_{\text{Paths}^n(C)} \subseteq \mathfrak{C}$. Thus, the monotone class theorem (Thm. 2.5, page 22) is applicable and states that $\sigma(\mathfrak{G}) \subseteq \mathfrak{C}$. Moreover, by definition of $\mathfrak{F}_{\text{Paths}^n(C)}$, it holds $\sigma(\mathfrak{G}) = \mathfrak{F}_{\text{Paths}^n(C)}$. Therefore we conclude that Eq. (6.44) holds for all $B \in \mathfrak{F}_{\text{Paths}^n}$. From here, the claim follows by the Ionescu-Tulcea extension theorem, which lifts the argument from finite measurable bases to the infinite product σ -field $\mathfrak{F}_{\text{Paths}^\omega}$. $\square$

Now we address the next question: Are there schedulers in $\mathcal{M}$ that induce a probability for the event $\text{sep}(\Pi)$ (where $\Pi \in \mathfrak{F}_{\text{Paths}^\omega(\mathcal{C})}$) that cannot be mimicked by a “native” scheduler $D^{\mathcal{C}}$ in the early CTMDP $\mathcal{C}$? We answer this question in the negative and use the one-to-one correspondence to apply Lemma 6.12 again:

Lemma 6.13. *Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ be a CTMDP and $\mathcal{M} = (\mathcal{S}', \text{Act}, \text{IT}, \text{MT}, \nu')$ be its induced IMC. Further, let $D \in \text{GM}(\mathcal{M})$ be a scheduler in $\mathcal{M}$. Define $D^{\mathcal{C}} \in \text{GM}(\mathcal{C})$ such that $D^{\mathcal{C}}(\pi) = D(\text{sep}(\pi))$ for all $\pi \in \text{Paths}^*(\mathcal{C})$. For all $\Pi \in \mathfrak{F}_{\text{Paths}^\omega(\mathcal{C})}$ it holds that*

$$Pr_{\nu', D}^\omega(\text{sep}(\Pi)) = Pr_{\nu, D^{\mathcal{C}}}^\omega(\Pi).$$

Proof. By Eq. (6.43), the scheduler $D^{\mathcal{M}}$ which corresponds to the early scheduler $D^{\mathcal{C}}$ is the scheduler D . Hence, Lemma 6.12 applies and yields the desired equality. $\square$

Corollary 6.1 (Measure preservation). *Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ be a CTMDP and let $\mathcal{M} = (\mathcal{S}', \text{Act}, \text{IT}, \text{MT}, \nu')$ be its induced IMC. For all $\Pi \in \mathfrak{F}_{\text{Paths}^\omega(\mathcal{C})}$ it holds that*

$$\sup_{D^{\mathcal{C}} \in \text{GM}(\mathcal{C})} Pr_{\nu, D^{\mathcal{C}}}^\omega(\Pi) = \sup_{D^{\mathcal{M}} \in \text{GM}(\mathcal{M})} Pr_{\nu', D^{\mathcal{M}}}^\omega(\text{sep}(\Pi)).$$

Proof. Direct consequence of Lemma 6.12 and Lemma 6.13. $\square$

Theorem 6.8 (Interval bounded reachability in $\mathcal{C}$ and $\mathcal{M}$). *Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \nu)$ be a CTMDP and $\mathcal{M} = (\mathcal{S}', \text{Act}, \text{IT}, \text{MT}, \nu')$ be its induced IMC. For a set $G \subseteq \mathcal{S}$ of goal states and a time interval $I \in \mathcal{I}$ define*

$$\begin{aligned} \diamond^I G &= \{\pi \in \text{Paths}^\omega(\mathcal{C}) \mid \exists t \in I. \pi @ t \in G\} \text{ and} \\ \diamond^I \overline{G} &= \{\pi \in \text{Paths}^\omega(\mathcal{M}) \mid \exists t \in I. \pi @ t \cap \overline{G} \neq \emptyset\}, \end{aligned}$$

where $\overline{G} = G \cup \{s^\alpha \mid s \in G \wedge \alpha \in \text{Act}(s)\}$. Then it holds

$$\sup_{D^{\mathcal{C}} \in \text{GM}(\mathcal{C})} Pr_{\nu, D^{\mathcal{C}}}^\omega(\diamond^I G) = \sup_{D^{\mathcal{M}} \in \text{GM}(\mathcal{M})} Pr_{\nu', D^{\mathcal{M}}}^\omega(\diamond^I \overline{G}). \quad (6.45)$$

Proof. First, observe that $Pr_{\nu', D^{\mathcal{M}}}^\omega(\diamond^I \overline{G}) = Pr_{\nu', D^{\mathcal{M}}}^\omega(\text{sep}(\diamond^I G))$ for all $D^{\mathcal{M}} \in \text{GM}(\mathcal{M})$. To see this, note that $\mathcal{M}$ is an alternating IMC where each interactive goal state is followed directly by a Markovian goal state. Then Cor. 6.1 implies Eq. (6.45). $\square$

6.8 Comparison of different scheduler classes

Consider the CTMDP $\mathcal{C}$ which is depicted in Fig. 6.8(a). To compute the maximum time-bounded reachability probability for state s_4 with respect to initial state s_0 , we apply Def. 6.11 to obtain the induced IMC of $\mathcal{C}$, which is depicted in Fig. 6.8(b).

By Thm. 6.8, we can compute the maximum time-interval bounded reachability probability for state s_4 in the early CTMDP $\mathcal{C}$ by applying the modified value iteration algorithm from Sec. 6.4 to its induced IMC $\mathcal{M}(\mathcal{C})$ and the set of goal states $\overline{G} = \{s_4, s_4^y\}$.

In Fig. 6.9, the curve for early schedulers depicts the results that we obtain for the maximum reachability probability for intervals of the form $[0, z]$ with $z \in \mathbb{Q}_{\geq 0}$.

Moreover, note that the example in Fig. 6.8 is constructed such that it is locally and globally uniform. This enables a comparison of all analysis methods and their underlying scheduler classes, that are currently available for CTMDPs. The results depicted in Fig. 6.9 can be explained as follows:

- As $\mathcal{C}$ is locally uniform, we can compute the maximum time-bounded reachability for late schedulers according to the approximation algorithm in Chapter 5. The results depicted in Fig. 6.9(b) coincide with our theoretical findings in Chapter 4: The class of late schedulers outperforms all other scheduler classes.
- For positional schedulers, the only relevant choice is between actions α and β in state s_1 ; Fig. 6.9 depicts the results for both choices. Hence, the maximum reachability probability for the class of positional schedulers is the maximum of the two curves labeled α and β , respectively.
- Finally, $\mathcal{C}$ is globally uniform; hence, the algorithm in [BHKH05] is applicable, which computes the maximum time-bounded reachability probability for the class of time-abstract schedulers. Due to the restricted scheduler class, the obtained maxima are considerably smaller compared to those that are obtained by time-dependent schedulers. In fact, in Fig. 6.9 they agree with the maximum that is achieved by positional schedulers. This is not surprising, as the only nondeterministic choice in $\mathcal{C}$ occurs in state s_1 , which is always entered along the trajectory $\pi = s_0 \xrightarrow{\alpha} s_1$.

6.9 Related work and conclusions

By providing an efficient and quantifiably precise approximation algorithm to compute interval bounded reachability probabilities, we solve the long standing open problem in the area of performance and dependability evaluation [BHKH05], that is, the CSL model checking problem on CTMDPs and on arbitrary IMCs.

In the setting of stochastic games, the time-bounded reachability problem has been studied extensively in [BFK⁺09], with extensions to timed automata in [BF09]. Closely

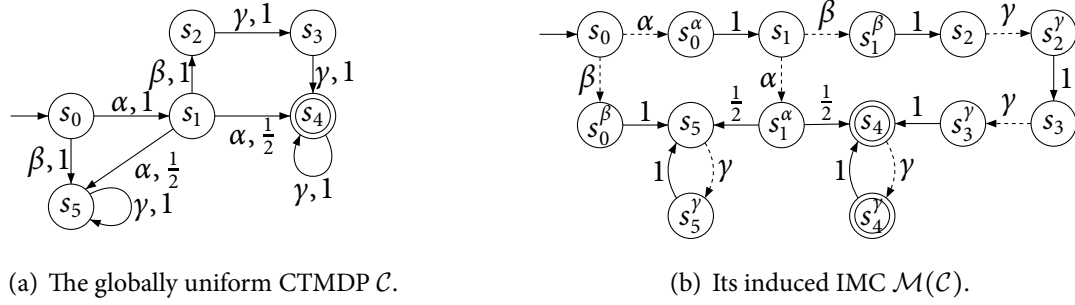


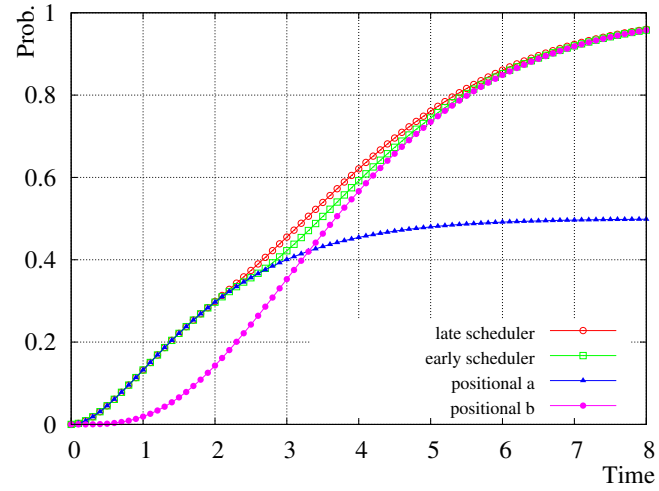
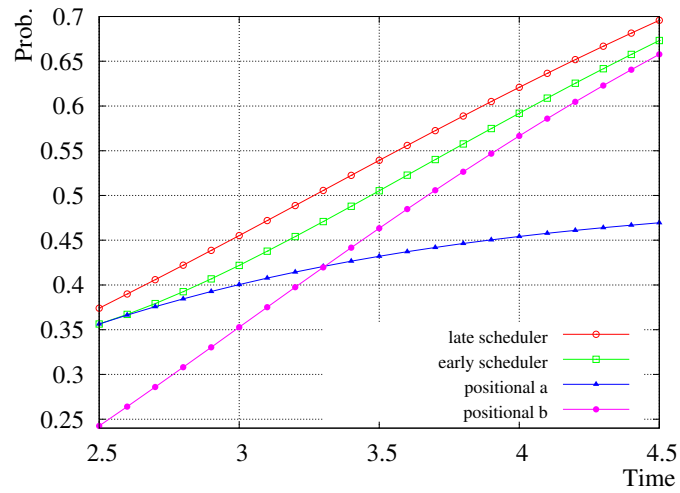
Figure 6.8: Transforming an early CTMDP into its induced IMC.

related to our results in this chapter is the work in [Joh07, BHH⁺09], where globally uniform IMCs — which require the sojourn times in all Markovian states to be equally distributed — are transformed into continuous-time Markov decision processes (CTMDPs). Subsequently, the algorithm in [BHKH05] is used to compute the maximum time-bounded reachability probability in the resulting globally uniform CTMDP. However, the applicability of this approach is severely restricted, as global uniformity is hard (and often impossible) to achieve on nondeterministic models.

Further, the above approaches rely on time-abstract schedulers. From [BHKH05] and Chapter 4 we know that they are strictly less powerful than the time-dependent ones that we consider in this thesis.

Section 6.7 is closely related to Chapter 5, where we analyze time-bounded reachability probabilities in locally uniform CTMDPs under *late schedulers*: From Chapter 4 we know that in locally uniform CTMDPs, late schedulers outperform *early schedulers*, which are the largest class of history- and time-dependent schedulers that is definable on general CTMDPs [Joh07].

Although the discretizations used in Chapters 5 and 6 may appear similar, the obtained results are complementary: In general, transforming IMCs to CTMDPs as done in [Joh07] does not yield locally (or globally) uniform CTMDPs. Hence, the approach in Chapter 5 is inapplicable for the analysis of general IMCs. Reversely however, we have proved in Sec. 6.7 that the problem of computing time-interval bounded reachability in CTMDPs with respect to *early schedulers* can be solved by the analysis of the CTMDP's induced IMC. In this way, this chapter not only solves the problem of model checking IMCs, but also yields a CSL model checking algorithm for early CTMDPs under time and history dependent schedulers.

(a) Reachability of state s_4 within $0 \leq z \leq 8$ time units.

(b) Maxima obtained for different scheduler classes.

Figure 6.9: Maximum time-bounded reachability for the CTMDP and IMC in Fig. 6.8.

7 Equivalences and logics for CTMDPs

The difference between the
right word and the almost
right word is the difference
between lightning and the
lightning bug.

(Mark Twain)

In Chapter 5, we have developed an algorithm to compute time-bounded reachability probabilities in locally uniform CTMDPs. Moreover, in Sec. 6.7, we have shown that similar ideas allow to model check CSL formulas on arbitrary CTMDPs by analyzing their induced IMCs. In fact, this is the first time that efficient and quantifiably precise model checking techniques are available for time-dependent schedulers on arbitrary CTMDPs and IMCs.

In practice however, both models are mostly used as the underlying semantics of high-level modeling formalism such as generalized stochastic Petri nets [CMBC93], stochastic activity networks [SM00] and dynamic fault trees [BCS07]. These formalism allow to represent complex models in a compact and structured way. Once the high-level model is finished, it is transformed into an equivalent CTMDP (or IMC) which is then the starting point for the analysis.

However, during this transformation, one usually encounters the *state space explosion problem*: The unfolding of a rather compact high-level model in many cases yields a CTMDP with an exponentially larger state space. For an example, we refer to the GSPN model of a workstation cluster that we analyze in Chapter 8.

Even though the approximation algorithms that we have developed in the previous chapters are all in PTIME, the state space explosion problem still renders them inapplicable for large scale applications. This is not surprising, as the same problem also arises in the classical setting, where CTL and LTL formulas are verified on Kripke structures. To address this problem, equivalence notions such as strong- and weak bisimulation have been proposed, which allow to minimize the state space by identifying states that have similar behavior.

This idea has carried over to the stochastic setting with great success: For example, bisimulation minimization has become a standard tool for reducing the state space when model checking CTMCs [BHHK03], DTMCs [LS91, BKH05] and MDPs [SL95]. Further, due to their process algebraic background, it comes as no surprise that strong and weak bisimulation are readily available for IMCs [HHK02]. In this setting, lumping (i.e.

bisimulation minimization) has been used to eliminate τ -transitions [MT06].

Such results do not exist for CTMDPs and a corresponding notion of strong bisimulation has not been defined yet. This chapter is meant to close this theoretical gap:

We define strong bisimulation on CTMDPs as a conservative extension of the existing notion of strong bisimulation on CTMCs [Buc94] and investigate which kind of logical properties it preserves. In particular, we show that bisimulation preserves the validity of CSL [ASSB00, BHHK03], which we already used in a slightly restricted version to reason about IMCs (cf. Sec. 6.5).

Accordingly, in this chapter, we provide a semantics of CSL on CTMDPs which is obtained in a similar way as the semantics of PCTL on MDPs [BK98, BdA95]. We show the semantic soundness of our definition by using measure-theoretic arguments to prove that bisimilar states preserve full CSL. Finally, we close the discussion by noting that similar to MDPs, CSL equivalence does not coincide with bisimulation: This observation corresponds to the discrete-time case [Bai98], where reasoning about the maximal and minimal achievable probabilities (as done by logics like PCTL) is not enough to fully characterize the model, either.

Organization of this chapter. In Sec. 7.1 we define strong bisimulation for CTMDPs and investigate its properties. In Sec. 7.2 we adapt CSL to reason about CTMDPs; in this context, we answer the question whether CSL path formulas induce measurable sets in the affirmative. Section 7.3 finally proves that CSL-formulas are preserved under strong bisimulation.

7.1 Strong bisimilarity

By definition, CSL is a *state based* logic which reasons about the labeling of the states of a CTMDP. As this chapter aims at establishing the relation between CSL and strong bisimulation, we extend the definition of CTMDPs (cf. Def. 3.11 on page 75) with a state labeling function $L : \mathcal{S} \rightarrow 2^{AP}$ that assigns each state of the CTMDP the set of *atomic propositions* from the set AP , that hold in that state.

Strong bisimilarity [BKHW05, LS91] is an equivalence on the set of states of a CTMDP which relates two states if they are equally labeled and exhibit the same stepwise behavior. As we will prove in Thm. 7.4, strong bisimilarity allows us to aggregate the state space while preserving transient and long run measures.

As usual, we denote the equivalence class of s under an equivalence relation $\mathcal{R} \subseteq \mathcal{S} \times \mathcal{S}$ by $[s]_{\mathcal{R}}$ and define $[s]_{\mathcal{R}} = \{s' \in \mathcal{S} \mid (s, s') \in \mathcal{R}\}$. If $\mathcal{R}$ is clear from the context, we also write $[s]$ instead of $[s]_{\mathcal{R}}$. Further, $\mathcal{S}_{\mathcal{R}} = \{[s]_{\mathcal{R}} \mid s \in \mathcal{S}\}$ is the quotient space of $\mathcal{S}$ under $\mathcal{R}$.

Definition 7.1 (Strong bisimulation relation). Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, AP, L, \nu)$ be a state labeled CTMDP. An equivalence relation $\mathcal{R} \subseteq \mathcal{S} \times \mathcal{S}$ is a strong bisimulation relation iff for all $(u, v) \in \mathcal{R}$ it holds that $L(u) = L(v)$ and $\mathbf{R}(u, \alpha, C) = \mathbf{R}(v, \alpha, C)$ for all $\alpha \in \text{Act}$ and all $C \in \mathcal{S}_{\mathcal{R}}$.

Two states u and v are strongly bisimilar (denoted $u \sim v$) iff there exists a strong bisimulation relation $\mathcal{R}$ such that $(u, v) \in \mathcal{R}$. Strong bisimilarity is the union of all strong bisimulation relations.

Theorem 7.1 (Strong bisimilarity). Strong bisimilarity is

- (a) an equivalence,
- (b) a strong bisimulation relation, and
- (c) the largest strong bisimulation relation.

Proof. As usual, we use $\sim = \bigcup \{ \mathcal{R} \mid \mathcal{R} \text{ is a strong bisimulation relation on } \mathcal{S} \}$ to denote strong bisimilarity. We prove each claim separately:

- (a) $\sim$ is an equivalence: Reflexivity and symmetry follow directly from the definition. For reflexivity, note that the identity relation is a strong bisimulation relation. For symmetry, it suffices to note that if $u \sim v$, then $(u, v) \in \mathcal{R}$ for some strong bisimulation relation $\mathcal{R}$. Hence $L(u) = L(v)$ and $\mathbf{R}(u, \alpha, C) = \mathbf{R}(v, \alpha, C)$ for all $\alpha \in \text{Act}$ and all $C \in \mathcal{S}_{\mathcal{R}}$. Then $\mathcal{R}^{-1} = \{ (v, u) \mid (u, v) \in \mathcal{R} \}$ is a strong bisimulation relation that proves $v \sim u$.

We need to show transitivity, that is $(u, v) \in \sim$ and $(v, w) \in \sim \implies (u, w) \in \sim$.

$(u, v) \in \sim \implies$ ex. strong bisimulation relation $\mathcal{R}_1 \subseteq \sim$ such that $(u, v) \in \mathcal{R}_1$.

$(v, w) \in \sim \implies$ ex. strong bisimulation relation $\mathcal{R}_2 \subseteq \sim$ such that $(v, w) \in \mathcal{R}_2$.

Let $\mathcal{R}$ denote the transitive closure of $\mathcal{R}_1 \cup \mathcal{R}_2$. Then $(u, w) \in \mathcal{R}$. Therefore it suffices to show that $\mathcal{R}$ is a strong bisimulation relation. As $\mathcal{R}$ obviously is an equivalence, it remains to show that for all $(u, v) \in \mathcal{R}$, $\alpha \in \text{Act}$ and $C \in \mathcal{S}_{\mathcal{R}}$ it holds $L(u) = L(v)$ and

$$\mathbf{R}(u, \alpha, C) = \mathbf{R}(v, \alpha, C). \quad (7.1)$$

The first condition, $L(u) = L(v)$ follows directly from the transitivity of the identity relation on 2^{AP} . For Cond. (7.1), let $C = \{s_1, \dots, s_n\} \in \mathcal{S}_{\mathcal{R}}$. Then it holds for $k = 1, 2$ that $C = \bigcup_{i=1}^n [s_i]_{\mathcal{R}_k}$; to see this, we prove both directions:

$\subseteq$: Let $s \in C$. Then $s \in [s_i]_{\mathcal{R}_k}$ for some $i \in \{1, \dots, n\}$. Hence $s \in \bigcup_{i=1}^n [s_i]_{\mathcal{R}_k}$.

$\ni$: Let $i \in \{1, \dots, n\}$. Then it holds:

$$\begin{aligned}
 s \in [s_i]_{\mathcal{R}_k} &\iff (s, s_i) \in \mathcal{R}_k & (* \text{ by definition } *) \\
 &\implies (s, s_i) \in \mathcal{R} & (* \mathcal{R}_k \subseteq \mathcal{R} *) \\
 &\iff s \in [s_i]_{\mathcal{R}} & (* \mathcal{R} \text{ is an equivalence relation } *) \\
 &\iff s \in C & (* [s_i]_{\mathcal{R}} = C *)
 \end{aligned}$$

Hence we can decompose C into equivalence classes with respect to $\mathcal{R}_1$ and $\mathcal{R}_2$ (see Fig. 7.1). As $\mathcal{R}_1$ is an equivalence relation, it induces a partitioning of C :

$$C = \bigcup \{ [s_{i_1}]_{\mathcal{R}_1}, [s_{i_2}]_{\mathcal{R}_1}, \dots, [s_{i_m}]_{\mathcal{R}_1} \} \text{ where } m \leq n. \quad (7.2)$$

Note that the same applies to $\mathcal{R}_2$ for a different set of indices $i'_1, \dots, i'_{m'}$. Now we are able to prove Property (7.1) by induction on the structure of $\mathcal{R}$. Therefore we provide an inductive definition of $\mathcal{R}$ as follows:

$$\begin{aligned}
 \mathcal{R}^0 &= \mathcal{R}_1 \cup \mathcal{R}_2 \quad \text{and} \\
 \mathcal{R}^{i+1} &= \{ (u, w) \mid \exists v \in \mathcal{S}. (u, v) \in \mathcal{R}^i \wedge (v, w) \in \mathcal{R}^i \} \quad \text{for } i \geq 0.
 \end{aligned}$$

By construction, the subset-ordering on $\mathcal{R}^i$ is bounded from above by $\mathcal{S} \times \mathcal{S}$. Further, $\mathcal{S}$ is finite, so that $\mathcal{R}^0 \subseteq \mathcal{R}^1 \subseteq \dots$ is an increasing sequence, that is, the transitive closure is reached after a finite number z of iterations such that $\mathcal{R}^{z+1} = \mathcal{R}^z$. Obviously, we then have $\mathcal{R} = \mathcal{R}^z$.

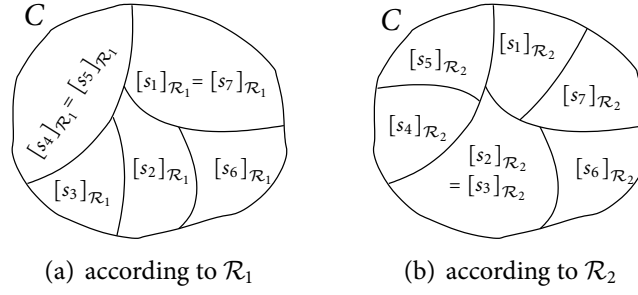
By induction on i , we prove that if $(u, v) \in \mathcal{R}^i$, then $\mathbf{R}(u, \alpha, C) = \mathbf{R}(v, \alpha, C)$ for all $\alpha \in \text{Act}$ and $C \in \mathcal{S}_{\mathcal{R}}$:

i. For the induction base ($i = 0$), we distinguish two cases:

- Let $(u, v) \in \mathcal{R}_1$:

$$\begin{aligned}
 (u, v) \in \mathcal{R}_1 &\implies \forall C' \in \mathcal{S}_{\mathcal{R}_1}. \forall \alpha \in \text{Act}. \mathbf{R}(u, \alpha, C') = \mathbf{R}(v, \alpha, C') \\
 &\implies \forall j \in \{1, \dots, m\}. \forall \alpha \in \text{Act}. \\
 &\quad \mathbf{R}(u, \alpha, [s_{i_j}]_{\mathcal{R}_1}) = \mathbf{R}(v, \alpha, [s_{i_j}]_{\mathcal{R}_1}) \\
 &\implies \forall \alpha \in \text{Act}. \sum_{j=1}^m \mathbf{R}(u, \alpha, [s_{i_j}]_{\mathcal{R}_1}) = \sum_{j=1}^m \mathbf{R}(v, \alpha, [s_{i_j}]_{\mathcal{R}_1}) \\
 &\implies \forall \alpha \in \text{Act}. \mathbf{R}(u, \alpha, \bigcup_{j=1}^m [s_{i_j}]_{\mathcal{R}_1}) = \mathbf{R}(v, \alpha, \bigcup_{j=1}^m [s_{i_j}]_{\mathcal{R}_1}) \\
 &\stackrel{(7.2)}{\implies} \forall \alpha \in \text{Act}. \mathbf{R}(u, \alpha, C) = \mathbf{R}(v, \alpha, C).
 \end{aligned}$$

- Let $(u, v) \in \mathcal{R}_2$: The argument is completely analogue to the first case.

Figure 7.1: Example partitioning of an equivalence class $C \in \mathcal{S}_{\mathcal{R}}$.

- ii. In the induction step ($i \rightsquigarrow i + 1$), assume $(u, w) \in \mathcal{R}^{i+1}$. By construction, we have $(u, v) \in \mathcal{R}^i$ and $(v, w) \in \mathcal{R}^i$. Applying the *induction hypothesis* we have $\mathbf{R}(u, \alpha, C) = \mathbf{R}(v, \alpha, C)$ and $\mathbf{R}(v, \alpha, C) = \mathbf{R}(w, \alpha, C)$ for all actions $\alpha \in \text{Act}$ and all $C \in \mathcal{S}_{\mathcal{R}}$. Therefore $\mathbf{R}(u, \alpha, C) = \mathbf{R}(w, \alpha, C)$ directly follows from the transitivity of $=$ on $\mathbb{R}_{\geq 0}$.

Now we can conclude that $\sim$ is indeed transitive: Given $(u, v) \in \mathcal{R}_1$ and $(v, w) \in \mathcal{R}_2$, there exists a strong bisimulation relation $\mathcal{R}$ such that $(u, w) \in \mathcal{R}$. By definition, $\mathcal{R} \subseteq \sim$ and therefore $u \sim w$.

- (b) $\sim$ is a strong bisimulation relation:

It remains to show for any $u \sim v$, that $L(u) = L(v)$ and $\mathbf{R}(u, \alpha, C) = \mathbf{R}(v, \alpha, C)$ holds for all $\alpha \in \text{Act}$ and $C \in \mathcal{S}_{\sim}$. Since $u \sim v$ implies the existence of a strong bisimulation relation $\mathcal{R} \subseteq \sim$ with $(u, v) \in \mathcal{R}$ it holds that $L(u) = L(v)$ and we may follow the idea in Eq. (7.2) and express C as finite union of equivalence classes of $\mathcal{S}_{\mathcal{R}}$. Since $\mathcal{R}$ is a strong bisimulation relation, the rates from u and v into those equivalence classes are equal and maintained by summation.

- (c) $\sim$ is the largest (i.e. the coarsest) strong bisimulation relation:

Clear from the fact that $\sim$ is the union of all strong bisimulation relations. $\square$

For the purpose of reducing the state space, the *quotient* CTMDP is essential: Instead of considering all states in $\mathcal{S}$, the quotient only retains their equivalence classes under strong bisimilarity:

Definition 7.2 (Quotient). Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, AP, L, \nu)$ be a state labeled CTMDP. The CTMDP $\tilde{\mathcal{C}} = (\tilde{\mathcal{S}}, \text{Act}, \tilde{\mathbf{R}}, AP, \tilde{L})$ where $\tilde{\mathcal{S}} = \mathcal{S}_{\sim}$, $\tilde{\mathbf{R}}([s], \alpha, C) = \mathbf{R}(s, \alpha, C)$ and $\tilde{L}([s]) = L(s)$ for all $s \in \mathcal{S}$, $\alpha \in \text{Act}$ and $C \in \tilde{\mathcal{S}}$ is the quotient of $\mathcal{C}$ under strong bisimilarity.

For states $[s], [t] \in \tilde{\mathcal{S}}$ of the quotient $\tilde{\mathcal{C}}$, let $\tilde{E}([s], \alpha) = \sum_{[s'] \in \tilde{\mathcal{S}}} \tilde{\mathbf{R}}([s], \alpha, [s'])$ be the exit rate of $[s]$ under action α . Further, if $\tilde{E}([s], \alpha) > 0$, then $\tilde{\mathbf{P}}([s], \alpha, [t]) = \frac{\tilde{\mathbf{R}}([s], \alpha, [t])}{\tilde{E}([s], \alpha)}$ is the

discrete branching probability from state $[s]$ to state $[t]$ under action α . For $\tilde{E}([s], \alpha) = 0$, we set $\tilde{\mathbf{P}}([s], \alpha, [t]) = 0$.

Example 7.1. Consider the CTMDP over the set $AP = \{a\}$ of atomic propositions depicted in Fig. 7.2(a). Its quotient under strong bisimilarity is outlined in Fig. 7.2(b). In this example, the states s_2 and s_3 are strongly bisimilar. The corresponding strong bisimulation relation is $\mathcal{R} = \{(s_0, s_0), (s_1, s_1), (s_2, s_2), (s_2, s_3), (s_3, s_3), (s_3, s_2)\}$. $\diamond$

In the quotient, exit rates and branching probabilities are preserved with respect to the underlying CTMDP as shown by the following two lemmas:

Lemma 7.1 (Preservation of exit rates). Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, AP, L, \nu)$ be a state labeled CTMDP and let $\tilde{\mathcal{C}}$ be its quotient under strong bisimilarity. Then $E(s, \alpha) = \tilde{E}([s], \alpha)$ for all $s \in \mathcal{S}$ and $\alpha \in \text{Act}$.

Proof. Let $\mathcal{S} = \bigcup_{k=0}^n [s_{i_k}]$ such that $[s_{i_j}] \cap [s_{i_k}] = \emptyset$ for all $j \neq k$. For all states $s \in \mathcal{S}$ it holds:

$$\begin{aligned} E(s, \alpha) &= \sum_{s' \in \mathcal{S}} \mathbf{R}(s, \alpha, s') = \sum_{k=0}^n \sum_{s' \in [s_{i_k}]} \mathbf{R}(s, \alpha, s') = \sum_{k=0}^n \mathbf{R}(s, \alpha, [s_{i_k}]) \\ &\stackrel{\text{Def. 7.2}}{=} \sum_{k=0}^n \tilde{\mathbf{R}}([s], \alpha, [s_{i_k}]) = \sum_{[s'] \in \tilde{\mathcal{S}}} \tilde{\mathbf{R}}([s], \alpha, [s']) = \tilde{E}([s], \alpha). \quad \square \end{aligned}$$

With Lemma 7.1 it directly follows that also the discrete transition probabilities are preserved under strong bisimulation:

Lemma 7.2 (Preservation of transition probabilities). Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, AP, L, \nu)$ be a state labeled CTMDP and let $\tilde{\mathcal{C}}$ be its quotient under strong bisimilarity. For all states $s, t \in \mathcal{S}$ and all actions $\alpha \in \text{Act}$ it holds

$$\tilde{\mathbf{P}}([s], \alpha, [t]) = \sum_{t' \in [t]} \mathbf{P}(s, \alpha, t').$$

Proof.

$$\begin{aligned} \tilde{\mathbf{P}}([s], \alpha, [t]) &= \frac{\tilde{\mathbf{R}}([s], \alpha, [t])}{\tilde{E}([s], \alpha)} \stackrel{\text{Def. 7.2}}{=} \frac{\mathbf{R}(s, \alpha, [t])}{\tilde{E}([s], \alpha)} \\ &= \frac{\sum_{t' \in [t]} \mathbf{R}(s, \alpha, t')}{\tilde{E}([s], \alpha)} \stackrel{\text{Lemma 7.1}}{=} \frac{\sum_{t' \in [t]} \mathbf{R}(s, \alpha, t')}{E(s, \alpha)} = \sum_{t' \in [t]} \mathbf{P}(s, \alpha, t'). \quad \square \end{aligned}$$

With these remarks, we conclude our definition of strong bisimulation for CTMDPs. To set its definition in a context, we adapt the continuous stochastic logic that we already used in Chapter 6 to reason about IMCs, to reason about CTMDPs.

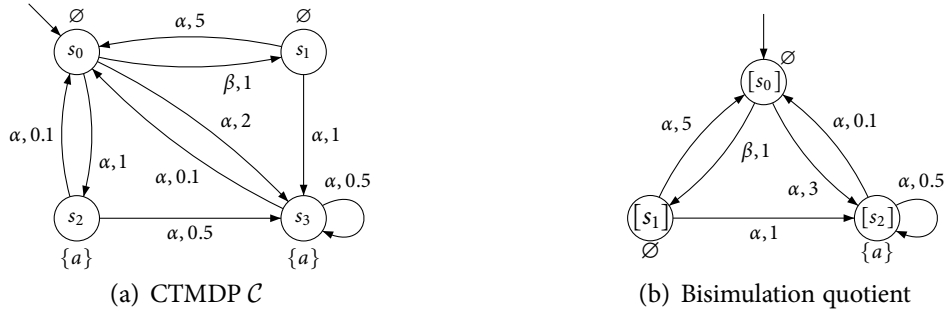


Figure 7.2: Quotient under strong bisimilarity.

7.2 Continuous Stochastic Logic

Continuous stochastic logic [ASSB00, BHHK03] is a state-based logic which was originally designed to reason about continuous-time Markov chains. In this context, its formulas characterize strong bisimilarity [DP03] as defined in [BHHK03]; moreover, strongly bisimilar states satisfy the same CSL formulas [BHHK03].

In this section, we extend CSL to CTMDPs along the lines of [BHHK04]. As steady states do not exist in CTMDPs, we further introduce a long-run average operator [dA97], which serves as a replacement of the steady state operator known from classical CSL. The semantics that we propose for CSL on CTMDPs is based on ideas from [BK98, BdA95] where variants of PCTL are extended to (discrete time) MDPs.

Definition 7.3 (CSL syntax). For $a \in AP$, $p \in [0, 1]$, $I \subseteq \mathbb{R}_{\geq 0}$ a nonempty interval and $\triangleleft \in \{<, \leq, \geq, >\}$, CSL state and CSL path formulas are defined according to the following grammar rules:

$$\Phi ::= a \mid \neg\Phi \mid \Phi \wedge \Phi \mid \forall^{\triangleleft p} \varphi \mid \mathsf{L}^{\triangleleft p} \Phi \quad \text{and} \quad \varphi ::= \mathsf{X}^I \Phi \mid \Phi \mathsf{U}^I \Phi.$$

The Boolean connectives $\vee$ and $\rightarrow$ are defined as usual; further we extend the syntax by deriving the timed modal operators “eventually” and “always” using the equalities $\Diamond^I \Phi \equiv \text{tt } \mathsf{U}^I \Phi$ and $\Box^I \Phi \equiv \neg \Diamond^I \neg \Phi$ where $\text{tt} := a \vee \neg a$ for some $a \in AP$. Similarly, the equality $\exists^{\triangleleft p} \varphi \equiv \neg \forall^{\triangleright p} \varphi$ defines an existentially quantified transient state operator, where $\triangleright$ denotes the negation of the comparison operator $\triangleleft$: For example, if $\triangleleft = <$, then $\triangleright = \geq$. The intuition for the probabilistic and the long-run average operators is given by an example:

Example 7.2. Reconsider the CTMDP depicted in Fig. 7.2(a). The transient state formula $\forall^{>0.1} (\Diamond^{[0,1]} a)$ states that the probability to reach an a -labeled state within at most one time unit exceeds 0.1, no matter how the nondeterministic choices in the current state are resolved.

Further, the long-run average formula $L^{<0.25}(\neg a)$ states that for all scheduling decisions, the system spends less than 25% of its execution time in non- a states, on average. $\diamond$

Formally, the long-run average is derived as follows: For $B \subseteq \mathcal{S}$, let $\mathbf{I}_B$ denote an indicator with $\mathbf{I}_B(s) = 1$ if $s \in B$ and 0 otherwise. Following the ideas of [dA97, LHK01], we compute the fraction of time spent in states from the set B on an infinite path π up to time bound $t \in \mathbb{R}_{\geq 0}$ and define

$$avg_{B,t}(\pi) = \frac{1}{t} \int_0^t \mathbf{I}_B(\pi@t') dt'.$$

As $avg_{B,t}$ is a random variable, its expectation can be derived given an initial distribution $\nu \in \text{Distr}(\mathcal{S})$ and a measurable scheduler $\mathcal{D} \in \text{GM}$. In this way, we obtain

$$E(avg_{B,t}) = \int_{\text{Paths}^\omega} avg_{B,t}(\pi) Pr_{\nu, \mathcal{D}}^\omega(d\pi).$$

Having defined the expectation for a fixed time bound $t \in \mathbb{R}_{\geq 0}$, we now take the limit $t \rightarrow \infty$ and obtain the long-run average as $\lim_{t \rightarrow \infty} E(avg_{B,t})$. This idea is made precise in the semantics of CSL:

Definition 7.4 (CSL semantics). Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, AP, L, \nu)$ be a state labeled CTMDP, $s, t \in \mathcal{S}$, $a \in AP$, $\trianglelefteq \in \{<, \leq, \geq, >\}$ and $\pi \in \text{Paths}^\omega$. Further let $\nu_s(t) := 1$ if $s = t$ and 0 otherwise. The semantics of state formulas is defined as follows:

$$\begin{aligned} s \models a & \iff a \in L(s) \\ s \models \neg \Phi & \iff \text{not } s \models \Phi \\ s \models \Phi \wedge \Psi & \iff s \models \Phi \text{ and } s \models \Psi \\ s \models \forall^{\trianglelefteq p} \varphi & \iff \forall \mathcal{D} \in \text{GM}. Pr_{\nu_s, \mathcal{D}}^\omega \{ \pi \in \text{Paths}^\omega \mid \pi \models \varphi \} \trianglelefteq p \\ s \models L^{\trianglelefteq p} \Phi & \iff \forall \mathcal{D} \in \text{GM}. \lim_{t \rightarrow \infty} \int_{\text{Paths}^\omega} avg_{\text{Sat}(\Phi), t}(\pi) Pr_{\nu_s, \mathcal{D}}^\omega(d\pi) \trianglelefteq p. \end{aligned}$$

The semantics of path formulas is defined such that

$$\begin{aligned} \pi \models X^I \Phi & \iff \pi[1] \models \Phi \wedge \delta(\pi, 0) \in I \\ \pi \models \Phi \mathcal{U}^I \Psi & \iff \exists t \in I. (\pi@t \models \Psi \wedge (\forall t' \in [0, t). \pi@t' \models \Phi)), \end{aligned}$$

where $\text{Sat}(\Phi) = \{s \in \mathcal{S} \mid s \models \Phi\}$.

In Def. 7.4, the transient-state operator $\forall^{\trianglelefteq p} \varphi$ is based on the measure of the set of paths that satisfy φ . However, in order to associate a probability to the set $\{\pi \in \text{Paths}^\omega \mid \pi \models \varphi\}$, we must prove that the set is measurable with respect to the σ -field $\mathfrak{F}_{\text{Paths}^\omega}$. This is the result of the next theorem:

Theorem 7.2 (Measurability of path formulas). The set $\{\pi \in \text{Paths}^\omega \mid \pi \models \varphi\}$ is measurable for all CSL path formula φ .

Proof. For next formulas, the proof is straightforward. For until formulas, let $\pi = s_0 \xrightarrow{\alpha_0, t_0} s_1 \xrightarrow{\alpha_1, t_1} \dots \in \text{Paths}^\omega$ and assume $\pi \models \Phi \mathcal{U}^I \Psi$. By Def. 7.4 it holds that $\pi \models \Phi \mathcal{U}^I \Psi$ iff $\exists t \in I. (\pi @ t \models \Psi \wedge \forall t' \in [0, t). \pi @ t' \models \Phi)$. As we may exclude Zeno behavior by Thm. 3.5 (see page 84), there exists $n \in \mathbb{N}$ with $\pi @ t = \pi[n] = s_n$ such that I and the period of time $[\sum_{i=0}^{n-1} t_i, \sum_{i=0}^n t_i)$ spent in state s_n overlap; further $s_n \models \Psi$ and $s_i \models \Phi$ for $i = 0, \dots, n-1$. Note however, that s_n must also satisfy Φ except for the case of *instantaneous arrival* where $\sum_{i=0}^{n-1} t_i \in I$. Accordingly, the set $\{\pi \in \text{Paths}^\omega \mid \pi \models \Phi \mathcal{U}^I \Psi\}$ can be represented by the union

$$\bigcup_{n=0}^{\infty} \left\{ \pi \in \text{Paths}^\omega \mid \sum_{i=0}^{n-1} t_i \in I \wedge \pi[n] \models \Psi \wedge \forall m < n. \pi[m] \models \Phi \right\} \quad (7.3)$$

$$\cup \bigcup_{n=0}^{\infty} \left\{ \pi \in \text{Paths}^\omega \mid \left(\sum_{i=0}^{n-1} t_i, \sum_{i=0}^n t_i \right) \cap I \neq \emptyset \wedge \pi[n] \models \Psi \wedge \forall m \leq n. \pi[m] \models \Phi \right\}. \quad (7.4)$$

It suffices to show that the subsets in the unions of Eq. (7.3) and Eq. (7.4) are measurable cylinders for all $n \in \mathbb{N}$. In the following, we give the proof for Eq. (7.4) and closed intervals $I = [a, b]$ and only note that the other cases are similar. For fixed $n \geq 0$, we show that the corresponding cylinder base is measurable using a discretization argument:

$$\begin{aligned} & \left\{ \pi \in \text{Paths}^{n+1} \mid \left(\sum_{i=0}^{n-1} t_i, \sum_{i=0}^n t_i \right) \cap [a, b] \neq \emptyset \wedge \pi[n] \models \Psi \wedge \forall m \leq n. \pi[m] \models \Phi \right\} \\ &= \bigcup_{k=1}^{\infty} \bigcup_{\substack{c_0 + \dots + c_n \geq ak \\ d_0 + \dots + d_{n-1} \leq bk \\ c_i < d_i}} \prod_{i=0}^{n-1} \left[\text{Sat}(\Phi) \times \text{Act} \times \left(\frac{c_i}{k}, \frac{d_i}{k} \right) \right] \times \text{Sat}(\Phi \wedge \Psi) \times \text{Act} \times \left(\frac{c_n}{k}, \infty \right) \times \mathcal{S} \end{aligned} \quad (7.5)$$

where $c_i, d_j \in \mathbb{N}$. To shorten notation, let $c = \sum_{i=0}^{n-1} t_i$ and $d = \sum_{i=0}^n t_i$.

We prove Eq. (7.5) in both directions separately:

$\subseteq$: Let $\pi = s_0 \xrightarrow{\alpha_0, t_0} s_1 \xrightarrow{\alpha_1, t_1} \dots \xrightarrow{\alpha_n, t_n} s_{n+1}$ be in the set on the left-hand side of Eq. (7.5). The intervals (c, d) and $[a, b]$ overlap, hence $c < b$ and $d > a$ (see top of Fig. 7.3). Further $\pi[i] \models \Phi$ for $i = 0, \dots, n$ and $\pi[n] \models \Psi$. To show that π is in the set on the right-hand side, let $c_i = \lceil t_i \cdot k - 1 \rceil$ and $d_i = \lfloor t_i \cdot k + 1 \rfloor$ for $k > 0$. Then $\frac{c_i}{k} < t_i < \frac{d_i}{k}$ approximates the sojourn times t_i as depicted in Fig. 7.3. Further let $\varepsilon = \sum_{i=0}^n t_i - a$ and choose k_0 such that $\frac{n+1}{k_0} \leq \varepsilon$ to obtain

$$a = \sum_{i=0}^n t_i - \varepsilon \leq \sum_{i=0}^n t_i - \frac{n+1}{k_0} \leq \sum_{i=0}^n \frac{c_i + 1}{k_0} - \frac{n+1}{k_0} = \sum_{i=0}^n \frac{c_i}{k_0}.$$

Thus $ak \leq \sum_{i=0}^n c_i$ for all $k \geq k_0$. Similarly, we obtain $k'_0 \in \mathbb{N}$ s.t. $\sum_{i=0}^{n-1} d_i \leq bk$ for all $k \geq k'_0$. Hence for large k , π is in the set on the right-hand side.

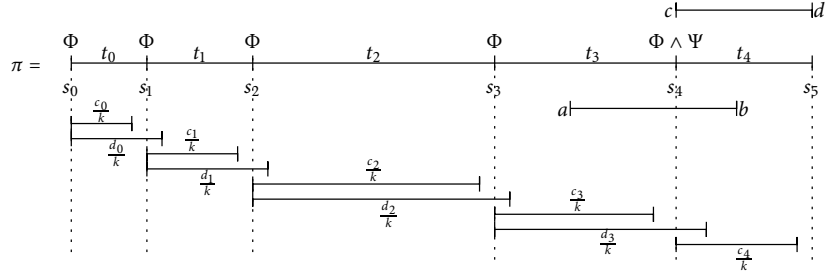


Figure 7.3: Discretization of intervals with $n = 4$ and $I = (a, b)$.

$\ni$: Let π be in the set on the right-hand side of Eq. (7.5) with corresponding values for c_i, d_i and k . Then $t_i \in (\frac{c_i}{k}, \frac{d_i}{k})$. Hence $a \leq \sum_{i=0}^n \frac{c_i}{k} < \sum_{i=0}^n t_i = d$ and $b \geq \sum_{i=0}^{n-1} \frac{d_i}{k} > \sum_{i=0}^{n-1} t_i = c$ so that the time-interval (c, d) of state s_n and the time interval $I = [a, b]$ of the formula overlap. Further, $\pi[m] \models \Phi$ for $m \leq n$ and $\pi[n] \models \Psi$; thus π is in the set on the left-hand side of Eq. (7.5).

The right-hand side of Eq. (7.5) is measurable, hence also the cylinder base. This extends to its cylinder and the countable union in Eq. (7.4). $\square$

7.3 Strong bisimilarity preserves CSL

We now come to the main contribution in this chapter. To prove that strong bisimilarity preserves CSL formulas, we establish a correspondence between certain sets of paths of a CTMDP and its quotient which is measure-preserving:

Definition 7.5 (Simple bisimulation closed). Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \text{AP}, L, \nu)$ be a state labeled CTMDP. A measurable rectangle $\Pi = S_0 \times A_0 \times T_0 \times \dots \times A_{n-1} \times T_{n-1} \times S_n$ is simple bisimulation closed iff $S_i \in (\tilde{\mathcal{S}} \cup \{\emptyset\})$ for $i = 0, \dots, n$. Further, let $\tilde{\Pi} = \{S_0\} \times A_0 \times T_0 \times \dots \times A_{n-1} \times T_{n-1} \times \{S_n\}$ be the corresponding rectangle in the quotient $\tilde{\mathcal{C}}$.

An essential step in our proof strategy is to obtain a scheduler on the quotient. The following example illustrates the intuition for such a scheduler.

Example 7.3. Let $\mathcal{C}$ be the CTMDP in Fig. 7.4(a) where $\nu(s_0) = \frac{1}{4}$, $\nu(s_1) = \frac{2}{3}$ and $\nu(s_2) = \frac{1}{12}$. Moreover, let D be the GM-scheduler such that $D(s_0, \{\alpha\}) = \frac{2}{3}$, $D(s_0, \{\beta\}) = \frac{1}{3}$, $D(s_1, \{\alpha\}) = \frac{1}{4}$ and $D(s_1, \{\beta\}) = \frac{3}{4}$. Intuitively, a scheduler D^ν that mimics D 's behavior on the quotient $\tilde{\mathcal{C}}$ (see Fig. 7.4(b)) can be defined by

$$D^\nu([s_0], \{\alpha\}) = \frac{\sum_{s \in [s_0]} \nu(s) \cdot D(s, \{\alpha\})}{\sum_{s \in [s_0]} \nu(s)} = \frac{\frac{1}{4} \cdot \frac{2}{3} + \frac{2}{3} \cdot \frac{1}{4}}{\frac{1}{4} + \frac{2}{3}} = \frac{4}{11} \quad \text{and}$$

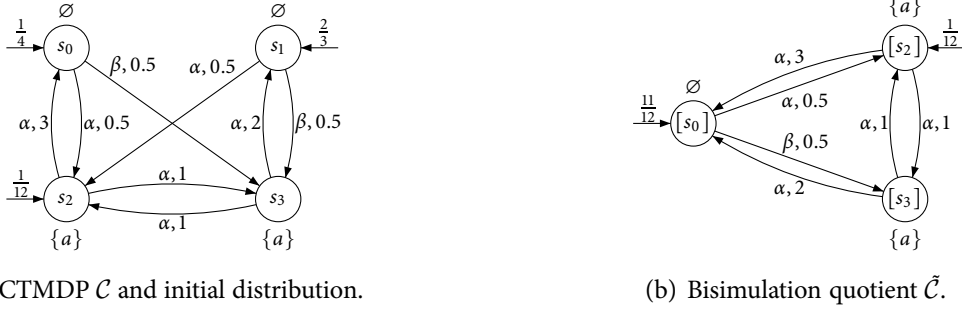
(a) CTMDP $\mathcal{C}$ and initial distribution.(b) Bisimulation quotient $\tilde{\mathcal{C}}$.

Figure 7.4: Derivation of the quotient scheduler.

$$D_{\sim}^v([s_0], \{\beta\}) = \frac{\sum_{s \in [s_0]} v(s) \cdot D(s, \{\beta\})}{\sum_{s \in [s_0]} v(s)} = \frac{\frac{1}{4} \cdot \frac{1}{3} + \frac{2}{3} \cdot \frac{3}{4}}{\frac{1}{4} + \frac{2}{3}} = \frac{7}{11}.$$

Even though s_0 and s_1 are bisimilar, the scheduler D decides differently for the histories $\pi_0 = s_0$ and $\pi_1 = s_1$. As π_0 and π_1 collapse into $\tilde{\pi} = [s_0]$ on the quotient, $D_{\sim}^v$ can no longer distinguish between π_0 and π_1 . Therefore D 's decision for any history $\pi \in \tilde{\pi}$ is weighted with respect to the total probability of $\tilde{\pi}$. $\diamond$

In order to formally derive the quotient scheduler, Def. 7.6 generalizes the ideas from Ex. 7.3 to histories of arbitrary (finite) length:

Definition 7.6 (Quotient scheduler). Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, \text{AP}, L, v)$ be a CTMDP and $D \in \text{GM}$. First, define the history weight of finite paths of length n inductively as follows:

$$\begin{aligned} hw_0(v, D, s_0) &= v(s_0) \text{ and} \\ hw_{n+1}(v, D, \pi \xrightarrow{\alpha_n, t_n} s_{n+1}) &= hw_n(v, D, \pi) \cdot D(\pi, \{\alpha_n\}) \cdot \mathbf{P}(\pi \downarrow, \alpha_n, s_{n+1}). \end{aligned}$$

Let $\tilde{\pi} = [s_0] \xrightarrow{\alpha_0, t_0} \dots \xrightarrow{\alpha_{n-1}, t_{n-1}} [s_n]$ be a timed history of $\tilde{\mathcal{C}}$ and $\Pi = [s_0] \times \{\alpha_0\} \times \{t_0\} \times \dots \times \{\alpha_{n-1}\} \times \{t_{n-1}\} \times [s_n]$ be the corresponding set of paths in $\mathcal{C}$. The quotient scheduler $D_{\sim}^v$ on $\tilde{\mathcal{C}}$ is then defined as follows:

$$D_{\sim}^v(\tilde{\pi}, \alpha_n) = \frac{\sum_{\pi \in \Pi} hw_n(v, D, \pi) \cdot D(\pi, \{\alpha_n\})}{\sum_{\pi \in \Pi} hw_n(v, D, \pi)}.$$

Further, let $\tilde{v}([s]) = \sum_{s' \in [s]} v(s')$ be the initial distribution on $\tilde{\mathcal{C}}$.

A history $\tilde{\pi}$ of $\tilde{\mathcal{C}}$ corresponds to a set of paths Π in $\mathcal{C}$; given $\tilde{\pi}$, the quotient scheduler decides by multiplying D 's decision on each path in Π with its corresponding weight and normalizing with the weight of Π afterwards. In this way, we obtain the first intermediate result: For CTMDP $\mathcal{C}$, if Π is a simple bisimulation closed set of paths, v an initial

distribution and $D \in GM$, the measure of Π in $\mathcal{C}$ coincides with the measure of $\tilde{\Pi}$ in $\tilde{\mathcal{C}}$ which is induced by $\tilde{\nu}$ and D^ν :

Theorem 7.3. *Let $\mathcal{C} = (\mathcal{S}, Act, \mathbf{R}, AP, L, \nu)$ be a CTMDP and $D \in GM(\mathcal{C})$ a scheduler. For all simple bisimulation closed sets of paths Π it holds that*

$$Pr_{\nu, D}^\omega(\Pi) = Pr_{\tilde{\nu}, D^\nu}^\omega(\tilde{\Pi}).$$

Proof. By induction on the length n of cylinder bases. The induction base holds for all $\nu \in Distr(\mathcal{S})$ since $Pr_{\nu, D}^0([s]) = \sum_{s' \in [s]} \nu(s') = \tilde{\nu}([s]) = Pr_{\tilde{\nu}, D^\nu}^0(\{[s]\})$. With the induction hypothesis that $Pr_{\nu, D}^n(\Pi) = Pr_{\tilde{\nu}, D^\nu}^n(\tilde{\Pi})$ for all $\nu \in Distr(\mathcal{S})$, $D \in GM$ and bisimulation closed $\Pi \subseteq Paths^n$ we obtain the induction step:

$$\begin{aligned} Pr_{\nu, D}^{n+1}([s_0] \times A_0 \times T_0 \times \Pi) &= \int_{[s_0] \times A_0 \times T_0} Pr_{\mathbf{P}(s, \alpha, \cdot), D(s \xrightarrow{\alpha, t} \cdot)}^n(\Pi) \mu_{\nu, D}(ds, d\alpha, dt) \\ &= \int_{s \in [s_0]} \nu(ds) \int_{\alpha \in A_0} D(s, d\alpha) \int_{T_0} Pr_{\mathbf{P}(s, \alpha, \cdot), D(s \xrightarrow{\alpha, t} \cdot)}^n(\Pi) \eta_{E(s, \alpha)}(dt) \\ &= \sum_{s \in [s_0]} \nu(s) \sum_{\alpha \in A_0} D(s, \{\alpha\}) \int_{T_0} Pr_{\mathbf{P}(s, \alpha, \cdot), D(s \xrightarrow{\alpha, t} \cdot)}^n(\Pi) \eta_{\tilde{E}([s_0], \alpha)}(dt) \quad (* \text{ Lemma 7.1 } *) \\ &\stackrel{\text{i.h.}}{=} \sum_{s \in [s_0]} \sum_{\alpha \in A_0} \int_{T_0} Pr_{\tilde{\mathbf{P}}([s_0], \alpha, \cdot), D^\nu([s_0] \xrightarrow{\alpha, t} \cdot)}^n(\tilde{\Pi}) \cdot \nu(s) \cdot D(s, \{\alpha\}) \eta_{\tilde{E}([s_0], \alpha)}(dt) \\ &= \sum_{\alpha \in A_0} \int_{T_0} Pr_{\tilde{\mathbf{P}}([s_0], \alpha, \cdot), D^\nu([s_0] \xrightarrow{\alpha, t} \cdot)}^n(\tilde{\Pi}) \cdot \sum_{s \in [s_0]} \left(\nu(s) \cdot D(s, \{\alpha\}) \right) \eta_{\tilde{E}([s_0], \alpha)}(dt) \\ &= \sum_{\alpha \in A_0} \int_{T_0} Pr_{\tilde{\mathbf{P}}([s_0], \alpha, \cdot), D^\nu([s_0] \xrightarrow{\alpha, t} \cdot)}^n(\tilde{\Pi}) \cdot \left(\sum_{s \in [s_0]} \nu(s) \right) \frac{\sum_{s \in [s_0]} \nu(s) \cdot D(s, \{\alpha\})}{\sum_{s \in [s_0]} \nu(s)} \eta_{\tilde{E}([s_0], \alpha)}(dt) \\ &= \sum_{\alpha \in A_0} \int_{T_0} Pr_{\tilde{\mathbf{P}}([s_0], \alpha, \cdot), D^\nu([s_0] \xrightarrow{\alpha, t} \cdot)}^n(\tilde{\Pi}) \cdot \tilde{\nu}([s_0]) \cdot D^\nu([s_0], \{\alpha\}) \eta_{\tilde{E}([s_0], \alpha)}(dt) \\ &= \int_{\{[s_0]\}} \tilde{\nu}(d[s]) \int_{A_0} D^\nu([s], d\alpha) \int_{T_0} Pr_{\tilde{\mathbf{P}}([s], \alpha, \cdot), D^\nu([s] \xrightarrow{\alpha, t} \cdot)}^n(\tilde{\Pi}) \eta_{\tilde{E}([s], \alpha)}(dt) \\ &= \int_{\{[s_0]\} \times A_0 \times T_0} Pr_{\tilde{\mathbf{P}}([s], \alpha, \cdot), D^\nu([s] \xrightarrow{\alpha, t} \cdot)}^n(\tilde{\Pi}) \tilde{\mu}_{\tilde{\nu}, D^\nu}(d[s], d\alpha, dt) \\ &= Pr_{\tilde{\nu}, D^\nu}^{n+1}(\{[s_0]\} \times A_0 \times T_0 \times \tilde{\Pi}) \end{aligned}$$

where $\tilde{\mu}_{\tilde{\nu}, D^\nu}$ is the extension of $\mu_{\nu, D}$ (Def. 3.16) to sets of initial triples in $\tilde{\mathcal{C}}$:

$$\begin{aligned} \tilde{\mu}_{\tilde{\nu}, D^\nu} : \mathfrak{F}_{\tilde{\mathcal{S}} \times Act \times \mathbb{R}_{\geq 0}} &\rightarrow [0, 1] : \\ I &\mapsto \int_{\tilde{\mathcal{S}}} \tilde{\nu}(d[s]) \int_{Act} D^\nu([s], d\alpha) \int_{\mathbb{R}_{\geq 0}} \mathbf{I}_I([s], \alpha, t) \eta_{\tilde{E}([s], \alpha)}(dt). \quad \square \end{aligned}$$

According to Thm. 7.3, the quotient scheduler preserves the measure for *simple* bisimulation closed sets of paths, i.e. for paths, whose state components are equivalence classes under strong bisimilarity. To generalize this to sets of paths that satisfy a CSL path formula, we introduce *general* bisimulation closed sets of paths:

Definition 7.7 (Bisimulation closed). Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, AP, L, \nu)$ be a CTMDP and $\tilde{\mathcal{C}}$ its quotient under strong bisimilarity. A measurable rectangle $\Pi = S_0 \times A_0 \times T_0 \times \dots \times A_{n-1} \times T_{n-1} \times S_n$ is bisimulation closed iff $S_i = \bigcup_{j=0}^{k_i} [s_{i,j}]$ for $k_i \in \mathbb{N}$ and $0 \leq i \leq n$. Let

$$\tilde{\Pi} = \bigcup_{j=0}^{k_0} \{[s_{0,j}]\} \times A_0 \times T_0 \times \dots \times A_{n-1} \times T_{n-1} \times \bigcup_{j=0}^{k_n} \{[s_{n,j}]\}$$

denote the corresponding rectangle in the quotient $\tilde{\mathcal{C}}$.

Lemma 7.3. Any bisimulation closed set of paths Π can be represented as a finite disjoint union of simple bisimulation closed sets of paths.

Proof. Direct consequence of Def. 7.7. □

Corollary 7.1. Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, AP, L, \nu)$ be a CTMDP. Then

$$Pr_{\nu, D}^{\omega}(\Pi) = Pr_{\tilde{\nu}, D_{\sim}^{\nu}}^{\omega}(\tilde{\Pi})$$

for all $D \in GM$ and all bisimulation closed sets of paths Π .

Proof. Follows directly from Lemma 7.3 and Thm. 7.3. □

Using these extensions, we are ready to prove the main result of this chapter:

Theorem 7.4 (Preservation theorem). Let $\mathcal{C} = (\mathcal{S}, \text{Act}, \mathbf{R}, AP, L, \nu)$ be a CTMDP. For all CSL state formulas Φ and for all states $u, v \in \mathcal{S}$ with $u \sim v$ it holds that

$$u \models \Phi \iff v \models \Phi.$$

Proof. By structural induction on Φ .

1. If $\Phi = a$ and $a \in AP$, the induction base follows as $L(u) = L(v)$.

2. In the induction step, conjunction and negation are obvious. Thus we only consider the transient state operator $\forall^{\leq p}$ and the long-run average operator:

- Let $\Phi = \forall^{\leq p} \varphi$ and $\Pi = \{\pi \in Paths^\omega \mid \pi \models \varphi\}$. To show $u \models \forall^{\leq p} \varphi$ implies $v \models \forall^{\leq p} \varphi$ it suffices to show that for any $\mathcal{V} \in GM$ there exists $\mathcal{U} \in GM$ with $Pr_{v_u, \mathcal{U}}^\omega(\Pi) = Pr_{v_v, \mathcal{V}}^\omega(\Pi)$. By Thm. 7.2 the set Π is measurable, hence $\Pi = \bigcup_{i=0}^\infty \Pi_i$ for disjoint $\Pi_i \in \mathfrak{F}_{Paths^\omega}$. By *induction hypothesis* for path formulas $X^I \Phi$ and $\Phi \mathcal{U}^I \Psi$ the sets $Sat(\Phi)$ and $Sat(\Psi)$ are disjoint unions of $\sim$ -equivalence classes. The same holds for any Boolean combination of Φ and Ψ . Hence $\Pi = \bigcup_{i=0}^\infty \Pi_i$ where the Π_i are bisimulation closed. For all $\mathcal{V} \in GM$ and $\pi = s_0 \xrightarrow{\alpha_0, t_0} \dots \xrightarrow{\alpha_{n-1}, t_{n-1}} s_n$ let $\mathcal{U}(\pi) := \mathcal{V}_{v_v}^{v_v}([s_0] \xrightarrow{\alpha_0, t_0} \dots \xrightarrow{\alpha_{n-1}, t_{n-1}} [s_n])$. Thus $\mathcal{U}$ mimics on π the decision of $\mathcal{V}_{v_v}^{v_v}$ on $\tilde{\pi}$. In fact $\mathcal{U}_{v_u}^{v_u} = \mathcal{V}_{v_v}^{v_v}$ since

$$\mathcal{U}_{v_u}^{v_u}(\tilde{\pi}, \alpha_n) = \frac{\sum_{\pi \in \Pi} hw_n(v_u, \mathcal{U}, \pi) \cdot \mathcal{V}_{v_v}^{v_v}(\tilde{\pi}, \alpha_n)}{\sum_{\pi \in \Pi} hw_n(v_u, \mathcal{U}, \pi)}$$

and $\mathcal{V}_{v_v}^{v_v}(\tilde{\pi}, \alpha_n)$ is independent of π . With $\tilde{v}_u = \tilde{v}_v$ and by Corollary 7.1 we obtain $Pr_{v_u, \mathcal{U}}^\omega(\Pi_i) = Pr_{\tilde{v}_u, \mathcal{U}_{v_u}^{v_u}}^\omega(\tilde{\Pi}_i) = Pr_{\tilde{v}_v, \mathcal{V}_{v_v}^{v_v}}^\omega(\tilde{\Pi}_i) = Pr_{v_v, \mathcal{V}}^\omega(\Pi_i)$ which carries over to Π for Π is a countable union of disjoint sets Π_i .

- Let $\Phi = L^{\leq p} \Psi$. Since $u \sim v$, it suffices to show that for all $s \in \mathcal{S}$ it holds $s \models L^{\leq p} \Psi$ iff $[s] \models L^{\leq p} \Psi$. The expectation of $avg_{Sat(\Psi), t}$ for $t \in \mathbb{R}_{\geq 0}$ can be expressed as follows:

$$\begin{aligned} \int_{Paths^\omega} \left(\frac{1}{t} \int_0^t \mathbf{I}_{Sat(\Psi)}(\pi @ t') dt' \right) Pr_{v_s, D}^\omega(d\pi) \\ = \frac{1}{t} \int_0^t Pr_{v_s, D}^\omega\{\pi \in Paths^\omega \mid \pi @ t' \models \Psi\} dt'. \end{aligned}$$

Further, the sets $\{\pi \in Paths^\omega \mid \pi @ t' \models \Psi\}$ and $\{\pi \in Paths^\omega \mid \pi \models \Diamond^{[t', t']} \Psi\}$ have the same measure and the *induction hypothesis* applies to Ψ . Applying the previous reasoning for the until case to the formula $\text{tt } \mathcal{U}^{[t', t']} \Psi$ once, we obtain

$$Pr_{v_s, D}^\omega\{\pi \in Paths^\omega(\mathcal{C}) \mid \pi \models \Diamond^{[t', t']} \Psi\} = Pr_{\tilde{v}_s, D_{\tilde{s}}}^\omega\{\tilde{\pi} \in Paths^\omega(\tilde{\mathcal{C}}) \mid \tilde{\pi} \models \Diamond^{[t', t']} \Psi\}$$

for all $t' \in \mathbb{R}_{\geq 0}$. Thus the expectations of $avg_{Sat(\Psi), t}$ on $\mathcal{C}$ and $\tilde{\mathcal{C}}$ are equal for all $t \in \mathbb{R}_{\geq 0}$ and the same holds for their limits if $t \rightarrow \infty$. This completes the proof as for $u \sim v$ we obtain $u \models L^{\leq p} \Psi$ iff $[u] \models L^{\leq p} \Psi$ iff $[v] \models L^{\leq p} \Psi$ iff $v \models L^{\leq p} \Psi$. $\square$

This theorem shows that bisimilar states satisfy the same CSL formulas.

The reverse direction, however, does not hold in general. One reason is obvious: The logic that we use throughout this thesis is purely state-based. However, the definition of strong bisimulation also accounts for action names. Therefore it comes as no surprise

that CSL cannot characterize strong bisimulation. However, there is another more profound reason which is analogous to the discrete-time setting where extensions of PCTL to Markov decision processes [SL95, Bai98] also cannot express strong bisimilarity: CSL and PCTL only allow to specify infima and suprema as probability bounds under a denumerable class of randomized schedulers; therefore intuitively, CSL cannot characterize exponential distributions which neither contribute to the supremum nor to the infimum of the probability measures of a given set of paths. Thus the counterexample from [Bai98, Fig. 9.5] interpreted as a CTMDP applies verbatim to our case.

7.4 Conclusion

In this chapter we define strong bisimulation on CTMDPs and adapt the continuous stochastic logic (CSL) to CTMDP such that it permits to reason about the maximum and minimum achievable performance and dependability measures in CTMDPs.

Using measure-theoretic arguments, we further prove that CSL path formulas induce measurable sets of paths. As this proof is done in the more general setting of CTMDPs, it applies to CSL-path formulas for CTMCs, as well. In this way, we close a gap in the theory of CSL, where the measurability of path formulas has not been discussed.

The main contribution of this chapter is the proof that strong bisimilarity preserves the validity of CSL formulas. In this way, we justify the definition of bisimulation that we use and embed it into the context of CSL. However, our logic is not capable of characterizing strong bisimilarity. This is not surprising, as similar limitations are also known for logics like PCTL in the discrete-time setting.

A promising approach to obtain a logic that is expressive enough to characterize CTMDPs are action based variants of CSL. To investigate such logics and their relation to scheduler classes remains for future research.

8 Model checking generalized stochastic Petri nets

Perfection is achieved, not
when there is nothing more to
add, but when there is
nothing left to take away.

(Antoine de Saint-Exupéry)

In a stochastic Petri net [Nat80, Mol82], all transitions are delayed according to an exponential distribution. Their associated token game induces a CTMC which represents the SPN's semantics.

This chapter considers generalized stochastic Petri nets [MCB84] (GSPNs) which extend SPNs with *immediate transitions*. Similar to the internal transitions in the closed IMCs of Chapter 6, immediate transitions in a GSPN fire instantaneously. Accordingly, a GSPN distinguishes exponentially delayed *timed transitions* from immediate transitions. Conflicts between immediate transitions lead to so-called “confused” GSPNs, where confusion arises if multiple immediate transitions are enabled at the same time. In principle, the choice which of them executes next is not specified and hence, nondeterministic.

However, at the time GSPNs were developed, no analysis techniques were available for nondeterministic and stochastically timed systems. Therefore, much work has been spent in order to rule out confused GSPNs [MCB84, CMBC93]. The solution that was chosen already in [MCB84] is to assign *weights* to immediate transitions. If multiple immediate transitions compete for execution, the proportion of their weights gives rise to a discrete probability distribution which resolves the nondeterminism probabilistically. Hence, all nondeterministic choices are replaced by probability distributions that are implicitly encoded in the syntax of the GSPN.

In this approach, the modeler has to assign weights “at the net level” [CDF91, CMBC93], that is, without knowing which immediate transitions actually get into conflict during the token game. As observed already in [MCB84], finding reasonable weight-assignments is difficult; for larger systems, it might even be practically impossible.

To mitigate against this shortcoming, the GSPN community tries to identify sets of immediate transitions that *might* get into conflict during the evolution of the GSPN. These *extended conflict sets* [CMBC93] rely on necessary conditions for a conflict and partition the set of immediate transitions accordingly. In this way, weights become local to each block of the ECS equivalence which facilitates the weight specification for the modeler.

The quest to find suitable necessary conditions for the occurrence of conflicts between immediate transitions led to extremely complex and technical definitions of *extended conflict sets*. Among others, this is testified by the research papers [MBCC87, CMBC93, MBC⁺91] and their further refinements in [CDF91, MBC⁺95, Bal00, Bal07]. However, despite all this work, the authors of [TFP99] and [TF03] still managed to disprove the correctness claim (i.e. the claim that immediate transitions in different extended conflict sets can never be in conflict) of the extended conflict set approach.

A further, more general shortcoming of weight-assignments is that weights only permit to formalize positional strategies to resolve the nondeterministic choices that occur in markings with competing immediate transitions. As we have seen in the previous chapters, depending on the measure of interest, positional schedulers are far from optimal.

Therefore, we do not follow this approach, but strive for a general semantics of GSPNs which accepts that nondeterminism occurs between competing immediate transitions. In this way, we obtain a new definition of GSPNs which avoids the use of weights while conservatively extending stochastic Petri nets [Mol82]. In this way, it resembles an earlier approach in [HHMR97] where compositional extensions of GSPN are discussed; in this context, immediate transitions are equipped with action names for synchronization purposes. This approach does not use the weight specification of the classical GSPN definition either, but relies on the fact that the precedence of competing immediate transitions is often resolved by synchronization with the environment. However, as mentioned already in [HHMR97, Sec. 4], nondeterminism cannot be ruled out completely. Instead, it generally occurs in the composed GSPNs due to competing immediate internal τ -transitions.

The same problem is also observed by the authors of [MH06b] and [MH06a]. In their work, they propose a framework for CSL model checking of *deterministic stochastic Petri nets*. The results in [MH06b] are closely related to the approach taken in this chapter. However, the technique that is proposed in [MH06b] is again restricted to deterministic stochastic Petri nets which induce a CTMC [MH06b, Sec. 3].

The results of this chapter overcome these limitations and enable an analysis of nondeterministic GSPNs that may occur in the frameworks [MH06b] and [HHMR97].

Opposed to earlier approaches, we describe the semantics of a GSPN by its marking graph, which is isomorphic to a closed IMC. Hence, our nondeterministic GSPNs can be analyzed by the approximation algorithm from Chapter 6.

Organization of this chapter. Section 8.1 introduces some basic notation. In Sec. 8.2, we define the syntax of GSPNs without weight-assignments. Section 8.3 introduces their semantics by interpreting their marking graph as an IMC. Finally, Sec. 8.4 provides a case study where we apply our GSPN semantics to analyze the dependability characteristics of a workstation cluster which is modeled by a nondeterministic GSPN.

8.1 Preliminaries

Our definition of GSPNs differs from that in [MCB84], as we do not support the specification of weights for immediate transitions. Specifically, we propose to completely abandon the idea of resolving the nondeterministic choices by weight-specifications.

To obtain a simple and semantically precise definition of our GSPNs, we only distinguish between *timed* and *immediate* transitions and do not allow for further priority specifications within the class of immediate transitions. Moreover, we do not care about marking dependent rates. Note however, that this is no severe restriction, as it is straightforward to adapt our approach to the aforementioned generalizations by extending the transformation from GSPNs to IMCs such that it reflects the priority levels and marking dependent rates in the induced *marking graph*.

As in Petri nets, a GSPN consists of finitely many *places* and *transitions*; each place can contain an unbounded finite number of *tokens*. Informally, the state of a GSPN — called a *marking* — is completely determined by the number of tokens in each place:

Definition 8.1 (Marking). *Let P be a nonempty, finite set of places. A marking m is a mapping $m : P \rightarrow \mathbb{N}$. Let $\mathcal{M} = \{m : P \rightarrow \mathbb{N}\}$ denote the set of all markings.*

8.2 The syntax of GSPNs

A GSPN consists of a finite, nonempty set of places and finitely many transitions that connect those places; transitions are further partitioned into the set of *immediate transitions* which execute instantaneously and the set of *timed transitions*, which are delayed by an exponentially distributed amount of time.

Example 8.1. *Consider the GSPN $\mathcal{G}$ in Fig. 8.1(a). It consists of the set of places (denoted by circular nodes) $\{p_0, \dots, p_3\}$; moreover, $\{t_0, t_1, t_2, t_8\}$ is its set of timed transitions (depicted as rectangles) and $\{t_3, t_4, t_5, t_6, t_7\}$ is the set of immediate transitions (solid bars).*

Each transition has a number of input, output and inhibition places¹, depicted as arcs in Fig. 8.1(a). Informally, a transition has concession if enough tokens are available in all its input places, while the corresponding inhibition places are empty. The effect of executing a transition is a new marking, which is obtained by removing a token from each input place and adding tokens to the transition's output places. Immediate transitions execute immediately upon becoming enabled, whereas timed transitions are delayed by an exponentially distributed duration, specified by the transition rate. $\diamond$

To define a GSPN formally, we encode its input, output and inhibition places as function $T \rightarrow (P \rightarrow \mathbb{N})$ which assign to each transition a mapping $P \rightarrow \mathbb{N}$, specifying the cardinality of the input, output or inhibition places.

¹Inhibition places may disable an otherwise enabled transition depending on the current marking.

Definition 8.2 (Generalized stochastic Petri net). A generalized stochastic Petri net (GSPN) is a tuple $\mathcal{G} = (P, T, \lambda, I, O, H, m_0)$ where

- P is a nonempty, finite set of places,
- $T = T_t \cup T_i$ is a finite set of transitions partitioned into the sets T_t and T_i of timed and immediate transitions,
- $\lambda : T_t \rightarrow \mathbb{R}_{>0}$ is a rate assignment,
- $I : T \rightarrow (P \rightarrow \mathbb{N})$ defines the transitions' input places,
- $O : T \rightarrow (P \rightarrow \mathbb{N})$ the transitions' output places and
- $H : T \rightarrow (P \rightarrow \mathbb{N})$ defines the transitions' inhibition places.

Finally, $m_0 \in \mathcal{M}$ is the initial marking.

For a given transition $t \in T$, we use I_t to denote t 's input places, that is, we define $I_t(p) = I(t)(p)$. Similarly, we use O_t and H_t to denote the output and inhibition places of t . Moreover, for any GSPN $\mathcal{G}$ and transition $t \in T$, we use

$$\begin{aligned} \text{pre}(t) &= \{p \in P \mid I_t(p) > 0\} \text{ and} \\ \text{post}(t) &= \{p \in P \mid O_t(p) > 0\} \end{aligned}$$

to define the sets of *input* and *output places* of transition t .

Example 8.2. The input places of the transitions t_6 and t_8 in Fig. 8.1(a) are represented as follows:

$$I_{t_6}(p) = \begin{cases} 1 & \text{if } p \in \{p_2, p_3\} \\ 0 & \text{otherwise} \end{cases} \quad I_{t_8}(p) = \begin{cases} 1 & \text{if } p = p_3 \\ 0 & \text{otherwise.} \end{cases}$$

Similarly, the formal description of the output places yields

$$O_{t_6}(p) = \begin{cases} 1 & \text{if } p = p_0 \\ 0 & \text{otherwise} \end{cases} \quad O_{t_8}(p) = \begin{cases} 2 & \text{if } p = p_1 \\ 0 & \text{otherwise.} \end{cases}$$

In the graphical notation, we do not label arcs that specify input or output places with cardinality 1. In Fig. 8.1(a), the initial marking $m_0 = (1, 0, 0, 0)$ is depicted by the number of tokens in each place. For notational convenience, we specify markings as vectors. $\diamond$

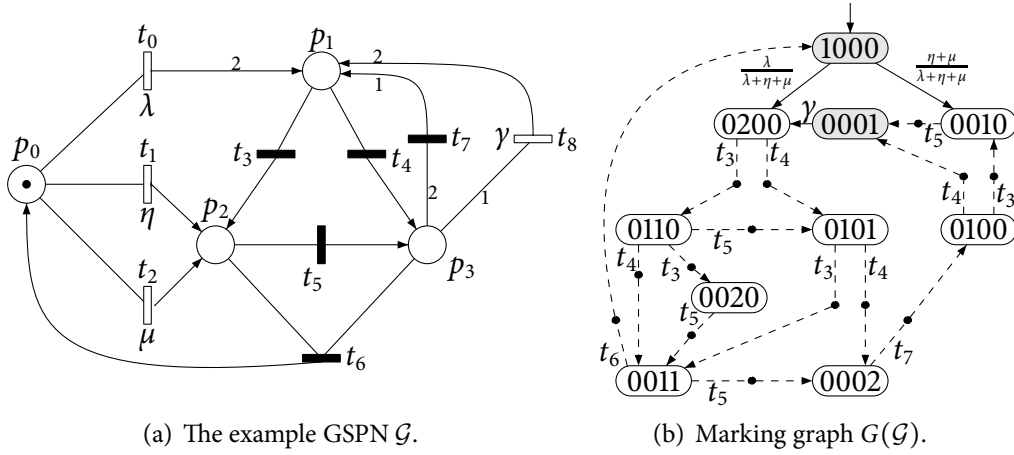


Figure 8.1: A confused GSPN and its induced marking graph.

8.3 A new semantics for GSPNs

The semantics of a GSPN is defined by its *marking graph*, which is informally obtained by playing the “token game”. To define this concept formally, we state the conditions that must be satisfied for a transition to execute:

Definition 8.3 (Concession and enabled transitions). Let $\mathcal{G} = (P, T, \lambda, I, O, H, m_0)$ be a GSPN and $m \in \mathcal{M}$. The set of transitions with concession in marking m is

$$\text{Conc}(m) = \{t \in T \mid \forall p \in P. m(p) \geq I_t(p) \wedge m(p) < H_t(p)\}.$$

The set of enabled transitions in marking m is

$$\text{en}(m) = \begin{cases} \text{Conc}(m) \cap T_i & \text{if } \text{Conc}(m) \cap T_i \neq \emptyset \\ \text{Conc}(m) & \text{otherwise.} \end{cases}$$

We distinguish transitions that have concession from those that are enabled: If a transition has concession in a marking, the number of tokens in its input and inhibition places is such that the transition could execute; however, GSPNs adopt the *maximal progress assumption* which states that immediate transitions take precedence over timed transitions. Therefore, if timed and immediate transitions have concession in a marking m , only the immediate transitions become enabled.

We classify markings according to their enabled transitions: If an immediate transition is enabled in a marking $m \in \mathcal{M}$, the marking changes immediately; we refer to such markings as *vanishing*. Otherwise, if only timed transitions are enabled, we call m a *tangible* marking.

Definition 8.4 (Tangible and vanishing markings). Let $\mathcal{G} = (P, T, \lambda, I, O, H, m_0)$ be a GSPN. A marking $m \in \mathcal{M}$ is vanishing if $\text{en}(m) \cap T_i \neq \emptyset$; otherwise, the marking m is tangible.

In a tangible marking m , only timed transitions are enabled. The residence time in m is then determined by a negative exponential distribution with rate $\sum_{t \in \text{en}(m)} \lambda(t)$. If m is vanishing instead, one of the immediate transitions executes directly, i.e. the sojourn time in m is deterministically zero. In this case, none of the timed transitions which have concession can execute. The effect of executing a transition is formally described by the transition execution relation:

Definition 8.5 (Transition execution). Let $\mathcal{G} = (P, T, \lambda, I, O, H, m_0)$ be a GSPN. We define the transition execution relation $[\cdot] \subseteq \mathcal{M} \times T \times \mathcal{M}$ such that for all markings $m, m' \in \mathcal{M}$ and transitions $t \in T$ it holds:

$$m[t]m' \iff t \in \text{en}(m) \wedge \forall p \in P. m'(p) = m(p) - I_t(p) + O_t(p).$$

Two markings m and m' are in the one-step successor relation $\rightsquigarrow_{\text{GSPN}}$ (denoted $m \rightsquigarrow_{\text{GSPN}} m'$) iff a transition $t \in \text{en}(m)$ exists such that $m[t]m'$ holds. Accordingly, the *reachability set* for marking $m \in \mathcal{M}$ is defined as

$$\text{Reach}(m) = \{m' \in \mathcal{M} \mid m \rightsquigarrow_{\text{GSPN}}^* m'\},$$

where $\rightsquigarrow_{\text{GSPN}}^*$ denotes the reflexive and transitive closure of the relation $\rightsquigarrow_{\text{GSPN}}$.

With Def. 8.5 and the reachability set, we are now ready to define the semantics of a GSPN. It is obtained by successively applying the transition execution relation to generate the (finite or infinite) *marking graph* of the GSPN:

Definition 8.6 (Marking graph). Let $\mathcal{G} = (P, T, \lambda, I, O, H, m_0)$ be a GSPN with immediate transitions in T_i and timed transitions in T_t . Then $\mathcal{G}$ induces the marking graph $\mathcal{M}(\mathcal{G}) = (M, T_i, \hookrightarrow, \dashrightarrow, m_0)$, where

- $M = \text{Reach}(m_0)$ is the set of reachable markings in $\mathcal{G}$,
- $\hookrightarrow \subseteq M \times \mathbb{R}_{>0} \times M$ is the timed transition relation where

$$m \xrightarrow{\mu} m' \iff \mu = \sum \{|\lambda(t)| \mid t \in T_t \wedge m[t]m'\} > 0$$

for all $m, m' \in M$ and $\mu \in \mathbb{R}_{>0}$. Further

- $\dashrightarrow \subseteq M \times \text{Act} \times M$ is the immediate transition relation where for all $m, m' \in M$ and $t \in T_i$ it holds $m \xrightarrow{t} m' \iff m[t]m'$.

Here we use the multiset $\{|\lambda(t)| \mid t \in T_t \wedge m[t]m'\}$ to sum up the rates of all Markovian transitions that lead from marking m to marking m' . As for classical Petri nets, we define the notion of k -boundedness: A GSPN $\mathcal{G}$ with initial marking m_0 is k -bounded iff the number of tokens in each place of all reachable markings is at most k . As a direct consequence, a k -bounded GSPN induces a finite marking graph. We do not discuss the details of determining whether a GSPN is bounded or not, but simply assume that all GSPNs that are intended for our analysis induce a finite marking graph.

Under this assumption, it is straightforward to define the induced IMC of a GSPN by simply interpreting its finite marking graph as an IMC. Informally, the GSPN's immediate transitions correspond to interactive transitions in a closed IMC. Similarly, timed transitions in the GSPN are turned into Markovian transitions in the induced IMC:

Definition 8.7 (Induced IMC). Let $\mathcal{G} = (P, T, \lambda, I, O, H, m_0)$ be a k -bounded GSPN with marking graph $\mathcal{M}(\mathcal{G}) = (M, T_i, \hookrightarrow, \dashrightarrow, m_0)$. Then $\mathcal{G}$ induces the closed IMC $\mathcal{I}(\mathcal{G}) = (\mathcal{S}, \text{Act}, IT, MT, \nu)$ where

- $\mathcal{S} = M$ is the finite set of states,
- $\text{Act} = \text{Act}_i \cup \text{Act}_e$ is the set of actions, where $\text{Act}_e = \emptyset$ and $\text{Act}_i = T_i$,
- $IT \subseteq \mathcal{S} \times \text{Act} \times \mathcal{S}$ with $(m, t, m') \in IT \iff m \xrightarrow{t} m'$ for $m, m' \in M$ and $t \in T_i$,
- $MT \subseteq \mathcal{S} \times \mathbb{R}_{>0} \times \mathcal{S}$ with $(m, \mu, m') \in MT \iff m \xrightarrow{\mu} m'$ for $m, m' \in M$ and $\mu \in \mathbb{R}_{>0}$ and
- $\nu = \{m_0 \mapsto 1\}$.

Stochastic Petri nets (SPNs) form a strict subclass of GSPNs which have a precisely defined semantics [Nat80, Mol81, Mol82]: Each marking in an SPN corresponds to a state of a CTMC; the set of enabled transitions in each marking determine the transition in the CTMC, where the rates of those SPN transitions that lead to the same successor marking are cumulated.

Corollary 8.1. The semantics of GSPN given in Def. 8.7 conservatively extends SPN.

Proof. Follows immediately by the definition of the SPN semantics in [Mol82]. $\square$

Hence, our definition of GSPNs is a conservative extension of stochastic Petri nets. However, our proposed semantics is different to that of [MCB84, CMBC93], as we do not permit to augment immediate transitions with weights but interpret the race between immediate transitions nondeterministically.

This allows us to define a semantics for all GSPNs. In particular, we do not have to restrict to well-defined GSPNs:

Example 8.3. Consider the GSPN $\mathcal{G}$ depicted in Fig. 8.1(a) and its marking graph $G(\mathcal{G})$ in Fig. 8.1(b). According to [TF03, Sec. 2.4], $\mathcal{G}$ is not well-defined: In marking $(0, 0, 1, 1)$, the set of reachable tangible markings is $\{(1, 0, 0, 0), (0, 0, 0, 1)\}$.

If t_5 is chosen, the tangible marking $(0, 0, 0, 1)$ is reached with probability 1; however, if t_6 is chosen, we enter the tangible marking $(1, 0, 0, 0)$ with probability 1. Hence, the distribution over next stable markings depends on the way, the nondeterminism in $(0, 0, 1, 1)$ is resolved. $\diamond$

In the next section, we model a dependable workstation cluster as a GSPN. As we will see, this GSPN model contains nondeterministic choices which correspond to the different strategies to repair failed components within the cluster.

8.4 Dependability analysis of a workstation cluster

In this section, we present our results for the analysis of a dependable workstation cluster which is modeled by a GSPN [HHK00]. The setting is depicted in Fig. 8.2: We consider two identical subclusters, each of which consists of $N \in \mathbb{N}_{>0}$ workstations that are interconnected by a switch. Moreover, via their switches and a central backbone, the workstations in the two subclusters can communicate with each other. For the dependability analysis, we use the failure rates of the components which are given in [HHK00] and restated in Table 8.1.

For our verification, we model the workstation cluster as the GSPN depicted in Fig. 8.3. The first two rows represent the N workstations in the left and right subcluster, respectively. Each single workstation fails after $500h$ of operation, on average. Hence, we associate a failure rate of $\frac{1}{500}$ to each workstation. Accordingly, the timed transitions *LeftWorkstationFail* and *RightWorkstationFail* are marking dependent: If n tokens are in place *LeftWorkstationUp*, each of them fails with rate $\frac{1}{500}$. Therefore, the timed transition *LeftWorkstationFail* has rate $\frac{n}{500}$. The same reasoning applies for *RightWorkstationFail*.

Once a component has failed, a single repair unit is available that can repair one failed component at a time. Depending on the type of component, the repair operation takes

event	duration	event	duration
<i>LeftWorkstationFail</i>	500h	<i>LeftWorkstationRepair</i>	0.5h
<i>RightWorkstationFail</i>	500h	<i>RightWorkstationRepair</i>	0.5h
<i>LeftSwitchFail</i>	4000h	<i>LeftSwitchRepair</i>	4h
<i>RightSwitchFail</i>	4000h	<i>RightSwitchRepair</i>	4h
<i>BackboneFail</i>	5000h	<i>BackboneRepair</i>	8h

Table 8.1: Average durations for component failures and repairs.

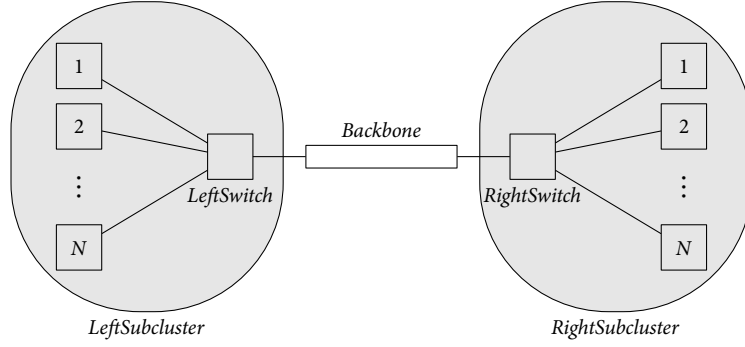


Figure 8.2: A dependable workstation cluster with $2N$ workstations [HHK00].

different average times, cf. Tab. 8.1.

Note that the GSPN model in Fig. 8.3 is confused: Whenever the repair unit is available and different components have failed, the choice which component to repair next is non-deterministic. In the GSPN model, this nondeterminism is represented by the immediate transitions *LeftWorkstationInspect*, *RightWorkstationInspect*, *LeftSwitchInspect*, etc.

By applying Def. 8.7, the GSPN model induces an IMC. As reported in [HHK00], the resulting state space of the IMC consists of 820 states if $N = 4$ and 2772 states for $N = 8$. In our prototypical implementation, we use bisimulation minimization on the obtained IMC to reduce the size of the state space. As can be seen in Table 8.2, the symmetry in the GSPN model yields enormous state space reductions in the bisimulation quotient. They are further amplified by the fact that for a time-bounded reachability analysis, we can make all goal states absorbing before computing the bisimulation quotient.

In the following, we analyze two of the dependability measures that are mentioned in [HHK00]. Therefore, we describe the minimum quality of service (QoS) criterion of a workstation cluster with $2N$ workstations by the number $k \in \{2, 3, \dots, 2N\}$ of workstations that are required to be operational and mutually connected.

For example, if $N = 4$ and $k = 5$, at least 5 of the 8 workstations must be up. Moreover, they must be able to communicate with each other; hence, satisfying the QoS criterion $k = 5$ implies that both switches and the backbone are operational.

For a marking $m \in M$ (which corresponds to a state $s \in \mathcal{S}$ of the IMC), let

$$\begin{aligned} \text{left}_k(m) &= m(\text{LeftSwitchUp}) > 0 \wedge m(\text{LeftWorkstationUp}) \geq k \\ \text{right}_k(m) &= m(\text{RightSwitchUp}) > 0 \wedge m(\text{RightWorkstationUp}) \geq k \\ \text{conn}(m) &= m(\text{LeftSwitchUp}) > 0 \wedge m(\text{RightSwitchUp}) > 0 \wedge m(\text{BackboneUp}) > 0 \\ \text{shared}_k(m) &= m(\text{LeftWorkstationUp}) + m(\text{RightWorkstationUp}) \geq k \wedge \text{conn}(m). \end{aligned}$$

With these definitions, we can assign an atomic propositions min_k to all states $s \in \mathcal{S}$ which correspond to a marking that meets the QoS requirement in the underlying GSPN:

$$\text{min}_k \in L(s) \iff \text{left}_k(s) \vee \text{right}_k(s) \vee \text{shared}_k(s).$$

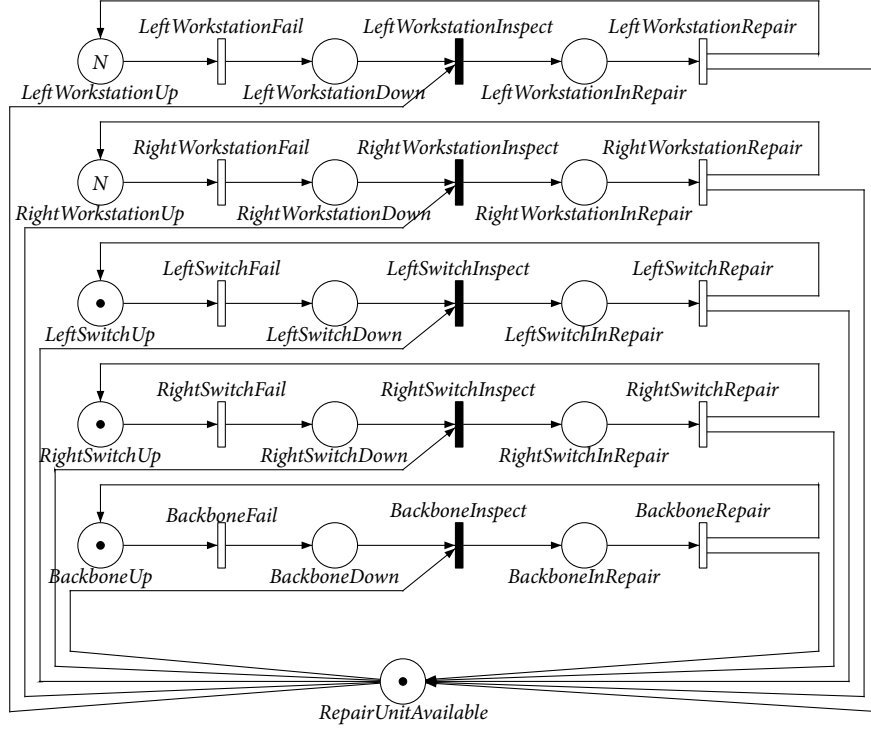


Figure 8.3: GSPN model of the fault tolerant workstation cluster [HHK00].

We analyze the following dependability measures for different parameters N and k :

1. “The chance that the QoS constraint k is violated within the next z time units is less than p ”:

This measure corresponds to the maximum time-bounded reachability probability for the set of goal states $S_{bad} = \{s \in \mathcal{S} \mid s \not\models \min_k\}$ that violate the QoS constraint k . It is formalized by the CSL state formula Φ_4 taken from [HHK00]:

$$\Phi_4 = \mathcal{P}_{\leq p}(\diamond^{\leq z}(\neg \min_k)).$$

To model check $s \models \Phi_4$, it suffices to compute

$$p_4(s) = \sup_{D \in GM} Pr_{v_s, D}^{\omega}(\diamond^{[0, z]} S_{bad})$$

and to decide whether $p_4(s) \leq p$ holds. In this section, we aim at computing the actual least upper bound on the achievable probability. Therefore, Table 8.2 lists the values $p_4(s)$ instead of the truth values for $s \models \Phi_4$.

We compute the probability p_4 for two different markings: The state s_{opt} denotes the marking where all components of the cluster are operational. On the other hand,

N	k	states	quot. states	z	measure	results		time	
						IMC	PRISM	IMC	PRISM
4	3	820	129	1000	$p_4(s_{opt})$	0.0009	0.0009	104h	73s
4	5	820	8	1000		0.5034	0.5034	3.1h	10s
8	8	2772	703	200		0.0076	0.0076	2.7h	18s
8	10	2772	14	100		0.0676	0.0676	196s	3s
8	10	2772	14	1000		0.5034	0.5034	5.3h	33s
4	3	820	130	1000	$p_4(s_{crit})$	0.0834	0.0437	91h	75s
8	8	2772	142	200		0.2275	0.1876	3.2h	18s
8	10	2772	15	200		0.1393	0.1393	2.2h	7s
4	3	820	424	20	$\max_{s \in S_{bad}} p_5(s)$	0.3797	0.3038	304s	4s
4	5	820	164	20		0.4219	0.3717	90s	4s
4	8	820	164	20		0.4278	0.4250	15m	4s
8	3	2772	1412	10		0.9319	0.7457	277s	6s
8	10	2772	316	10		0.9805	0.9178	45s	7s
8	16	2772	316	20		0.6147	0.6089	36m	123s

Table 8.2: Results of the dependability analysis.

s_{crit} is a marking with the minimum number of working components to satisfy the QoS constraint k . For example, if $N = 4$ and $k = 3$, s_{crit} is the state where k workstations and the switch of the left (or right) subcluster are working, whereas all other components have failed. Hence s_{crit} barely fulfills the QoS requirements.

2. “If the QoS constraint k is violated, the probability to face the same problem after z time units is less than p ”:

This measure corresponds to a time-interval bounded reachability probability. For a single state $s \in \mathcal{S}$, it is specified in [HHK00] by the CSL state formula Φ_5 :

$$\Phi_5 = \neg \min_k \rightarrow \mathcal{P}_{\leq p}(\diamond^{[z,z]}(\neg \min_k)).$$

Obviously, all states $s \in (\mathcal{S} \setminus S_{bad})$ satisfy Φ_5 . Therefore, we aim at deciding whether $S_{bad} \models \Phi_5$, where $A \models \Phi_5$ holds iff all states in $A \subseteq \mathcal{S}$ satisfy Φ_5 . Let $p_5(s) = \sup_{D \in GM} Pr_{v_s, D}^\omega(\diamond^{[z,z]} S_{bad})$ be the maximal probability of the event $\diamond^{[z,z]} S_{bad}$, starting from initial state s . Then $\max_{s \in S_{bad}} p_5(s)$ is the desired dependability measure.

Note that in theory (cf. Sec. 6.3.2), we cannot compute the probability in the induced IMC for a point-interval $[z, z]$. Therefore, we approximate the event by using a short time-interval $[z, z+\varepsilon]$, where $\varepsilon = 10^{-5}$.

In the following, we compare the results that we obtain by our prototypical implementation of the GSPN semantics from Sec. 8.3 and the IMC approximation algorithm (Chapter 6) to the probabilities that are obtained by the PRISM model checker [KNP02, HKNP06] on the classical GSPN model with weight specifications as given in [HHK00].

As pointed out earlier, nondeterminism occurs in the workstation cluster whenever different components have failed and the repair unit has to choose which one to repair next. However, PRISM is not capable of analyzing nondeterministic and randomly timed models such as CTMDPs and IMCs. Instead, the nondeterminism in the PRISM model² is resolved by assigning high rates to the immediate transitions. In this way, the GSPN is transformed into a CTMC, which is then analyzed. The outcomes are shown in Table 8.2.

Some remarks concerning this comparison are in order:

In the first block of Table 8.2, the probabilities $p_4(s_{opt})$ that are computed by our implementation of the IMC-based semantics are very close to those obtained by analyzing the weighted GSPN model.

This is no longer true if we consider the initial state s_{crit} : Here, the worst case probabilities in the nondeterministic GSPN semantics are approximately 4% higher than those obtained by the weighted GSPN, which resolves the nondeterminism by equi-probability. This is explained as follows:

Only k workstations and the left switch remain operational in state s_{crit} . In this situation, the scheduling strategy for the *RepairUnit* matters: In the worst case, all faulty workstations in the right subcluster are repaired first; however, as long as the right switch and the backbone are defective, this does not improve the dependability probability. The uniform probability distribution used in classical GSPN model does not reflect this worst case scenario, effectively producing *false positives*.

This phenomenon is not observed for initial state s_{opt} , as the probability to reach a state such as s_{crit} that is badly degraded, is extremely low. As the repair time is short compared to the failure rate, only states with few failed components occur with considerable probability. Therefore, the degree of nondeterminism is low for initial state s_{opt} .

If $k > N$, the QoS constraint is violated as soon as one switch or the backbone fail. Hence, in this case, the strategy of the repair unit does not matter. Accordingly, the results agree for the case $N = 8$, $k = 10$ and initial state s_{crit} .

For Φ_5 , the dependability measures differ considerably: In the worst case, the dependability is 18% worse than predicted by the classical GSPN model. This difference is explained as follows:

Assume that s_{down} is the state where both switches, the backbone and all N workstations in the right subcluster have failed, whereas in the left subcluster, all workstations are operational. To compute $p_5(s_{down})$, we have to select the worst schedule possible. Therefore, note that if $k \leq N$, repairing the left switch establishes QoS. Thus, the desired worst case probability is obtained if all workstations in the right subcluster are repaired — which does not establish QoS — before the left switch.

However, in the classical GSPN model, each immediate transition has weight 1; therefore, the probability to repair the switch in the otherwise intact left subcluster is $\frac{1}{5}$. Obviously, this implicit strategy does not reflect the worst case scenario, which is needed to

²The source code of the PRISM model is available online on the PRISM website:
<http://www.prismmodelchecker.org/casestudies/cluster.php>

decide Φ_5 .

Again, no difference occurs if $k = 2N$: In this case, all components must be operational in order to satisfy QoS. Hence, the scheduler is irrelevant and the resulting probabilities coincide (up to rounding errors).

Further, note that our prototypical implementation is not optimized yet; for example, it relies on an arbitrary precision floating point library (the MPFR library) that does not make use of the underlying floating point hardware. Therefore, it is realistic to expect improvements in the performance of our model checking tool. All measurements were carried out on a 2.2GHz Xeon CPU with 16GB RAM.

In [Joh07], the dependable workstation cluster [HHK00] has been modeled as an IMC, directly. More precisely, the IMC model is obtained by composing (untimed) labeled transition systems that model the cluster's components with corresponding time constraints that are specified as IMCs (see [Joh07, Fig. 10.3]). The approach taken in [Joh07] is to transform the composed IMC model into a globally uniform CTMDP which is then subject to a time-bounded reachability probability analysis. In order to obtain a globally uniform CTMDP, the approach relies on the assumption that the underlying IMC is globally uniform, as well. From a modeling point of view, this is not the case in the workstation cluster. Hence, to still achieve global uniformity, the time-constraints that are weaved into the IMC model in [Joh07] are uniformized *before* the composition. In this way, the resulting IMC is globally uniform; however, it contains self-loops that are introduced artificially by the uniformization of the time-constraints [Joh07, Fig. 10.4].

In contrast to our results, [Joh07] computes time-bounded reachability probabilities for time-abstract scheduler classes. However, as shown before in [BHKH05] and in Sec. 4.3, the implicit uniformization that is used in [Joh07] is not measure preserving for the class of time-abstract schedulers: Intuitively, a history dependent but time-abstract scheduler can estimate the amount of time that has passed by observing which states have been visited. Introducing artificial self loops as done in [Joh07] exposes additional information to such schedulers: By counting the number of times such a self loop is taken, the otherwise time-abstract scheduler can improve (as proved in [BHKH05] and in Sec. 4.3) its decisions considerably. Thus it may exploit the structural changes in the CTMDP that are induced by uniformization. Due to these differences, the results of [Joh07] are not directly comparable to ours.

As expected from a theoretical point of view, all probabilities that are computed in our IMC model are larger or equal to those that are obtained by the PRISM model. This stands in contrast to the surprising result in [Joh07, p. 187], where the probabilities that are obtained by analyzing the CTMC model are larger than those of the IMC model [Joh07, Sec. 10.1.3]. The reason for this phenomenon remains unclear; however, our results do not support the claim in [Joh07] that imprecisions in the PRISM model lead to probabilities that are too large.

8.5 Conclusion

Motivated by the development of our approximation algorithm for the analysis of IMCs (cf. Chapter 6), we propose a nondeterministic semantics for generalized stochastic Petri nets and omit the weight-specification that has been used in the classical GSPN definitions. In this way, all static (qualitative) analyses such as k -boundedness, reachability and coverability are also applicable to our modified definition of GSPNs.

It remains an interesting question for future research to apply the results in this chapter to analyse the compositional extensions of GSPN models that are proposed in [HHMR97]. When [HHMR97] was published, the analysis of compositional GSPNs was restricted to deterministic instances. We expect that applying the results of this chapter to the compositional modeling framework permits the analysis of a much broader class of compositional GSPNs.

If a GSPN is k -bounded, it induces a closed IMC with a finite state space on which important performance and dependability measures can be computed.

We apply our definition to a case study from the literature and compare the results of our technique to those that are obtained by the classical weighted GSPN semantics. Thereby it turns out, that the reliability estimates that are obtained by analyzing the classical GSPN model are up to 18% higher than those that might actually occur.

These false positives clearly prove that nondeterministic modeling is essential in the area of dependability analysis.

9 Conclusion

When my supervisor Joost-Pieter introduced me to CTMDPs, I hardly had a background in stochastic modeling. However, with his guidance and our joint research on bisimulation minimization for CTMDPs, I slowly got more confident in my understanding of stochastic processes and probability & measure theory. The results of this early work are the definition of bisimulation for CTMDPs in Chapter 7 and the proof that it preserves not only CSL, but all quantitative measures.

In the sequel, I gave a talk on this topic at the University of Twente, when Mariëlle asked an elementary question: *“Wouldn’t it be better for the scheduler if it was allowed to decide later, when the state is actually left?”*

The subsequent research of Mariëlle, Joost-Pieter and myself led to the results in Chapter 4, where we study a hierarchy of scheduler classes and characterize their relationships. Our motivation was to delay the scheduling decisions in CTMDPs. Therefore, we investigated local uniformity and defined late schedulers. In retrospect, the latter turned out to be the most influential idea for the achievements in this thesis.

When I visited his group in Saarbrücken, Holger asked me to give a talk about local uniformity and late schedulers. The following discussion with Lijun was the most revealing of my entire PhD time. When we were finished, we had sketched the discretization for locally uniform CTMDPs which is the basis of the time-bounded reachability analysis in Chapter 5. In the following months, we proved that our approximation is quantifiably correct, that is, it determines the maximal or minimal reachability probability in a locally uniform CTMDP up to an error which can be made arbitrarily small.

This result encouraged further research: We adapted the idea behind our discretization technique to IMCs and extended it to also account for lower time-interval bounds. The result is the first model checking algorithm for CSL on IMCs. It is presented in Chapter 6.

At roughly the same time, Holger, Lijun, Sven and I discussed about a new semantics for GSPNs. However, at that time, no model checking algorithms were available that would have made our proposal attractive to a broader audience. Luckily, this has changed by now: With the achievements in Chapters 5 and 6, we are able to model check nondeterministic GSPNs. This is the topic of Chapter 8 that proposes a new semantics for GSPNs that overcomes the shortcomings in modeling nondeterminism of the former approaches. By means of a case study which considers dependability characteristics of a workstation cluster, we show that nondeterministic modeling indeed makes a difference: As it turns out, earlier reliability predictions that were obtained in the classical GSPN semantics are up to 18% too optimistic. These false positives clearly prove the necessity of analyzing nondeterministic and randomly timed systems.

To conclude the thesis, we summarize our achievements and propose directions for future research:

- We define a hierarchy of time-dependent scheduler classes and investigate their expressive power. Moreover, we propose local uniformization and identify the scheduler classes for which it is measure preserving. This culminates in the discovery of late schedulers that are more expressive than the scheduler classes considered previously and in the literature.

Future research: The definition of late schedulers is limited to locally uniform CTMDPs. To bridge this gap and to define corresponding schedulers for arbitrary CTMDPs is an important further step. In the same context, the question whether local uniformization is measure preserving w.r.t. such a new scheduler definition is another interesting starting point for future research.

- We develop an efficient and quantifiably precise algorithm that computes time bounded reachability probabilities in locally uniform CTMDPs with respect to time- and history-dependent late schedulers. To the best of our knowledge, this is the first time that such an analysis becomes feasible.

Future research: The definition of late schedulers on arbitrary CTMDPs is an open problem. We believe that in combination with the results on local uniformization from Chapter 4, such a definition will allow us to model check non-locally uniform CTMDPs with respect to late schedulers.

- Along similar lines, we derive a model checking algorithm that verifies a broad class of CSL formulas on IMCs. It is the first algorithm that is not restricted to specific subclasses but enables the analysis of arbitrary IMCs.

Future research: Model checking long run average properties and specific instances of until formulas remain unsolved problems which must be tackled.

- We introduce strong bisimulation minimization for CTMDPs and prove that it preserves all quantitative measures. Moreover, we define CSL on CTMDPs and prove its measure theoretic soundness.

Future research: Chapter 7 is based on time- and history dependent schedulers. It is an open question whether its results also apply to less powerful schedulers. Considering action-based variants of CSL is another promising approach to obtain a logical characterization for strong bisimilarity.

- We define a new semantics for GSPNs that allows nondeterminism to occur in the model. Via a transformation which turns a GSPN into an equivalent IMC, we can model check CSL formulas on GSPNs. Finally, we show the applicability of this approach by means of a larger case study.

Bibliography

- [ADD00] R. B. Ash and C. A. Doléans-Dade. *Probability & Measure Theory*. Academic Press, 2nd edition, 2000.
- [And02] S. Andova. *Probabilistic Process Algebra*. PhD thesis, Eindhoven University of Technology, Eindhoven, The Netherlands, 2002.
- [ASSB96] A. Aziz, K. Sanwal, V. Singhal, and R. K. Brayton. Verifying continuous time Markov chains. In *Proceedings of the 8th International Conference on Computer Aided Verification (CAV)*, volume 1102 of *Lecture Notes in Computer Science*, pages 269–276. Springer, 1996.
- [ASSB00] A. Aziz, K. Sanwal, V. Singhal, and R. K. Brayton. Model-checking continuous-time Markov chains. *ACM Transactions on Computational Logic (TOCL)*, 1(1):162–170, 2000.
- [Baa08] S. Baase. *A Gift of Fire: Social, Legal and Ethical Issues for Computing and the Internet*. Pearson Education, 3rd edition, 2008.
- [Bai98] C. Baier. On Algorithmic Verification Methods for Probabilistic Systems. Habilitation Thesis, 1998. University of Mannheim.
- [Bal00] G. Balbo. Introduction to stochastic Petri nets. In *European Educational Forum: School on Formal Methods and Performance Analysis*, volume 2090 of *Lecture Notes in Computer Science*, pages 84–155. Springer, 2000.
- [Bal07] G. Balbo. Introduction to generalized stochastic Petri nets. In *7th International School on Formal Methods for the Design of Computer, Communication, and Software Systems*, volume 4486 of *Lecture Notes in Computer Science*, pages 83–131. Springer, 2007.
- [BCH⁺08] H. Boudali, P. Crouzen, B. R. Haverkort, M. Kuntz, and M. I. A. Stoelinga. Architectural dependability evaluation with Arcade. In *38th Annual International Conference on Dependable Systems and Networks (DSN)*, pages 512–521. IEEE Computer Society, 2008.
- [BCS07] H. Boudali, P. Crouzen, and M. I. A. Stoelinga. Dynamic fault tree analysis using input/output interactive Markov chains. In *Proceedings of the 37th Annual International Conference on Dependable Systems and Networks (DSN)*. IEEE Computer Society, 2007.

- [BdA95] A. Bianco and L. de Alfaro. Model checking of probabilistic and nondeterministic systems. In *15th Conference on Foundations of Software Technology and Theoretical Computer Science (FSTTCS)*, volume 1026 of *Lecture Notes in Computer Science*, pages 499–513. Springer, 1995.
- [BDF81] J. L. Bruno, P. J. Downey, and G. N. Frederickson. Sequencing tasks with exponential service times to minimize the expected flow time or makespan. *Journal of the ACM*, 28(1):100–113, 1981.
- [Bel57] R. E. Bellman. A Markovian decision process. *Indiana University Mathematics Journal*, 6(4):679–684, 1957.
- [Ben76] J. Benedetto. *Real Variable and Integration*. Teubner Verlag, 1976.
- [Ber95] D. Bertsekas. *Dynamic Programming and Optimal Control*, volume II. Athena Scientific, 1995.
- [BF09] P. Bouyer and V. Forejt. Reachability in stochastic timed games. In *36th International Colloquium on Automata, Languages and Programming (ICALP), Part II*, volume 5556 of *Lecture Notes in Computer Science*, pages 103–114. Springer, 2009.
- [BFK⁺09] T. Brázdil, V. Forejt, J. Krcal, J. Kretinsky, and A. Kucera. Continuous-time stochastic games with time-bounded reachability. In *Annual Conference on Foundations of Software Technology and Theoretical Computer Science (FSTTCS)*, volume 4 of *Leibniz International Proceedings in Informatics*, pages 61–72. Schloss Dagstuhl–Leibniz-Zentrum für Informatik, Germany, 2009.
- [BG98] M. Bernardo and R. Gorrieri. A tutorial on EMPA: A theory of concurrent processes with nondeterminism, priorities, probabilities and time. *Theoretical Computer Science*, 202(1-2):1–54, 1998.
- [BG01] M. Bernardo and R. Gorrieri. Corrigendum to “A tutorial on EMPA: A theory of concurrent processes with nondeterminism, priorities, probabilities and time” - [TCS 202 (1998) 1–54]. *Theoretical Computer Science*, 254(1-2):691–694, 2001.
- [BHH⁺09] E. Böde, M. Herbstritt, H. Hermanns, S. Johr, T. Peikenkamp, R. Pulungan, J. Rakow, R. Wimmer, and B. Becker. Compositional dependability evaluation for STATEMATE. *IEEE Transactions on Software Engineering*, 35(2):274–292, 2009.
- [BHHK03] C. Baier, B. R. Haverkort, H. Hermanns, and J.-P. Katoen. Model-checking algorithms for continuous-time Markov chains. *IEEE Transactions on Software Engineering*, 29(6):524–541, 2003.

- [BHHK04] C. Baier, B. R. Haverkort, H. Hermanns, and J.-P. Katoen. Nonuniform CTMDPs. Unpublished manuscript, 2004.
- [BHK06] M. Bravetti, H. Hermanns, and J.-P. Katoen. YMCA: Why Markov chain algebra? In *Proceedings of the Workshop Essays on Algebraic Process Calculi*, volume 162 of *Electronic Notes in Theoretical Computer Science*, pages 107–112. Elsevier, 2006.
- [BHKH05] C. Baier, H. Hermanns, J.-P. Katoen, and B. R. Haverkort. Efficient computation of time-bounded reachability probabilities in uniform continuous-time Markov decision processes. *Theoretical Computer Science*, 345(1):2–26, 2005.
- [Bil95] P. Billingsley. *Probability and Measure*. John Wiley & Sons, 3rd edition, 1995.
- [BK98] C. Baier and M. Z. Kwiatkowska. Model checking for a probabilistic branching time logic with fairness. *Distributed Computing*, 11(3):125–155, 1998.
- [BK08] C. Baier and J.-P. Katoen. *Principles of Model Checking*. MIT Press, 1st edition, 2008.
- [BKHW05] C. Baier, J.-P. Katoen, H. Hermanns, and V. Wolf. Comparative branching-time semantics for Markov chains. *Information and Computation*, 200(2):149–214, 2005.
- [BR04] M. R. Blaha and J. R. Rumbaugh. *Object-Oriented Modeling and Design with UML*. Prentice Hall, 2nd edition, 2004.
- [Buc94] P. Buchholz. Exact and ordinary lumpability in finite Markov chains. *Journal of Applied Probability*, 31(1):59–75, 1994.
- [BW90] J. C. M. Baeten and W. P. Weijland. *Process Algebra*, volume 18. Cambridge University Press, 1990.
- [CCGR00] A. Cimatti, E. M. Clarke, F. Giunchiglia, and M. Roveri. NuSMV: A new symbolic model checker. *International Journal on Software Tools for Technology Transfer (STTT)*, 2(4):410–425, 2000.
- [CDF91] G. Chiola, S. Donatelli, and G. Franceschinis. GSPNs versus SPNs: What is the actual role of immediate transitions? In *Proceedings of the 4th International Workshop on Petri Nets and Performance Models (PNPM)*, pages 20–31. IEEE Computer Society, 1991.

- [CDHS06] D. Cerotti, S. Donatelli, A. Horváth, and J. Sproston. CSL model checking for generalized stochastic Petri nets. In *3rd International Conference on the Quantitative Evaluation of Systems (QEST)*, pages 199–210. IEEE Computer Society, 2006.
- [CES86] E. M. Clarke, E. A. Emerson, and A. P. Sistla. Automatic verification of finite-state concurrent systems using temporal logic specifications. *ACM Transactions on Programming Languages and Systems (TOPLAS)*, 8(2):244–263, 1986.
- [CG89] A. E. Conway and N. D. Georganas. *Queueing networks—exact computational algorithms: a unified theory based on decomposition and aggregation*. MIT Press, 1989.
- [CGH⁺08] N. Coste, H. Garavel, H. Hermanns, R. Hersemeule, Y. Thonnart, and M. Zidouni. Quantitative evaluation in embedded system design: Validation of multiprocessor multithreaded architectures. In *Design, Automation and Test in Europe (DATE)*, pages 88–89. IEEE Computer Society, 2008.
- [CHLS09] N. Coste, H. Hermanns, E. Lantreibeacq, and W. Serwe. Towards performance prediction of compositional models in industrial GALS designs. In *Proceedings of the 21st International Conference on Computer Aided Verification (CAV)*, volume 5643, pages 204–218. Springer, 2009.
- [CMBC93] G. Chiola, M. A. Marsan, G. Balbo, and G. Conte. Generalized stochastic Petri nets: A definition at the net level and its implications. *IEEE Transactions on Software Engineering*, 19(2):89–107, 1993.
- [CW87] D. Coppersmith and S. Winograd. Matrix multiplication via arithmetic progressions. In *Proceedings of the 19th Annual ACM Symposium on Theory of Computing (STOC)*. ACM Press, 1987.
- [dA97] L. de Alfaro. *Formal Verification of Probabilistic Systems*. PhD thesis, Stanford University, 1997.
- [DJJL01] P. R. D’Argenio, B. Jeannet, H. E. Jensen, and K. G. Larsen. Reachability analysis of probabilistic systems by successive refinements. In *Proceedings of the Joint International Workshop on Process Algebra and Probabilistic Methods, Performance Modeling and Verification (PAPM-PROBMIV)*, volume 2165 of *Lecture Notes in Computer Science*, pages 39–56. Springer, 2001.
- [DP03] J. Desharnais and P. Panangaden. Continuous stochastic logic characterizes bisimulation of continuous-time Markov processes. *Journal of Logic and Algebraic Programming*, 56(1–2):99–115, 2003.

- [Dud02] R. M. Dudley. *Real Analysis and Probability*. Cambridge University Press, 2002.
- [GHLPR06] X. P. Guo, O. Hernández-Lerma, and T. Prieto-Rumeau. A survey of recent results on continuous-time Markov decision processes. *TOP*, 14:177–261, 2006.
- [GHR93] N. Götz, U. Herzog, and M. Rettelbach. Multiprocessor and distributed system design: The integration of functional specification and performance analysis using stochastic process algebras. In *Proceedings of the 16th International Symposium on Computer Performance Modelling, Measurement and Evaluation (PERFORMANCE)*, volume 729 of *Lecture Notes in Computer Science*, pages 121–146. Springer, 1993.
- [GM84] D. Gross and D. R. Miller. The randomization technique as a modeling tool and solution procedure for transient Markov processes. *Operations Research*, 32(2):343–361, 1984.
- [Gra91] W. K. Grassmann. Finding transient solutions in Markovian event systems through randomization. In *Numerical Solutions of Markov Chains*, pages 357–371. Marcel Dekker, 1991.
- [Har87] D. Harel. Statecharts: a visual formalism for complex systems. *Science of Computer Programming*, 8(3):231–274, 1987.
- [Hav98] B. R. Haverkort. *Performance of Computer Communication Systems: A Model-Based Approach*. John Wiley & Sons, 1998.
- [Hav00] B. R. Haverkort. Markovian models for performance and dependability evaluation. In *European Educational Forum: School on Formal Methods and Performance Analysis*, volume 2090 of *Lecture Notes in Computer Science*, pages 38–83. Springer, 2000.
- [Her02] H. Hermanns. *Interactive Markov Chains: And the Quest for Quantified Quality*, volume 2428 of *Lecture Notes in Computer Science*. Springer, 2002.
- [HHK00] B. R. Haverkort, H. Hermanns, and J.-P. Katoen. On the use of model checking techniques for dependability evaluation. In *Proceedings of the 19th Symposium on Reliable Distributed Systems (SRDS)*, pages 228–237. IEEE Computer Society, 2000.
- [HHK02] H. Hermanns, U. Herzog, and J.-P. Katoen. Process algebra for performance evaluation. *Theoretical Computer Science*, 274(1-2):43–87, 2002.

- [HHMR97] H. Hermanns, U. Herzog, V. Mertsiotakis, and M. Rettelbach. Exploiting stochastic process algebra achievements for generalized stochastic Petri nets. In *Proceedings of the 7th International Workshop on Petri Nets and Performance Models (PNPM)*, pages 183–192. IEEE Computer Society, 1997.
- [Hil96] J. Hillston. *A Compositional Approach to Performance Modelling*. Cambridge University Press, 1996.
- [HJ94] H. Hansson and B. Jonsson. A logic for reasoning about time and reliability. *Formal Aspects of Computing*, 6(5):512–535, 1994.
- [HJ07] H. Hermanns and S. Johr. Uniformity by construction in the analysis of nondeterministic stochastic systems. In *37th Annual International Conference on Dependable Systems and Networks (DSN)*, pages 718–728. IEEE Computer Society, 2007.
- [HK00] H. Hermanns and J.-P. Katoen. Automated compositional Markov chain generation for a plain-old telephone system. *Science of Computer Programming*, 36(1):97–127, 2000.
- [HKNP06] A. Hinton, M. Z. Kwiatkowska, G. Norman, and D. Parker. PRISM: A tool for automatic verification of probabilistic systems. In *Proceedings of the 12th International Conference on Tools and Algorithms for the Construction and Analysis of Systems (TACAS)*, volume 3920 of *Lecture Notes in Computer Science*, pages 441–444. Springer, 2006.
- [Hoa85] C. A. R. Hoare. *Communicating Sequential Processes*. Prentice Hall, 1985.
- [Hol04] G. J. Holzmann. *The SPIN Model Checker: Primer and Reference Manual*. Addison-Wesley, 2004.
- [How60] R. A. Howard. *Dynamic Programming and Markov Process*. MIT Press, 1960.
- [How71] R. A. Howard. *Dynamic Probabilistic Systems*. John Wiley & Sons, 1971.
- [IR90] A. Itai and M. Rodeh. Symmetry breaking in distributed networks. *Information and Computation*, 88(1):60–87, 1990.
- [Jan03] D. N. Jansen. *Extensions of Statecharts with Probability, Time, and Stochastic Timing*. PhD thesis, University of Twente, Enschede, 2003.
- [Jen53] A. Jensen. Markov chains as an aid in the study of Markov processes. *Skandinavisk Aktuarietidskrift*, 3:87–91, 1953.

- [Joh07] S. Johr. *Model Checking Compositional Markov Systems*. PhD thesis, Saarland University, Saarbrücken, Germany, 2007.
- [Kan91] V. G. Kanovei. Cardinality of the set of Vitali equivalence classes. *Mathematical Notes*, 49(4):55–62, 1991.
- [KKN09] J.-P. Katoen, D. Klink, and M. R. Neuhäuser. Compositional abstraction for stochastic systems. In *Proceedings of the 7th International Conference on Formal Modeling and Analysis of Timed Systems (FORMATS)*, volume 5813 of *Lecture Notes in Computer Science*, pages 195–211. Springer, 2009.
- [KNP02] M. Z. Kwiatkowska, G. Norman, and D. Parker. PRISM: Probabilistic symbolic model checker. In *Proceedings of the 12th International Conference on Computer Performance Evaluation, Modelling Techniques and Tools (TOOLS)*, volume 2324 of *Lecture Notes in Computer Science*, pages 200–204. Springer, 2002.
- [KS76] J. G. Kemeny and J. L. Snell. *Finite Markov Chains*. Springer, 1976.
- [Kul95] V. G. Kulkarni. *Modeling and Analysis of Stochastic Systems*. Chapman & Hall, 1995.
- [KZH⁺09] J.-P. Katoen, I. S. Zapreev, E. M. Hahn, H. Hermanns, and D. N. Jansen. The ins and outs of the probabilistic model checker MRMC. In *6th International Conference on the Quantitative Evaluation of Systems (QEST)*, pages 167–176. IEEE Computer Society, 2009.
- [LHK01] G. G. Infante López, H. Hermanns, and J.-P. Katoen. Beyond memoryless distributions: Model checking semi-Markov chains. In *Proceedings of the Joint International Workshop on Process Algebra and Probabilistic Methods, Performance Modeling and Verification (PAPM-PROBMIV)*, volume 2165 of *Lecture Notes in Computer Science*, pages 57–70. Springer, 2001.
- [LS91] K. G. Larsen and A. Skou. Bisimulation through probabilistic testing. *Information and Computation*, 94(1):1–28, 1991.
- [MBC⁺91] M. A. Marsan, G. Balbo, G. Chiola, G. Conte, S. Donatelli, and G. Franceschinis. An introduction to generalized stochastic Petri nets. *Microelectronics and Reliability*, 31(4):699–725, 1991.
- [MBC⁺95] M. A. Marsan, G. Balbo, G. Conte, S. Donatelli, and G. Franceschinis. *Modelling with Generalized Stochastic Petri Nets*. John Wiley & Sons, 1995.
- [MBCC87] M. A. Marsan, G. Balbo, G. Chiola, and G. Conte. Generalized stochastic Petri nets revisited: Random switches and priorities. In *Proceedings of the*

- 2nd International Workshop on Petri Nets and Performance Models (PNPM)*, pages 44–53. IEEE Computer Society, 1987.
- [MCB84] M. A. Marsan, G. Conte, and G. Balbo. A class of generalized stochastic Petri nets for the performance evaluation of multiprocessor systems. *ACM Transactions on Computer Systems (TOCS)*, 2(2):93–122, 1984.
- [MH06a] J. M. Martínez and B. R. Haverkort. CSL model checking of deterministic and stochastic Petri nets. In *Proceedings of the 13th GI/ITG Conference on Measuring, Modelling and Evaluation of Computer and Communication Systems (MMB)*, pages 265–282. VDE Verlag, 2006.
- [MH06b] J. M. Martínez and B. R. Haverkort. MathMC: A Mathematica-based tool for CSL model checking of deterministic and stochastic Petri nets. In *3rd International Conference on the Quantitative Evaluation of Systems (QEST)*, pages 133–134. IEEE Computer Society, 2006.
- [Mil68a] B. L. Miller. Finite state continuous time Markov decision processes with a finite planning horizon. *SIAM Journal of Control and Optimization*, 6(2):266–280, 1968.
- [Mil68b] B. L. Miller. Finite state continuous time Markov decision processes with an infinite planning horizon. *Journal of Mathematical Analysis and Applications*, 22:552–569, 1968.
- [Mil82] R. Milner. *A Calculus of Communicating Systems*. Springer, 1982.
- [Mil99] R. Milner. *Communicating and Mobile Systems: the Pi-Calculus*. Cambridge University Press, 1999.
- [Mol81] M. K. Molloy. *On the integration of delay and throughput measures in distributed processing models*. PhD thesis, University of California, Los Angeles, 1981.
- [Mol82] M. K. Molloy. Performance analysis using stochastic Petri nets. *IEEE Transactions on Computers*, 31(9):913–917, 1982.
- [MP90] R. Mathar and D. Pfeifer. *Stochastik für Informatiker*. Teubner Verlag, 1990.
- [MT06] J. Markovski and N. Trčka. Lumping Markov chains with silent steps. In *3rd International Conference on the Quantitative Evaluation of Systems (QEST)*, pages 221–232. IEEE Computer Society, 2006.
- [MVCR08] H. Maciá, V. Valero, F. Cuartero, and M. C. Ruiz. sPBC: A Markovian extension of Petri box calculus with immediate multiactions. *Fundamenta Informaticae*, 87(3-4):367–406, 2008.

- [Nat80] S. Natkin. *Les réseaux de Petri stochastiques et leur application à l'évaluation des systèmes informatiques*. PhD thesis, Conservatoire National des Arts et Metier, Paris, 1980.
- [NK07] M. R. Neuhäuser and J.-P. Katoen. Bisimulation and logical preservation for continuous-time Markov decision processes. In *Proceedings of the 18th International Conference on Concurrency Theory (CONCUR)*, volume 4703 of *Lecture Notes in Computer Science*, pages 412–427. Springer, 2007.
- [NN07] M. R. Neuhäuser and T. Noll. Abstraction and model checking of Core Erlang programs in Maude. *Electronic Notes in Theoretical Computer Science*, 176(4):147–163, 2007.
- [NSK09] M. R. Neuhäuser, M. Stoelinga, and J.-P. Katoen. Delayed nondeterminism in continuous-time Markov decision processes. In *Proceedings of the 12th International Conference on Foundations of Software Science and Computational Structures (FoSSaCS)*, volume 5504 of *Lecture Notes in Computer Science*, pages 364–379. Springer, 2009.
- [NZ09] M. R. Neuhäuser and L. Zhang. Time-bounded reachability in continuous-time Markov decision processes. Technical report, RWTH Aachen University, 2009.
- [Pnu77] A. Pnueli. The temporal logic of programs. In *Proceedings of the 18th Annual Symposium on Foundations of Computer Science (FOCS)*, pages 46–57. IEEE Computer Society, 1977.
- [Pul09] R. Pulungan. *Reduction of Acyclic Phase-Type Representations*. PhD thesis, Universität des Saarlandes, Saarbrücken, Germany, 2009.
- [Put94] M. L. Puterman. *Markov Decision Processes: Discrete Stochastic Dynamic Programming*. John Wiley & Sons, 1994.
- [Rei85] W. Reisig. *Petri nets: An introduction*. Springer, 1985. ISBN 0-387-13723-8.
- [Ros00] J. S. Rosenthal. *A First Look at Rigorous Probability Theory*. World Scientific, 2000.
- [Seg95] R. Segala. *Modeling and Verification of Randomized Distributed Real-Time Systems*. PhD thesis, Laboratory for Computer Science, Massachusetts Institute of Technology, 1995.
- [Seg97] R. Segala. Compositional verification of randomized distributed algorithms. In *International Symposium on Compositionality: The Significant Difference (COMPOS)*, volume 1536 of *Lecture Notes in Computer Science*, pages 515–540. Springer, 1997.

- [SL95] R. Segala and N. Lynch. Probabilistic simulations for probabilistic processes. *Nordic Journal of Computing*, 2(2):250–273, 1995.
- [SM00] W. H. Sanders and J. F. Meyer. Stochastic activity networks: Formal definitions and concepts. In *European Educational Forum: School on Formal Methods and Performance Analysis*, volume 2090 of *Lecture Notes in Computer Science*, pages 315–343. Springer, 2000.
- [SV99] M. Stoelinga and F. W. Vaandrager. Root contention in IEEE 1394. In *Proceedings of the 5th International AMAST Workshop on Formal Methods for Real-Time and Probabilistic Systems*, volume 1601 of *Lecture Notes in Computer Science*, pages 53–74. Springer, 1999.
- [TF03] E. Teruel and G. Franceschinis. Well-defined generalized stochastic Petri nets: A net-level method to specify priorities. *IEEE Transactions on Software Engineering*, 29(11):962–973, Nov 2003.
- [TFP99] E. Teruel, G. Franceschinis, and M. De Pierro. Clarifying the priority specification of GSPN: Detached priorities. In *Proceedings of the 8th International Workshop on Petri Nets and Performance Models (PNPM)*, pages 114–123. IEEE Computer Society, 1999.
- [Var85] M. Y. Vardi. Automatic verification of probabilistic concurrent finite-state programs. In *26th Annual Symposium on Foundations of Computer Science (FOCS)*, pages 327–338. IEEE Computer Society, 1985.
- [vRS92] A. van Rooij and J. Smit. Dictaat bij het college maat & integraal. Lecture notes, Radboud Universiteit Nijmegen, 1992.
- [WJ06] N. Wolovick and S. Johr. A characterization of meaningful schedulers for continuous-time Markov decision processes. In *Proceedings of the 4th International Conference on Formal Modeling and Analysis of Timed Systems (FORMATS)*, volume 4202 of *Lecture Notes in Computer Science*, pages 352–367. Springer, 2006.
- [Zap08] I. S. Zapreev. *Model Checking Markov Chains: Techniques and Tools*. PhD thesis, University of Twente, Enschede, 2008.
- [ZN10] L. Zhang and M. R. Neuhäuser. Model checking interactive Markov chains, 2010. Accepted at the 16th International Conference on Tools and Algorithms for the Construction and Analysis of Systems (TACAS).

Curriculum Vitae

Martin R. Neuhäuser was born on September 1, 1979 in Kulmbach, Germany. After his Abitur in 1999 he began to study computer science at RWTH Aachen University. In 2005 he received his Diploma. Since then, he holds a PhD position at the *Formal Methods and Tools* group at the University of Twente (The Netherlands) and works as a research assistant at Prof. Joost-Pieter Katoen's *Software Modeling and Verification* group at RWTH Aachen University.

Titles in the IPA Dissertation Series since 2005

- E. Ábrahám.** *An Assertional Proof System for Multithreaded Java -Theory and Tool Support-*. Faculty of Mathematics and Natural Sciences, UL. 2005-01
- R. Ruimerman.** *Modeling and Remodeling in Bone Tissue*. Faculty of Biomedical Engineering, TU/e. 2005-02
- C.N. Chong.** *Experiments in Rights Control - Expression and Enforcement*. Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2005-03
- H. Gao.** *Design and Verification of Lock-free Parallel Algorithms*. Faculty of Mathematics and Computing Sciences, RUG. 2005-04
- H.M.A. van Beek.** *Specification and Analysis of Internet Applications*. Faculty of Mathematics and Computer Science, TU/e. 2005-05
- M.T. Ionita.** *Scenario-Based System Architecting - A Systematic Approach to Developing Future-Proof System Architectures*. Faculty of Mathematics and Computing Sciences, TU/e. 2005-06
- G. Lenzini.** *Integration of Analysis Techniques in Security and Fault-Tolerance*. Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2005-07
- I. Kurtev.** *Adaptability of Model Transformations*. Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2005-08
- T. Wolle.** *Computational Aspects of Treewidth - Lower Bounds and Network Reliability*. Faculty of Science, UU. 2005-09
- O. Tveretina.** *Decision Procedures for Equality Logic with Uninterpreted Functions*. Faculty of Mathematics and Computer Science, TU/e. 2005-10
- A.M.L. Liekens.** *Evolution of Finite Populations in Dynamic Environments*. Faculty of Biomedical Engineering, TU/e. 2005-11
- J. Eggermont.** *Data Mining using Genetic Programming: Classification and Symbolic Regression*. Faculty of Mathematics and Natural Sciences, UL. 2005-12
- B.J. Heeren.** *Top Quality Type Error Messages*. Faculty of Science, UU. 2005-13
- G.F. Frehse.** *Compositional Verification of Hybrid Systems using Simulation Relations*. Faculty of Science, Mathematics and Computer Science, RU. 2005-14
- M.R. Mousavi.** *Structuring Structural Operational Semantics*. Faculty of Mathematics and Computer Science, TU/e. 2005-15
- A. Sokolova.** *Coalgebraic Analysis of Probabilistic Systems*. Faculty of Mathematics and Computer Science, TU/e. 2005-16
- T. Gelsema.** *Effective Models for the Structure of π -Calculus Processes with Replication*. Faculty of Mathematics and Natural Sciences, UL. 2005-17
- P. Zoetewij.** *Composing Constraint Solvers*. Faculty of Natural Sciences, Mathematics, and Computer Science, UvA. 2005-18
- J.J. Vinju.** *Analysis and Transformation of Source Code by Parsing and Rewriting*. Faculty of Natural Sciences, Mathematics, and Computer Science, UvA. 2005-19
- M. Valero Espada.** *Modal Abstraction and Replication of Processes with Data*. Faculty

of Sciences, Division of Mathematics and Computer Science, VUA. 2005-20

A. Dijkstra. *Stepping through Haskell.* Faculty of Science, UU. 2005-21

Y.W. Law. *Key management and link-layer security of wireless sensor networks: energy-efficient attack and defense.* Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2005-22

E. Dolstra. *The Purely Functional Software Deployment Model.* Faculty of Science, UU. 2006-01

R.J. Corin. *Analysis Models for Security Protocols.* Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2006-02

P.R.A. Verbaan. *The Computational Complexity of Evolving Systems.* Faculty of Science, UU. 2006-03

K.L. Man and R.R.H. Schiffelers. *Formal Specification and Analysis of Hybrid Systems.* Faculty of Mathematics and Computer Science and Faculty of Mechanical Engineering, TU/e. 2006-04

M. Kyas. *Verifying OCL Specifications of UML Models: Tool Support and Compositionality.* Faculty of Mathematics and Natural Sciences, UL. 2006-05

M. Hendriks. *Model Checking Timed Automata - Techniques and Applications.* Faculty of Science, Mathematics and Computer Science, RU. 2006-06

J. Ketema. *Böhm-Like Trees for Rewriting.* Faculty of Sciences, VUA. 2006-07

C.-B. Breunesse. *On JML: topics in tool-assisted verification of JML programs.* Faculty of Science, Mathematics and Computer Science, RU. 2006-08

B. Markvoort. *Towards Hybrid Molecular Simulations.* Faculty of Biomedical Engineering, TU/e. 2006-09

S.G.R. Nijssen. *Mining Structured Data.* Faculty of Mathematics and Natural Sciences, UL. 2006-10

G. Russello. *Separation and Adaptation of Concerns in a Shared Data Space.* Faculty of Mathematics and Computer Science, TU/e. 2006-11

L. Cheung. *Reconciling Nondeterministic and Probabilistic Choices.* Faculty of Science, Mathematics and Computer Science, RU. 2006-12

B. Badban. *Verification techniques for Extensions of Equality Logic.* Faculty of Sciences, Division of Mathematics and Computer Science, VUA. 2006-13

A.J. Mooij. *Constructive formal methods and protocol standardization.* Faculty of Mathematics and Computer Science, TU/e. 2006-14

T. Krilavicius. *Hybrid Techniques for Hybrid Systems.* Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2006-15

M.E. Warnier. *Language Based Security for Java and JML.* Faculty of Science, Mathematics and Computer Science, RU. 2006-16

V. Sundramoorthy. *At Home In Service Discovery.* Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2006-17

B. Gebremichael. *Expressivity of Timed Automata Models.* Faculty of Science, Mathematics and Computer Science, RU. 2006-18

L.C.M. van Gool. *Formalising Interface Specifications.* Faculty of Mathematics and Computer Science, TU/e. 2006-19

C.J.F. Cremers. *Scyther - Semantics and Verification of Security Protocols.* Faculty of Mathematics and Computer Science, TU/e. 2006-20

J.V. Guillen Scholten. *Mobile Channels for Exogenous Coordination of Distributed Systems: Semantics, Implementation and Composition.* Faculty of Mathematics and Natural Sciences, UL. 2006-21

H.A. de Jong. *Flexible Heterogeneous Software Systems.* Faculty of Natural Sciences, Mathematics, and Computer Science, UvA. 2007-01

N.K. Kavaldjiev. *A run-time reconfigurable Network-on-Chip for streaming DSP applications.* Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2007-02

M. van Veelen. *Considerations on Modeling for Early Detection of Abnormalities in Locally Autonomous Distributed Systems.* Faculty of Mathematics and Computing Sciences, RUG. 2007-03

T.D. Vu. *Semantics and Applications of Process and Program Algebra.* Faculty of Natural Sciences, Mathematics, and Computer Science, UvA. 2007-04

L. Brandán Briones. *Theories for Model-based Testing: Real-time and Coverage.* Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2007-05

I. Loeb. *Natural Deduction: Sharing by Presentation.* Faculty of Science, Mathematics and Computer Science, RU. 2007-06

M.W.A. Streppel. *Multifunctional Geometric Data Structures.* Faculty of Mathematics and Computer Science, TU/e. 2007-07

N. Trčka. *Silent Steps in Transition Systems and Markov Chains.* Faculty of Mathematics and Computer Science, TU/e. 2007-08

R. Brinkman. *Searching in encrypted data.* Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2007-09

A. van Weelden. *Putting types to good use.* Faculty of Science, Mathematics and Computer Science, RU. 2007-10

J.A.R. Noppen. *Imperfect Information in Software Development Processes.* Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2007-11

R. Boumen. *Integration and Test plans for Complex Manufacturing Systems.* Faculty of Mechanical Engineering, TU/e. 2007-12

A.J. Wijs. *What to do Next?: Analysing and Optimising System Behaviour in Time.* Faculty of Sciences, Division of Mathematics and Computer Science, VUA. 2007-13

C.F.J. Lange. *Assessing and Improving the Quality of Modeling: A Series of Empirical Studies about the UML.* Faculty of Mathematics and Computer Science, TU/e. 2007-14

T. van der Storm. *Component-based Configuration, Integration and Delivery.* Faculty of Natural Sciences, Mathematics, and Computer Science, UvA. 2007-15

B.S. Graaf. *Model-Driven Evolution of Software Architectures.* Faculty of Electrical Engineering, Mathematics, and Computer Science, TUD. 2007-16

A.H.J. Mathijssen. *Logical Calculi for Reasoning with Binding.* Faculty of Mathematics and Computer Science, TU/e. 2007-17

D. Jarnikov. *QoS framework for Video Streaming in Home Networks.* Faculty of Mathematics and Computer Science, TU/e. 2007-18

M. A. Abam. *New Data Structures and Algorithms for Mobile Data.* Faculty of Mathematics and Computer Science, TU/e. 2007-19

W. Pieters. *La Volonté Machinale: Understanding the Electronic Voting Controversy.* Faculty of Science, Mathematics and Computer Science, RU. 2008-01

A.L. de Groot. *Practical Automaton Proofs in PVS.* Faculty of Science, Mathematics and Computer Science, RU. 2008-02

M. Bruntink. *Renovation of Idiomatic Crosscutting Concerns in Embedded Systems.* Faculty of Electrical Engineering, Mathematics, and Computer Science, TUD. 2008-03

A.M. Marin. *An Integrated System to Manage Crosscutting Concerns in Source Code.* Faculty of Electrical Engineering, Mathematics, and Computer Science, TUD. 2008-04

N.C.W.M. Braspenning. *Model-based Integration and Testing of High-tech Multidisciplinary Systems.* Faculty of Mechanical Engineering, TU/e. 2008-05

M. Bravenboer. *Exercises in Free Syntax: Syntax Definition, Parsing, and Assimilation of Language Conglomerates.* Faculty of Science, UU. 2008-06

M. Torabi Dashti. *Keeping Fairness Alive: Design and Formal Verification of Optimistic Fair Exchange Protocols.* Faculty of Sciences, Division of Mathematics and Computer Science, VUA. 2008-07

I.S.M. de Jong. *Integration and Test Strategies for Complex Manufacturing Machines.* Faculty of Mechanical Engineering, TU/e. 2008-08

I. Hasuo. *Tracing Anonymity with Coalgebras.* Faculty of Science, Mathematics and Computer Science, RU. 2008-09

L.G.W.A. Cleophas. *Tree Algorithms: Two Taxonomies and a Toolkit.* Faculty of Mathematics and Computer Science, TU/e. 2008-10

I.S. Zapreev. *Model Checking Markov Chains: Techniques and Tools.* Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2008-11

M. Farshi. *A Theoretical and Experimental Study of Geometric Networks.* Faculty of Mathematics and Computer Science, TU/e. 2008-12

G. Gulesir. *Evolvable Behavior Specifications Using Context-Sensitive Wildcards.* Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2008-13

F.D. Garcia. *Formal and Computational Cryptography: Protocols, Hashes and Commitments.* Faculty of Science, Mathematics and Computer Science, RU. 2008-14

P. E. A. Dürr. *Resource-based Verification for Robust Composition of Aspects.* Faculty

of Electrical Engineering, Mathematics & Computer Science, UT. 2008-15

E.M. Bortnik. *Formal Methods in Support of SMC Design.* Faculty of Mechanical Engineering, TU/e. 2008-16

R.H. Mak. *Design and Performance Analysis of Data-Independent Stream Processing Systems.* Faculty of Mathematics and Computer Science, TU/e. 2008-17

M. van der Horst. *Scalable Block Processing Algorithms.* Faculty of Mathematics and Computer Science, TU/e. 2008-18

C.M. Gray. *Algorithms for Fat Objects: Decompositions and Applications.* Faculty of Mathematics and Computer Science, TU/e. 2008-19

J.R. Calamé. *Testing Reactive Systems with Data - Enumerative Methods and Constraint Solving.* Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2008-20

E. Mumford. *Drawing Graphs for Cartographic Applications.* Faculty of Mathematics and Computer Science, TU/e. 2008-21

E.H. de Graaf. *Mining Semi-structured Data, Theoretical and Experimental Aspects of Pattern Evaluation.* Faculty of Mathematics and Natural Sciences, UL. 2008-22

R. Brijder. *Models of Natural Computation: Gene Assembly and Membrane Systems.* Faculty of Mathematics and Natural Sciences, UL. 2008-23

A. Koprowski. *Termination of Rewriting and Its Certification.* Faculty of Mathematics and Computer Science, TU/e. 2008-24

U. Khadim. *Process Algebras for Hybrid Systems: Comparison and Development.* Faculty of Mathematics and Computer Science, TU/e. 2008-25

J. Markovski. *Real and Stochastic Time in Process Algebras for Performance Evaluation.* Faculty of Mathematics and Computer Science, TU/e. 2008-26

H. Kastenbergh. *Graph-Based Software Specification and Verification.* Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2008-27

I.R. Buhan. *Cryptographic Keys from Noisy Data Theory and Applications.* Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2008-28

R.S. Marin-Perianu. *Wireless Sensor Networks in Motion: Clustering Algorithms for Service Discovery and Provisioning.* Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2008-29

M.H.G. Verhoef. *Modeling and Validating Distributed Embedded Real-Time Control Systems.* Faculty of Science, Mathematics and Computer Science, RU. 2009-01

M. de Mol. *Reasoning about Functional Programs: Sparkle, a proof assistant for Clean.* Faculty of Science, Mathematics and Computer Science, RU. 2009-02

M. Lormans. *Managing Requirements Evolution.* Faculty of Electrical Engineering, Mathematics, and Computer Science, TUD. 2009-03

M.P.W.J. van Osch. *Automated Model-based Testing of Hybrid Systems.* Faculty of Mathematics and Computer Science, TU/e. 2009-04

H. Sozer. *Architecting Fault-Tolerant Software Systems.* Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2009-05

M.J. van Weerdenburg. *Efficient Rewriting Techniques.* Faculty of Mathematics and Computer Science, TU/e. 2009-06

H.H. Hansen. *Coalgebraic Modelling: Applications in Automata Theory and Modal Logic.* Faculty of Sciences, Division of Mathematics and Computer Science, VUA. 2009-07

A. Mesbah. *Analysis and Testing of Ajax-based Single-page Web Applications.* Faculty of Electrical Engineering, Mathematics, and Computer Science, TUD. 2009-08

A.L. Rodriguez Yakushev. *Towards Getting Generic Programming Ready for Prime Time.* Faculty of Science, UU. 2009-9

K.R. Olmos Joffré. *Strategies for Context Sensitive Program Transformation.* Faculty of Science, UU. 2009-10

J.A.G.M. van den Berg. *Reasoning about Java programs in PVS using JML.* Faculty of Science, Mathematics and Computer Science, RU. 2009-11

M.G. Khatib. *MEMS-Based Storage Devices. Integration in Energy-Constrained Mobile Systems.* Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2009-12

S.G.M. Cornelissen. *Evaluating Dynamic Analysis Techniques for Program Comprehension.* Faculty of Electrical Engineering, Mathematics, and Computer Science, TUD. 2009-13

D. Bolzoni. *Revisiting Anomaly-based Network Intrusion Detection Systems.* Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2009-14

H.L. Jonker. *Security Matters: Privacy in Voting and Fairness in Digital Exchange.* Faculty of Mathematics and Computer Science, TU/e. 2009-15

M.R. Czenko. *TuLiP - Reshaping Trust Management.* Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2009-16

T. Chen. *Clocks, Dice and Processes.* Faculty of Sciences, Division of Mathematics and Computer Science, VUA. 2009-17

C. Kaliszyk. *Correctness and Availability: Building Computer Algebra on top of Proof Assistants and making Proof Assistants available over the Web.* Faculty of Science, Mathematics and Computer Science, RU. 2009-18

R.S.S. O'Connor. *Incompleteness & Completeness: Formalizing Logic and Analysis in Type Theory.* Faculty of Science, Mathematics and Computer Science, RU. 2009-19

B. Ploeger. *Improved Verification Methods for Concurrent Systems.* Faculty of Mathematics and Computer Science, TU/e. 2009-20

T. Han. *Diagnosis, Synthesis and Analysis of Probabilistic Models.* Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2009-21

R. Li. *Mixed-Integer Evolution Strategies for Parameter Optimization and Their Applications to Medical Image Analysis.* Faculty of Mathematics and Natural Sciences, UL. 2009-22

J.H.P. Kwisthout. *The Computational Complexity of Probabilistic Networks.* Faculty of Science, UU. 2009-23

T.K. Cocx. *Algorithmic Tools for Data-Oriented Law Enforcement.* Faculty of Mathematics and Natural Sciences, UL. 2009-24

A.I. Baars. *Embedded Compilers.* Faculty of Science, UU. 2009-25

M.A.C. Dekker. *Flexible Access Control for Dynamic Collaborative Environments.* Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2009-26

J.F.J. Laros. *Metrics and Visualisation for Crime Analysis and Genomics.* Faculty of Mathematics and Natural Sciences, UL. 2009-27

C.J. Boogerd. *Focusing Automatic Code Inspections.* Faculty of Electrical Engineering, Mathematics, and Computer Science, TUD. 2010-01

M.R. Neuhäuser. *Model Checking Non-deterministic and Randomly Timed Systems.* Faculty of Electrical Engineering, Mathematics & Computer Science, UT. 2010-02