

Strategies in Robust and Stochastic Model Predictive Control



Diego Muñoz Carpintero
St Edmund Hall
University of Oxford

A thesis submitted for the degree of
Doctor of Philosophy
Trinity 2014

Abstract

The presence of uncertainty in model predictive control (MPC) has been accounted for using two types of approaches: robust MPC (RMPC) and stochastic MPC (SMPC). Ideal RMPC and SMPC formulations consider closed-loop optimal control problems whose exact solution, via dynamic programming, is intractable for most systems. Much effort then has been devoted to find good compromises between the degree of optimality and computational tractability. This thesis expands on this effort and presents robust and stochastic MPC strategies with reduced online computational requirements where the conservativeness incurred is made as small as conveniently possible.

Two RMPC strategies are proposed for linear systems under additive uncertainty. They are based on a recently proposed approach which uses a triangular prediction structure and a non-linear control policy. One strategy considers a transference of part of the computation of the control policy to an offline stage. The other strategy considers a modification of the prediction structure so that it has a striped structure and the disturbance compensation extends throughout an infinite horizon. An RMPC strategy for linear systems with additive and multiplicative uncertainty is also presented. It considers polytopic dynamics that are designed so as to maximize the volume of an invariant ellipsoid, and are used in a dual-mode prediction scheme where constraint satisfaction is ensured by an approach based on a variation of Farkas' Lemma.

Finally, two SMPC strategies for linear systems with additive uncertainty are presented, which use an affine-in-the-disturbances control policy with a striped structure. One strategy considers an offline sequential design of the gains of the control policy, while these are variables in the online optimization in the other.

Control theoretic properties, such as recursive feasibility and stability, are studied for all the proposed strategies. Numerical comparisons show that the proposed algorithms can provide a convenient compromise in terms of computational demands and control authority.

Acknowledgements

I want to thank everyone who somehow helped me write this thesis and go all the way through my DPhil.

My deepest gratitude goes to my supervisors, Basil Kouvaritakis and Mark Cannon, for giving me the opportunity to work with them. Their guidance, constant support, encouragement and advice has been invaluable. It has been an honor, an incredibly rewarding experience and a real pleasure.

I want to thank Saša Raković for showing me a new way of thinking about robust control. Special thanks also go to my fellow control group members Qifeng Cheng and James Fleming for the great disposition and many constructive discussions.

But not everything is about research, so I want to express my gratitude to all the friends I made during my stay in Oxford for helping me have truly a magnificent and unforgettable experience (and survive the writing process!). Beth, Christian, Dave, John, Lucy, Phil, Rosie, Sarah, Suzanne, Tam, Julian, Rodolfo and Tamas, thanks for being my very large Teddy Hall Family, for all the good times and constant support. To Aparajita and Shuaishuai, for all the good times in the last portion of my stay in Oxford. And last but not least, to all the members of the Teddy Hall MCR Football Team.

Felipe, thanks to your advice I came to Oxford. You will have my eternal gratitude.

I want to thank my parents, Teresa and Marcelo, and my sister, Natalia, for their love, full support and encouragement throughout this adventure.

Thank you all for this wonderful journey.

Contents

Abstract	i
Acknowledgements	ii
1 Introduction	1
1.1 Motivation and objectives	1
1.2 Contributions and outline	4
1.3 Publications	6
1.4 Notation and basic definitions	7
1.5 Acronyms	9
2 Background on Classical, Robust and Stochastic Model Predictive Control	10
2.1 Classical MPC	11
2.1.1 Formulation	11
2.1.2 Invariance, feasibility and stability in classical MPC	16
2.1.3 Different types of cost	24
2.2 Robust Model Predictive Control	25
2.2.1 Invariance and stability for Robust MPC	26
2.2.2 Open-loop versus feedback Robust MPC	31
2.2.3 Review of Robust MPC with multiplicative uncertainty	32
2.2.3.1 Optimized state linear feedback	32
2.2.3.2 Augmented dynamics for efficient RMPC	33
2.2.3.3 Optimized dynamics RMPC	37
2.2.4 Review of Robust MPC with additive uncertainty	39
2.2.4.1 Feedback RMPC strategies	39
2.2.4.2 Tube MPC strategies	43
2.3 Stochastic Model Predictive Control	52
2.3.1 Main aspects of Stochastic MPC	53

2.3.2	Review of Stochastic MPC	57
2.3.2.1	Probabilistic constraints via ellipsoidal bounding . .	57
2.3.2.2	Tight constraints based on explicit use of probabilistic distributions	60
3	Offline tube design for efficient implementation of Parameterized Tube Model Predictive Control	66
3.1	Introduction	67
3.2	System description and PTMPC background	69
3.3	PTMPC: An offline Simplification	77
3.3.1	Offline design of initial partial tube cross sections	84
3.3.2	Simplified PTMPC: Online implementation and relevant prop- erties	89
3.4	Simulation Results	100
3.5	Conclusions	104
4	Striped Parameterized Tube Model Predictive Control	106
4.1	Introduction	107
4.2	System description and separable prediction scheme	109
4.3	Striped PTMPC: Predictions and constraints	114
4.3.1	Prediction strategy	114
4.3.2	Constraints	117
4.4	Striped PTMPC: Online implementation and properties - Two cases .	123
4.4.1	Cost based on state distance to an invariant set - strong stability	124
4.4.2	Nominal cost and input-to-state stability	135
4.4.3	On the selection of $\mathbb{X}_0$	140
4.5	Simulation results	141
4.6	Conclusions	145
5	Dual-Mode Robust MPC with optimized polytopic dynamics	147
5.1	Introduction	148
5.2	System description and background on optimized polytopic dynamics	150
5.3	Necessity and sufficiency conditions for the optimized polytopic dynamics	155
5.4	The online robust MPC strategy	160
5.4.1	The control policy and constraints	160
5.4.1.1	On the design of H , H_i and V	167
5.4.2	Cost function	170

5.4.3	Overall MPC algorithm and properties	173
5.5	Simulation results	178
5.6	Conclusions	180
6	Striped affine-in-the-disturbances compensation in Stochastic MPC	184
6.1	Introduction	185
6.2	System description and background	187
6.3	Closed-loop prediction strategies to compensate for the effects of disturbances	193
6.3.1	Connection with the control policy of Chapter 4	199
6.4	SMPC with offline design of the disturbance compensation gains	200
6.4.1	Offline design of the disturbance compensation gains	200
6.4.2	Constraints	204
6.4.3	Cost function	206
6.4.4	SMPC implementation and properties	209
6.5	SMPC with online optimization of the disturbance compensation gains	212
6.5.1	Constraints and cost function	212
6.5.2	SMPC online implementation and properties	215
6.6	Simulation results	219
6.7	Concluding remarks	223
7	Concluding remarks and further research	225
7.1	Concluding remarks: a summary	225
7.2	Further research	228
	Bibliography	231

Chapter 1

Introduction

1.1 Motivation and objectives

Model predictive control (MPC) is a family of control strategies that consist of finding the solution, at each sample time, of an optimal control problem where the current state of the plant is the initial state. This yields an optimal control input sequence, but only the first control input of the sequence is actually applied to the plant. Given that the optimization is performed at each sample time and the horizon of the optimal control problem is receding one step at each sample time, MPC is also known as receding horizon control (RHC). MPC has enjoyed wide success because of its ability to handle constraints in the optimal control formulation, which allows processes to operate close to the constraint boundaries, thereby providing the potential for more effective control.

The idea of using linear programming for the control of linear systems was first proposed in 1963 by [71]. Then, MPC was first advocated in the following decade by [79], [80], and then a plethora of works followed, passing through the success in the process industry in the 1980s until the present (for a review see [62], [46], [17], [78]). The field can be considered mature in respect of developments of stability results [62].

Initial developments considered deterministic prediction models, where no uncer-

tainty was considered. However, for most applications, the mathematical models can only provide an approximate description of the system dynamics, and/or the effect of the uncertainties is too significant to be neglected. Hence the use of MPC can result in poor performances or even infeasibility of the online optimization. For this reason, robust MPC (RMPC) emerged as a family of MPC strategies that preserve stability and performance for all possible realizations of bounded uncertainties, which have been usually modelled as model uncertainties (of the system parameters) and/or additive disturbances.

Early formulations of RMPC considered open-loop optimal control problems [95], [64], just like in the case of deterministic models, where only a single sequence of control inputs is determined for all possible realizations of the uncertainties or disturbances. This, however, is conservative [9] because it does not acknowledge the fact that the future information of the states (and therefore of the uncertainties) will be known and ready to use by the controller. Furthermore, not including these in the formulation could even lead to infeasibility [86]. This information can be effectively used in closed-loop optimal control problems; the optimization now is performed over a control policy, which is a sequence of control laws (which are functions of the state), instead of a single control sequence, and therefore each predicted control input will depend on the uncertainty. The exact solution of such closed-loop optimal control problems can be obtained via dynamic programming [9] or feedback optimization [86]. However, the size of the problem (number of constraints and variables) grows exponentially with the prediction horizon of the optimal control problem, and is therefore intractable for most problems. For this reason most developments have consisted of simplifications of the control policies (e.g. see [45], [51], [61], [63], [58], [38], [73]) so as to decrease the size of the problems so that they can be applied online. However, some of these strategies can sometimes prove to be too conservative [45], [51], [61], [63]. Others are less conservative [58], [38], [73] but the number of variables and

constraints grows quadratically with the prediction horizon, which may still require significant computational effort and restrict application to systems with low-order models and short predictions horizons. Based on this, one objective of this thesis is to develop RMPC strategies with reduced computational complexity (namely, linear growth of the number of constraints and variables with the prediction horizon) while having as much control authority as conveniently possible.

RMPC only takes into account the bounds of the uncertainties, but not their probabilistic distribution. This is a problem when probabilistic constraints are present. Additionally, the min-max optimization often adopted in RMPC, which consists of minimizing the performance of the worst cases of the uncertainties, can be extremely conservative because the uncertainties near the bounds can be extremely unlikely. For this reason, stochastic MPC (SMPC) has emerged as a family of techniques that incorporate the stochastic descriptions of the uncertainties for handling the probabilistic constraints and using more realistic probabilistic measures of the performance. Early works in this area [57], [88], [89] considered these aspects, but the exploration of issues such as the stability and feasibility of the closed-loop implementation was explored later, e.g. [23], [49], [21]. These strategies, however, only consider quasi closed-loop optimal control formulations (although there is feedback in the predictions in the form of a static feedback term Kx in the predicted inputs, no special treatment is considered for the unknown future uncertainties), which can be conservative. Therefore the second objective of this thesis is to design SMPC strategies that use control policies with the aim of having increased control authority, but that at the same time impose reduced computational requirements; to be understood again as having a linear growth of the number of constraints and variables with the prediction horizon.

1.2 Contributions and outline

Based on the discussion above, this thesis presents robust and stochastic MPC strategies that require reduced online computational effort, i.e. such that they have a linear growth of the number of constraints and variables with the prediction horizon. This will be achieved by transferring part of the computations to an offline stage and/or modifying the structure of the control policies. These strategies will be developed with the intention that the reduction in the computational burden is such that the degree of conservativeness incurred is made as small as conveniently possible. Furthermore, a reduction in computational burden implies that with an adequate selection of parameters it might be possible that the proposed methods have larger domains of attractions whilst being faster than the established methods in the literature.

Chapter 2 provides a background of the most relevant concepts and earlier works that are needed for the strategies presented in this thesis. The background is divided into classical MPC, RMPC and SMPC.

Chapter 3 presents an RMPC strategy for linear systems under additive uncertainty based on the strategy of [73], called parameterized tube MPC (PTMPC), such that part of the computation of the predicted control policy is transferred to the offline stage, hence this strategy will be called offline PTMPC (OPTMPC). This yields a time varying online optimization where the number of variables and constraints grows linearly with the prediction horizon as opposed to quadratically as in [73], and therefore is comparatively much faster than PTMPC as N increases. Thus in a considerable number of cases, one might use larger values of N in order to obtain larger domains of attraction than with PTMPC while still having to solve, online, problems with less computational burden.

Chapter 4 presents an RMPC strategy that modifies the prediction structure of [73] and will be called striped PTMPC (SPTMPC). This modification results in a

striped control policy such that the disturbance compensation extends throughout the infinite prediction horizon, rather than being limited to a finite horizon. On account of this, SPTMPC has the potential to outperform PTMPC. Additionally, in SPTMPC the online optimization has a number of variables and constraints that grows linearly with the prediction horizon, and as with the strategy of Chapter 3, this means that one can use larger values of N in order to obtain larger domains of attraction than with PTMPC while still requiring less computational burden.

Chapter 5 presents an RMPC strategy for linear systems under multiplicative and additive uncertainty. In common with [34], it extends the strategy of [19], which is for systems subject only to multiplicative uncertainty, and considers a formulation based on optimized polytopic dynamics. Necessary and sufficient conditions for the invariance of polytopic dynamics are provided, whereas only sufficient conditions are provided in [34]. Additionally, the optimized polytopic dynamics are effectively used in a dual-mode prediction based MPC strategy, whereas the strategy of [34] only presents a dual-mode prediction based MPC strategy under the assumption that the multiplicative uncertainty is known in the near future.

Chapter 6 presents two SMPC strategies based on the strategy of [49], that effectively use a striped affine-in-the-disturbances control policy to obtain improved control authority. As in SPTMPC, the compensation extends throughout an infinite horizon and yields an online optimization problem with a number of variables and constraints that increases linearly with the prediction horizon N . As in [49], the satisfaction of the probabilistic constraints is ensured with conditions on the nominal state and input predictions through the introduction of slack variables that account for the effect of the uncertainty. In one of the strategies, the disturbance feedback gains are computed offline, whereas in the other these are optimized online.

Finally, Chapter 7 provides a summary of the results presented and of the concluding remarks presented at the end of chapters 3, 4 5 and 6, and suggests possible

future research directions.

1.3 Publications

The work presented in this thesis is largely based on published or submitted material.

Chapter 3 is based on

S.V. Rakovic, D. Muñoz-Carpintero, M. Cannon and B. Kouvaritakis. Offline tube design for efficient implementation of parameterized tube model predictive control. In *Decision and Control (CDC), 51st IEEE Conference on, Maui, Hawaii*, pages 5176-5181, Dec. 2012. [76]

Chapter 4 contains material to appear in

D. Muñoz-Carpintero, B. Kouvaritakis and M. Cannon. Striped Parameterized Tube Model Predictive Control. In *19th IFAC World Congress, Cape Town, SA*, Aug. 2014. [67]

Chapter 5 is based on

D. Muñoz-Carpintero, M. Cannon and B. Kouvaritakis. Recursively feasible robust MPC for linear systems with additive and multiplicative uncertainty using optimized polytopic dynamics. In *Decision and Control (CDC), 52nd IEEE Conference on, Florence, Italy*, pages 1101-1106, Dec. 2013. [65]

Chapter 6 is based on

B. Kouvaritakis, M. Cannon and D. Muñoz-Carpintero. Efficient prediction strategies for disturbance compensation in stochastic MPC. *Int. Journal of Systems Science*, 44(7):1344-1353, 2013. [48]

D. Muñoz-Carpintero, B. Kouvaritakis and M. Cannon. On prediction strategies in stochastic MPC. In *Control and Automation (ICCA), 9th IEEE International Conference on, Santiago, Chile*, pages 249-254, Dec. 2011. [66]

1.4 Notation and basic definitions

Frequently used notation and basic definitions are given next.

- x_t , u_t and w_t denote the state, control input and disturbance acting at time t .
- $(\cdot)_{t+k|t}$ denotes the prediction of the variable at time $t+k$ made at time t , so that $x_{t+k|t}$ is the predicted state at time $t+k$ made at time t .
- 0 denotes a real matrix of appropriate dimensions with all elements zero.
- $\underline{1}$ denotes the vector $[1, 1, \dots, 1]^T$ of appropriate dimension.
- $A \succ 0$ denotes that the matrix A is positive definite, and $A \succeq 0$ denotes that A is semi-positive definite.
- $A > 0$ denotes that all the elements of A are greater than zero, and $A \geq 0$ denotes that all the elements of A are greater or equal to zero.
- For $a, b \in \mathbb{R}^n$ the inequality $a \geq b$ denotes that $a_i \geq b_i$, for $i = 1, \dots, n$, where a_i denotes the i th component of a (i.e. the inequality $a \geq b$ applies elementwise).
- $\text{tr}(A)$ is the trace of A .
- $\mathbb{N}$ is the set of natural numbers (non-negative integers) and $\mathbb{N}_+$ is the set of positive natural numbers. $\mathbb{R}$ is the set of real numbers and $\mathbb{R}_+$ is the set of non-negative real numbers.
- Given two sets $X, Y \subset \mathbb{R}^n$, the Minkowski set addition is defined by $X \oplus Y := \{x + y : x \in X, y \in Y\}$. Given a vector $x \in \mathbb{R}^n$, $x \oplus X$ denotes $\{x\} \oplus X$. Given a sequence of sets $\{X_i\}_{i \in \mathbb{N}}$, $\bigoplus_{i=a}^b X_i$ denotes $X_a \oplus \dots \oplus X_b$.
- $\text{conv}(X)$ denotes the convex hull of X .

- Given a convex compact set $L \subset \mathbb{R}^n$ that contains the origin in its interior, the function $\mathcal{G}(L, x) = \min_{\mu} \{\mu : x \in \mu L, \mu \in \mathbb{R}_+\}$ for $x \in \mathbb{R}^n$ is the gauge function of L .
- Given a symmetric convex compact set $L \subset \mathbb{R}^n$ that contains the origin in its interior, the gauge function of L induces the vector norm $|x|_L = \mathcal{G}(L, x)$.
- Given a symmetric convex compact set $L \subset \mathbb{R}^n$ that contains the origin in its interior, and a non-empty closed set $X \subset \mathbb{R}^n$, $\text{dist}(L, y, X) := \inf_x \{|x - y|_L : x \in X\}$ for $y \in \mathbb{R}^n$ is called the distance function associated with the set X .
- A polyhedron is the intersection of a finite number of open and/or closed half-spaces and a polytope is a closed and bounded polyhedron.
- For $x \in \mathbb{X}$ and $y \in \mathbb{Y}$, (x, y) denotes an element in the product space $\mathbb{X} \times \mathbb{Y}$.
- Given a convex set $X \subset \mathbb{R}^n$ and two matrices $A, B \in \mathbb{R}^{m \times n}$, $(A, B)X$ denotes the set $\text{conv}(\{(Ax, Bx) : x \in X\})$. Given two sets $X = \text{conv}(\{x_1, \dots, x_n\})$, $Y = \text{conv}(\{y_1, \dots, y_n\})$, then (X, Y) denotes the set $\text{conv}(\{(x_1, y_1), \dots, (x_n, y_n)\})$.
- A continuous function $\phi : \mathbb{R}_+ \rightarrow \mathbb{R}_+$ is said to be a $\mathcal{K}$ -function if it is continuous, strictly increasing and $\phi(0) = 0$, and it is a $\mathcal{K}_\infty$ -function if it is a $\mathcal{K}$ -function and $\phi(s) \rightarrow \infty$ as $s \rightarrow \infty$.
- A continuous function $\phi : \mathbb{R}_+ \times \mathbb{R}_+ \rightarrow \mathbb{R}_+$ is a $\mathcal{KL}$ -function if for all $t \in \mathbb{R}_+$, $\phi(\cdot, t)$ is a $\mathcal{K}$ -function and for all $s \in \mathbb{R}_+$, $\phi(s, \cdot)$ is decreasing with $\phi(s, t) \rightarrow 0$ as $t \rightarrow \infty$.
- $\Pr\{E\}$ denotes the probability of the event E .
- Given a stochastic process X_k , $k = 0, 1, \dots$, $\mathbb{E}\{X_k\}$ and $\mathbb{E}_t\{X_k\}$ denote the expected value of X_k and the expected value of X_k conditional on the information available at time t , respectively.

As far as possible use of the above notation will be made throughout the thesis. However this will not always be possible, mainly in the case of the symbols, due to the different mathematical tools used in the different chapters.

1.5 Acronyms

BMI	Bilinear matrix inequality
CFD	Current future disturbances
DP	Dynamic programming
KFD	Known future disturbances
LMI	Linear matrix inequality
LP	Linear program
MPC	Model predictive control
PTMPC	Parameterized tube model predictive control
PWA	Piecewise Affine
PWQ	Piecewise quadratic
QP	Quadratic program
RMPC	Robust model predictive control
SMPC	Stochastic model predictive control
SDP	Semidefinite program

Chapter 2

Background on Classical, Robust and Stochastic Model Predictive Control

This chapter contains a review of the relevant concepts and developments for the study of robust and stochastic MPC. Firstly, the fundamental concepts and ideas behind classical MPC (that uses linear deterministic systems), which are the basis for robust and stochastic MPC, are presented. This is followed with a review of robust MPC (RMPC). The differences with classical MPC in terms of feasibility, invariance and stability, and the differences between open-loop and feedback prediction strategies are highlighted. This review also describes the main concepts and the most important strategies for handling different types of uncertainty: additive uncertainty (or disturbances) and multiplicative uncertainty (also known as polytopic uncertainty or difference inclusions). The section concludes by presenting the main concepts and a review of the most relevant contributions in stochastic MPC (SMPC) for systems with additive disturbances.

The presentation of the section on classical MPC is the most detailed among the

three sections in this chapter. This is because it introduces concepts and ideas that are used again (or extended) for the robust and stochastic cases. This higher level of detail is reflected mainly in the proofs; these are presented only in the classical MPC section. This is convenient for two reasons: (i) the ideas and concepts used in the proofs for the deterministic case are simpler than in the robust and stochastic cases (yet it is noted that the concepts and ideas are very similar for the classical, robust and stochastic cases) ; and (ii) similar proofs for the robust and stochastic cases are performed in the latter chapters.

2.1 Classical MPC

2.1.1 Formulation

Classical MPC refers to MPC strategies that consider linear time-invariant systems that are not subject to any type of uncertainties.

A plethora of different formulations have been proposed since the introduction of MPC in the 1960s, going through the success in the process industry in the 1980s, until the present, and it is now considered a mature technique. Thus focus will be given next to what is currently considered the most typical formulation for linear systems [62], [78], [25], [94].

Consider a linear, time-invariant, discrete-time system

$$x_{t+1} = Ax_t + Bu_t, \tag{2.1}$$

where $A \in \mathbb{R}^{n \times n}$, $B \in \mathbb{R}^{n \times n_u}$, and x_t and u_t are the state and the control input. The system is subject to state and input constraints.

$$x_t \in \mathbb{X}, \quad u_t \in \mathbb{U}, \tag{2.2}$$

where $\mathbb{X}$ is a convex, closed subset of $\mathbb{R}^n$ and $\mathbb{U}$ is a convex, compact subset of $\mathbb{R}^{n_u}$. $\mathbb{X}$ and $\mathbb{U}$ contain the origin in their interiors and are described by linear inequalities

$$\mathbb{X} = \{x : Fx \leq \underline{1}\}, \quad \mathbb{U} = \{u : Gu \leq \underline{1}\}, \quad (2.3)$$

where $F \in \mathbb{R}^{n_{cx} \times n}$, $G \in \mathbb{R}^{n_{cu} \times n_u}$.

Assumption 2.1. (i) The pair (A, B) is stabilizable. (ii) The state is directly measurable. (iii) x_t is known at the moment of choosing u_t , and therefore control laws that are functions of the current state are considered.

The regulation problem is considered, where the control objective is to steer the state to the origin, or to an equilibrium state $\bar{x}$, for which there exists a control input $\bar{u}$ that satisfies $\bar{x} = A\bar{x} + B\bar{u}$. A suitable change of coordinates reduces the second case to the first, so without loss of generality only steering the state to the origin is considered in this thesis.

The main idea in MPC is to optimize the predicted performance of the system, such that the system constraints are satisfied, over a given prediction horizon N by finding an optimal sequence of control inputs, and then only the first of these inputs is actually applied to the system. This procedure is repeated at each sample time, and the prediction horizon slides along with time t such that the predictions go until time $t + N$ in what is called a *receding horizon* strategy, which lends MPC the alternative name of *receding horizon control*.

Early formulations considered finite horizon formulations that could not guarantee stability, or that in order to do so require the use of very restrictive constraints (such as setting the terminal state to be zero), whereas schemes based on infinite horizons can guarantee stability [59]. Dual prediction mode MPC [62] provides practical means to handle infinite horizons, and is described next.

Dual prediction mode MPC is a type of infinite horizon formulations of MPC

where the infinite prediction horizon is divided into (i) the first N prediction instants, known as the near horizon or Mode 1; (ii) and the remaining prediction instants, known as the far horizon or Mode 2. The N inputs in Mode 1 $\{u_t, \dots, u_{t+N-1}\}$ are free (optimization variables), while the remaining inputs of Mode 2 are given by a fixed terminal control law. This terminal control law is usually (and during the remainder of this chapter, unless mentioned otherwise) given by a stabilizing linear state feedback law $u = Kx$. The idea is that the optimization of a predicted cost is performed by using explicitly only the variables of Mode 1. The length of Mode 1, N , is regarded as the prediction horizon.

The state dynamics and system constraints to be satisfied by the predictions are

$$\begin{aligned} x_{t|t} &= x_t, \\ x_{t+k+1|t} &= Ax_{t+k|t} + Bu_{t+k|t}, \quad k = 0, \dots, N-1 \end{aligned} \tag{2.4}$$

and

$$\begin{aligned} x_{t+k|t} &\in \mathbb{X}, \quad k = 0, \dots, N-1, \\ u_{t+k|t} &\in \mathbb{U}, \quad k = 0, \dots, N-1, \end{aligned} \tag{2.5}$$

where x_t is the current state at time t and the sub-index $(\cdot)_{t+k|t}$ denotes the prediction of the variable at time $t+k$ made at time t .

The predicted cost of the system at time t depends on the predicted sequences of states $\mathbf{x}_t = \{x_{t|t}, x_{t+1|t}, \dots, x_{t+N|t}\}$ and control inputs $\mathbf{u}_t = \{u_{t|t}, u_{t+1|t}, \dots, u_{t+N-1|t}\}$, such that only the variables of Mode 1 and a terminal cost are considered:

$$J(\mathbf{x}_t, \mathbf{u}_t) = \sum_{k=0}^{N-1} \ell(x_{t+k|t}, u_{t+k|t}) + F(x_{t+N|t}), \tag{2.6}$$

where $\ell(x, u) = x^T Q x + u^T R u$ is the stage cost, $F(x) = x^T P x$ is the terminal cost that accounts for the performance in Mode 2, $Q, P \in \mathbb{R}^{n \times n}$, $R \in \mathbb{R}^{n_u \times n_u}$, $Q, P \succeq 0$, $R \succ 0$, and the pair $(A, Q^{\frac{1}{2}})$ is detectable.

The following terminal constraint is also used:

$$x_{t+N|t} \in \mathbb{X}_f = \{x : Hx \leq \underline{1}\}, \quad (2.7)$$

where $H \in \mathbb{R}^{n_f \times n}$ and $\mathbb{X}_f \subseteq \mathbb{X}$ is such that (2.7) guarantees that the constraints are satisfied throughout all of Mode 2.

More details about the terminal cost and the terminal constraints are provided in Section 2.1.2 when the stability, feasibility and invariance of the MPC strategy is analysed.

At every sample time t , the MPC strategy requires the solution of the optimal control problem $\mathbb{P}(x_t)$:

$$\begin{aligned} J^*(x_t) &= \min_{(\mathbf{x}_t, \mathbf{u}_t)} \{J(\mathbf{x}_t, \mathbf{u}_t) \mid (\mathbf{x}_t, \mathbf{u}_t) \text{ satisfy (2.4), (2.5), (2.7)}\}, \\ (\mathbf{x}_t^*(x_t), \mathbf{u}_t^*(x_t)) &= \arg \min_{(\mathbf{x}_t, \mathbf{u}_t)} \{J(\mathbf{x}_t, \mathbf{u}_t) \mid (\mathbf{x}_t, \mathbf{u}_t) \text{ satisfy (2.4), (2.5), (2.7)}\}. \end{aligned} \quad (2.8)$$

Then, only the first control input $u_t = u_{t|t}^*(x_t)$ in the optimal sequence is applied to the system, which defines the implicit MPC law

$$\kappa(x_t) = u_{t|t}^*(x_t), \quad (2.9)$$

which is time-invariant since the stage and terminal costs and the constraints are time-invariant.

The optimization in $\mathbb{P}(x_t)$ has linear constraints and a quadratic cost, therefore it

is quadratic program (QP), as shown next.

$$\begin{aligned}
 & \underset{\mathbf{x}_t, \mathbf{u}_t}{\text{minimize}} && \sum_{k=0}^{N-1} (x_{t+k|t}^T Q x_{t+k|t} + u_{t+k|t}^T R u_{t+k|t}) + x_{t+N|t}^T P x_{t+N|t} \\
 & \text{subject to} && x_{t|t} = x_t, \\
 & && x_{t+k+1|t} = A x_{t+k|t} + B u_{t+k|t}, \quad k = 0, \dots, N-1, \\
 & && F x_{t+k|t} \leq \underline{1}, \quad k = 0, \dots, N-1, \\
 & && G u_{t+k|t} \leq \underline{1}, \quad k = 0, \dots, N-1, \\
 & && H x_{t+N|t} \leq \underline{1}.
 \end{aligned} \tag{2.10}$$

The optimization problem presented above has a number of variables and constraints that grow linearly with the prediction horizon N , the state dimension n and the input dimension n_u (when each varies individually), and therefore is tractable. The problem can be formulated to have $x_{t+k|t}$ as variables, as shown above, and in this case the dynamic equations are included as constraints. Alternatively, the predicted $x_{t+k|t}$ can be written as functions of the inputs $u_{t+k|t}$ and x_t

$$x_{t+k|t} = A^k x_t + \begin{bmatrix} A^{k-1} B & A^{k-2} B & \dots & B \end{bmatrix} \begin{bmatrix} u_{t|t} \\ u_{t+1|t} \\ \vdots \\ u_{t+N-1|t} \end{bmatrix}, \tag{2.11}$$

therefore reducing the number of variables and constraints since $x_{t+k|t}$ and the dynamic equations are not considered explicitly, and the constraints $F x_{t+k|t} \leq \underline{1}$, $k =$

$0, \dots, N-1$, are replaced by

$$F \left(A^k x_t + \begin{bmatrix} A^{k-1}B & A^{k-2}B & \dots & B \end{bmatrix} \begin{bmatrix} u_{t|k} \\ u_{t+1|t} \\ \vdots \\ u_{t+N-1|t} \end{bmatrix} \right) \leq \underline{1}, \quad k = 0, \dots, N-1. \quad (2.12)$$

The former structure however, offers a numerically more robust formulation by usually being better conditioned [77].

Under the control law of (2.9), the dynamics of system (2.1) become

$$x_{t+1} = Ax_t + B\kappa(x_t), \quad (2.13)$$

which are non-linear since $\kappa(x_t)$ is in general non-linear. More precisely, $\kappa(x_t)$ is a piecewise affine (PWA) function of x_t [5].

2.1.2 Invariance, feasibility and stability in classical MPC

An analysis of the feasibility, stability and convergence properties of the classical MPC controller is performed here.

The concept of positively invariant sets is fundamental for the analysis of feasibility of dual-mode prediction based MPC controllers.

Definition 2.1 (Positively invariant set). *A set $\Omega \subseteq \mathbb{X}$ is said to be positively invariant for system (2.1) and constraints (2.2) under the control law $u = \kappa(x)$ if and only if, for all $x \in \Omega$, $Ax + B\kappa(x) \in \Omega$ and $\kappa(x) \in \mathbb{U}$.*

This means that, for a positively invariant set Ω , for system (2.1) and constraints (2.2) under the control law $u = \kappa(x)$: (i) if the state is in Ω then it will remain in this set at the subsequent time instant, which recursively guarantees that the state will remain in Ω for all future instants; and (ii) if $x \in \Omega$ then $x \in \mathbb{X}$ and $\kappa(x) \in \mathbb{U}$.

Combining (i) and (ii) it follows that if $x \in \Omega$ then the constraints $x \in \mathbb{X}$ and $\kappa(x) \in \mathbb{U}$ will be satisfied for all future instants.

Definition 2.2 (Maximal positively invariant set). *A set $\bar{\Omega} \subseteq \mathbb{X}$ is said to be the maximal positively invariant for system (2.1) and constraints (2.2) under the control law $u = \kappa(x)$ if and only if it is positively invariant and $\Omega \subseteq \bar{\Omega}$ for all positively invariant sets Ω .*

The following assumption is made about the terminal constraint and control law.

Assumption 2.2. *$\mathbb{X}_f$ is a positively invariant set for system (2.1) and constraints (2.2) under the stabilizing control law $u = Kx$.*

Remark 2.1. (i) *Assumption 2.2 ensures that the terminal control law is stable and that the satisfaction of the terminal constraint (2.7) guarantees the satisfaction of the system constraints (2.2) in Mode 2.*

(ii) *While $\mathbb{X}_f$ can be any positively invariant set under $u = Kx$, it might be convenient to use the maximal positively invariant set under this control law in order to reduce conservativeness. This set however might be defined by many inequalities, so a compromise between the size of the set and the online computational burden might be desirable.*

Let $\mathcal{X}$ be the feasibility domain of $\mathbb{P}(x)$,

$$\mathcal{X} = \{x : \mathbb{P}(x) \text{ is feasible}\}. \quad (2.14)$$

The feasibility domain is also called the domain of attraction.

Given that the control law requires the solution of $\mathbb{P}(x)$, a basic requirement for the stability is recursive feasibility of the optimization.

Definition 2.3 (Recursive feasibility). *The optimal control problem $\mathbb{P}(x)$ is said to be recursively feasible if and only if feasibility at the current time guarantees feasibility at*

the next instant (and recursively for all future instants), or equivalently, $Ax + B\kappa(x) \in \mathcal{X}$ for all $x \in \mathcal{X}$.

It is clear from this definition that if $\mathbb{P}(x)$ is recursively feasible, then $\mathcal{X}$ is a positively invariant set under the control law of (2.9), $\kappa(x)$.

Theorem 2.2. *The optimal control problem $\mathbb{P}(x_t)$ of the control law of (2.9) is recursively feasible.*

Proof. If $\mathbf{u}_t^*(x_t) = \{u_{t|t}^*, \dots, u_{t+N-1|t}^*\}$ is the optimal solution of $\mathbb{P}(x_t)$, then it defines a feasible sequence of states in Mode 1, $\mathbf{x}_t^*(x_t) = \{x_t, x_{t+1|t}^*, \dots, x_{t+N|t}^*\}$, such that $\mathbf{x}_t^*(x_t)$ and $\mathbf{u}_t^*(x_t)$ satisfy (2.5) and (2.7). Based on the *receding horizon* strategy, candidate solutions at time $t + 1$ are built from the optimal solutions at time t by discarding the first element of the sequence, which has already been applied to the system, and adding one more element at the end, as would be obtained by the terminal control law. In this thesis this process is called *receding* operation, and yields a *receded* solution given by $\tilde{\mathbf{u}}_{t+1} = \{\tilde{u}_{t+1|t+1}, \dots, \tilde{u}_{t+N-1|t+1}, \tilde{u}_{t+N|t+1}\} = \{u_{t+1|t}^*, \dots, u_{t+N-1|t}^*, Kx_{t+N|t}^*\}$ and $\tilde{\mathbf{x}}_{t+1} = \{\tilde{x}_{t+1|t+1}, \dots, \tilde{x}_{t+N|t+1}, \tilde{x}_{t+N+1|t+1}\} = \{x_{t+1|t}^*, \dots, x_{t+N|t}^*, (A + BK)x_{t+N|t}^*\}$, which satisfy (2.5). Satisfaction of (2.5) for $k = 0, \dots, N - 2$ comes directly from the feasibility of Mode 1 at the previous instant, and satisfaction for $k = N - 1$ follows because $x_{t+N|t}^* \in \mathbb{X}_f$, which guarantees that $x_{t+N|t}^* \in \mathbb{X}$ and $Kx_{t+N|t}^* \in \mathbb{U}$. Finally satisfaction of (2.7) follows from the invariance of $\mathbb{X}_f$, which implies that $(A + BK)x_{t+N|t}^* \in \mathbb{X}_f$. $\tilde{\mathbf{u}}_{t+1}$ is then a feasible solution for the optimal control problem $\mathbb{P}(x)$ at time $t + 1$, provided that the problem was feasible at time t . By a recursive argument, it remains feasible for all future instants. $\square$

The effect of N on the size of the domain of attraction is analysed next. To make the effect of N explicit, the domain of attraction and the optimal problem when using a Mode 1 horizon of length N are denoted by $\mathcal{X}_N$ and $\mathbb{P}_N(x)$.

Theorem 2.3.

- (i) $\mathbb{X}_f \subseteq \mathcal{X}_N$ for all $N = 0, 1, 2, \dots$
- (ii) $\mathcal{X}_N \subseteq \mathcal{X}_{N+1}$ for all $N = 0, 1, 2, \dots$
- (iii) If $\mathcal{X}_N = \mathcal{X}_{N+1}$, then $\mathcal{X}_N = \mathcal{X}_{N+k}$ for all $k = 0, 1, 2, \dots$

Proof. (i) Since $\mathbb{X}_f$ is positively invariant under $u = Kx$ for system (2.1) and constraints (2.2), if $x_t \in \mathbb{X}_f$ then $\mathbf{u}_t = \{Kx_t, Kx_{t+1|t}, \dots, Kx_{t+N-1|t}\}$ is a feasible input sequence for problem $\mathbb{P}_N(x)$. Then $x \in \mathcal{X}_N$ which yields $\mathbb{X}_f \subseteq \mathcal{X}_N$.

(ii) If $x_t \in \mathcal{X}_N$ then there exists a $\mathbf{u}_t = \{u_{t|t}, u_{t+1|t}, \dots, u_{t+N-1|t}\}$ that is a feasible solution of $\mathbb{P}_N(x_t)$. According to the dual-mode prediction paradigm, the first implicit input of that sequence in Mode 2 is $u_{t+N|t} = Kx_{t+N|t}$. Because $\mathbb{X}_f$ is positively invariant and satisfies $\mathbb{X}_f \subseteq \mathbb{X}$ and $K\mathbb{X}_f \subseteq \mathbb{U}$, then $x_{t+N|t} \in \mathbb{X}$, $u_{t+N|t} \in \mathbb{U}$ and $x_{t+N+1|t} \in \mathbb{X}_f$, which implies that the sequence $\mathbf{u}_t = \{u_{t|t}, u_{t+1|t}, \dots, u_{t+N-1|t}, Kx_{t+N|t}\}$ satisfies the constraints of $\mathbb{P}_{N+1}(x)$. Then $x_t \in \mathcal{X}_{N+1}$ which yields the desired result.

(iii) $\mathcal{X}_{N+1}$ is also the set of states x for which there exists an input u such that $x \in \mathbb{X}$, $u \in \mathbb{U}$ and $Ax + Bu \in \mathcal{X}_N$. If $\mathcal{X}_{N+1} = \mathcal{X}_N$, it follows that $\mathcal{X}_{N+2} = \mathcal{X}_{N+1}$ and then applying this recursively yields $\mathcal{X}_{N+t} = \mathcal{X}_{N+t-1} = \dots = \mathcal{X}_N$. $\square$

This theorem implies that (i) the domain of attraction is at least as big as the terminal set; (ii) if N is increased the domain of attraction of the MPC strategy becomes larger, or at worst remains the same; and (iii) if N is increased by 1 and the domain of attraction remains the same, then further increasing N will not change the domain of attraction.

Before analysing the stability and convergence of the system, the basic notions of stability are introduced and a theorem that guarantees stability for a system is presented. These will be presented in general for non-linear systems since the controlled dynamics are non-linear as shown in (2.13).

Consider a system of the form

$$x_{t+1} = f(x_t), \quad (2.15)$$

such that the origin is an equilibrium (i.e. $f(0) = 0$).

Definition 2.4 (Uniform local stability [12]). *System (2.15) is said to be Uniformly Locally Stable if for all $\epsilon \geq 0$ there exists δ such that $\|x_0\| \leq \delta$ implies $\|x_t\| \leq \epsilon$ for all $t = 0, 1, 2, \dots$*

Let $\mathcal{S}$ be a neighbourhood of the origin. System (2.15) is said to be Uniformly Locally Asymptotically Stable with domain of attraction $\mathcal{S}$ if it is Uniformly Locally Stable and, for all $x_0 \in \mathcal{S}$, $x_t \rightarrow 0$ as $t \rightarrow \infty$.

Definition 2.5 (Exponential stability). *System (2.15) is said to be Exponentially Stable with domain of attraction $\mathcal{S}$ if and only if there exists $\mu > 0$ and $0 < \lambda < 1$ such that $\|x_t\| \leq \mu \lambda^t \|x_0\|$ for all $x_0 \in \mathcal{S}$.*

Definition 2.6 (Lyapunov function inside a set). *The continuous positive definite function V is said to be a Lyapunov function inside $\mathcal{S}$ for system (2.15) if there exists ν such that $\mathcal{S} \subseteq \{x : V(x) \leq \nu\}$ and for all $x \in \{x : V(x) \leq \nu\}$ the inequality*

$$V(f(x)) - V(x) \leq -\phi(\|x\|) \quad (2.16)$$

holds for some $\mathcal{K}$ -function ϕ .

Theorem 2.4. *Assume system (2.15) admits a Lyapunov function V in $\mathcal{S}$. Then system (2.15) is Uniformly Locally stable. If V is also upper and lower polynomially bounded, then system (2.15) is Uniformly Locally Asymptotically (and exponentially) stable with domain of attraction $\mathcal{S}$.*

The following assumptions made about the terminal cost and control law are relevant for the stability of the MPC strategy.

Assumption 2.3. *The terminal cost given by $F(x) = x^T P x$, with $P \succeq 0$, satisfies the Lyapunov condition*

$$(A + BK)^T P (A + BK) - P \leq -(Q + K^T R K). \quad (2.17)$$

The main results regarding stability and convergence of the MPC strategy are presented next.

Theorem 2.5. *For system (2.1) under the control law of (2.9):*

(i) *For all $x_t \in \mathcal{X}$, $J^*(x_t)$ is monotonically decreasing such that*

$$J^*(Ax_t + B\kappa(x_t)) - J^*(x_t) \leq -(x_t^T Q x_t + \kappa(x_t)^T R \kappa(x_t)). \quad (2.18)$$

(ii) *For all $x_0 \in \mathcal{X}$, the closed-loop performance satisfies the bound*

$$\sum_{t=0}^{\infty} (x_t^T Q x_t + \kappa(x_t)^T R \kappa(x_t)) \leq J^*(x_0). \quad (2.19)$$

(iii) *For all $x_0 \in \mathcal{X}$, $(x_t^T Q x_t + \kappa(x_t)^T R \kappa(x_t)) \rightarrow 0$ as $t \rightarrow \infty$.*

Proof. (i) From Theorem 2.2 the *receded* solutions are feasible and satisfy:

$$\begin{aligned} J(x_{t+1}, \tilde{\mathbf{u}}_{t+1}) &= \sum_{k=0}^{N-1} \ell(\tilde{x}_{t+1+k|t+1}, \tilde{u}_{t+1+k|t+1}) + F(\tilde{x}_{t+N+1|t+1}) \\ &= \sum_{k=1}^N \ell(x_{t+k|t}^*, u_{t+k|t}^*) + F((A + BK)x_{t+N|t}^*) \\ &= \sum_{k=0}^{N-1} \ell(x_{t+k|t}^*, u_{t+k|t}^*) - \ell(x_{t|t}^*, u_{t|t}^*) + \ell(x_{t+N|t}^*, Kx_{t+N|t}^*) \\ &\quad + F((A + BK)x_{t+N|t}^*) \end{aligned}$$

The Lyapunov condition of (2.17) then implies that

$$\begin{aligned} J(x_{t+1}, \tilde{\mathbf{u}}_{t+1}) &\leq \sum_{k=0}^{N-1} \ell(x_{t+k|t}^*, u_{t+k|t}^*) - \ell(x_{t|t}^*, u_{t|t}^*) + F(x_{t+N|t}^*) \\ &= J(x_t, \mathbf{u}_t^*) - \ell(x_{t|t}, u_{t|t}^*). \end{aligned}$$

But since $J^*(x_t) = J(\mathbf{x}_t^*, \mathbf{u}_t^*)$ and an optimization is performed to yield the optimal value of the cost, it follows that $J^*(x_{t+1}) = J(\mathbf{x}_{t+1}^*, \mathbf{u}_{t+1}^*) \leq J(\tilde{\mathbf{x}}_{t+1}, \tilde{\mathbf{u}}_{t+1})$. Thus $J^*(x_{t+1}) - J(x_t) \leq -\ell(x_{t|t}, u_{t|t}^*)$, which is equivalent to (2.18).

(ii)-(iii) Summing (2.18) for $t = 0, 1, 2, \dots$, it is obtained that $\lim_{\bar{t} \rightarrow \infty} \sum_0^{\bar{t}} (x_t^T Q x_t + u_t^{*T} R u_t^*) + J^*(x_{\bar{t}}) \leq J^*(x_0)$. The desired property of (ii) follows directly since $J^*(\cdot)$ is non-negative. Also, since the optimization is feasible, $J^*(x_0)$ must be finite, which together with (ii) imply that $(x_t^T Q x_t + \kappa(x_t)^T R \kappa(x_t)) \rightarrow 0$ as $t \rightarrow \infty$. $\square$

Theorem 2.6. *System (2.1) under the control law of (2.9):*

- (i) *Satisfies $\lim_{t \rightarrow \infty} x_t = 0$ and $\lim_{t \rightarrow \infty} u_t = 0$, if $x_0 \in \mathcal{X}$.*
- (ii) *Is asymptotically stable with domain of attraction $\mathcal{X}$.*
- (iii) *If $Q \succ 0$, then it is exponentially stable with domain of attraction $\mathcal{X}$.*

Proof. (i) From Theorem 2.5 $\kappa(x_t)^T R \kappa(x_t) \rightarrow 0$ as $t \rightarrow \infty$, and since $R \succ 0$, it follows that $\kappa(x_t) \rightarrow 0$ as $t \rightarrow \infty$. Also from Theorem 2.5, $x_t^T Q x_t \rightarrow 0$ as $t \rightarrow \infty$, and from the detectability of $(A, Q^{\frac{1}{2}})$ it follows that $x_t \rightarrow 0$ as $t \rightarrow \infty$.

(ii) From the facts that $\ell(x, u)$ and $F(x)$ are quadratic with $Q, P \succeq 0$ and $Q \succ 0$ and that the pair $(A, Q)^{\frac{1}{2}}$ is detectable, it follows that the objective function of $\mathbb{P}(x_t)$ is strictly convex and positive definite. And since condition (2.12) is polyhedral, it follows that the value function $J^*(x_t)$ is a continuous, strictly convex, piecewise quadratic function of x_t [7]. It also satisfies (2.18) which implies that it is a Lyapunov function, and therefore the origin is stable. Additionally, since $x_t \rightarrow 0$ as $t \rightarrow \infty$ whenever $x_0 \in \mathcal{X}$, it follows that the origin is asymptotically stable with domain of attraction $\mathcal{X}$.

(iii) Since $J^*(x_t)$ is a positive definite, piecewise quadratic function of x_t , it can be bounded from above and below, for all $x_t \in \mathcal{X}$, by

$$\alpha \|x_t\|^2 \leq J^*(x_t) \leq \beta \|x_t\|^2 \quad (2.20)$$

where $\alpha, \beta > 0$ since $J^*(x_t)$ is positive definite. If the smallest eigenvalue of Q is $\underline{\lambda}(Q)$, it follows from (2.18) and (2.20) that, for all $x_0 \in \mathcal{X}$,

$$\|x_t\|^2 \leq \frac{1}{\alpha} \left| 1 - \frac{\underline{\lambda}(Q)}{\beta} \right|^t J^*(x_0),$$

which indicates an exponential decrease in the norm of x_t as t increases, and therefore the origin is exponentially stable with domain of attraction $\mathcal{X}$. $\square$

Remark 2.7. *It is common that K is chosen to be the LQR optimal gain for the cost given by $J_\infty(x_0) = \sum_{t=0}^{\infty} x_t^T Q x_t + u_t^T R u_t$, and that P satisfies the Lyapunov equation $P - (A + BK)^T P (A + BK) = Q + K^T R K$ (so that $J(x) = J_\infty(x)$). In this case, since $\mathbb{X}_f$ is positively invariant and guarantees feasibility for $u = Kx$, it is clear that for $x_t \in \mathbb{X}_f$ the sequence of inputs given by $\mathbf{u}_t = \{Kx_t, Kx_{t+1|t}, \dots, Kx_{t+N-1|t}\}$ is feasible, and since $u = Kx$ is optimal in absence of constraints, it will be the optimal solution of (2.8) and therefore by (2.9) $\kappa(x_t) = Kx_t$.*

Note that since the origin is asymptotically stable and that $\mathbb{X}_f$ contains the origin in its interior there exists a time $\bar{t}$ such that $x_t \in \mathbb{X}_f$ for all $t \geq \bar{t}$. This means that there will always exist an instant when the control law becomes the unconstrained optimal, namely $u_t = \kappa(x_t) = Kx_t$ for all $t \geq \bar{t}$.

2.1.3 Different types of cost

The analysis presented in the previous section considers a quadratic cost function. A similar analysis can be performed for a cost function

$$J(\mathbf{x}_t, \mathbf{u}_t) = \sum_{k=0}^{N-1} \ell(x_{t+k|t}, u_{t+k|t}) + F(x_{t+N|t}),$$

where the stage cost, $\ell(x, u)$, and the terminal cost, $F(\cdot)$, are more generic. The stage cost is convex and satisfies $\ell(0, 0) = 0$ and $\ell(x, u) > 0$ if $x \neq 0$ or $u \neq 0$, and the terminal cost is convex and satisfies $F(x) \geq 0$ for all x .

The feasibility analysis is the same, because only the cost changes, and then Theorems 2.2 and 2.3 still hold. For the stability and convergence analysis, if the terminal cost satisfies the decrease condition

$$F((A + BK)x) - F(x) \leq -\ell(x, u), \quad (2.21)$$

then the same results presented in Theorems 2.5 and 2.6 can be obtained.

The most typical choice is a cost with 1-norm or ∞ -norm terms. In this case $\ell(x, u) = |Qx|_p + |Ru|_p$, $F(x) = |Px|_p$ where $p = 1$ or $p = \infty$, $Q \in \mathbb{R}^{n \times n}$ and $R \in \mathbb{R}^{n_u \times n_u}$ are non-singular and $P \in \mathbb{R}^{r \times n}$ has full column rank. The stage function and the terminal cost can be reformulated as linear functions with the addition of extra variables and linear constraints. Then the cost is linear and the constraints are linear too, thus defining a linear program (LP) as shown next for the case of $p = \infty$:

$$\begin{aligned}
& \underset{\mathbf{x}_t, \mathbf{u}_t}{\text{minimize}} && \sum_{k=0}^{N-1} \alpha_k + \beta_k + \alpha_N \\
& \text{subject to} && x_{t|t} = x_t \\
& && x_{t+k+1|t} = Ax_{t+k|t} + Bu_{t+k|t}, \quad k = 0, \dots, N-1 \\
& && Fx_{t+k|t} \leq 1, \quad k = 0, \dots, N-1, \\
& && Gu_{t+k|t} \leq 1, \quad k = 0, \dots, N-1, \\
& && Hx_{t+k|t} \leq 1 \\
& && \begin{bmatrix} Q \\ -Q \end{bmatrix} x_{t+k|t} \leq \underline{1}\alpha_k, \quad k = 0, \dots, N-1 \\
& && \begin{bmatrix} R \\ -R \end{bmatrix} x_{t+k|t} \leq \underline{1}\beta_k, \quad k = 0, \dots, N-1 \\
& && \begin{bmatrix} P \\ -P \end{bmatrix} x_{t+N|t} \leq \underline{1}\alpha_N.
\end{aligned} \tag{2.22}$$

2.2 Robust Model Predictive Control

The classical MPC formulation described in Section 2.1 does not consider uncertainties, and is also referred to as nominal MPC. Although the closed-loop application of nominal MPC inherently provides some robustness, which has been analysed in several works [93], [30], [60], not accounting explicitly for the presence of uncertainty can severely affect the performance of the controlled system and the satisfaction of constraints [86]. Consequently, robust MPC formulations (RMPC) have emerged to guarantee performance, constraint satisfaction and stability of the systems in the presence of constraints and uncertainties. Early developments [18], [2], [95] considered single-input, single-output (SISO) frequency impulse-response (FIR) plants, which for certain choices of the objective function, yield online problems that can be reduced to

linear programs. However, these approaches can only be applied to stable plants and, in order to simplify to online computational complexity one must choose simplistic and many times unrealistic models.

Thus typical modern formulations consider state-space linear systems with additive and multiplicative uncertainty defined by

$$x_{t+1} = A_t x_t + B_t u_t + D_t w_t. \quad (2.23)$$

Future realizations of the uncertainty $(A_{t+k}, B_{t+k}, D_{t+k}, w_{t+k})$ are unknown at time t (where $k = 0, 1, 2, \dots$), but they are known to satisfy

$$(A_t, B_t, D_t, w_t) \in \Theta, \quad t = 0, 1, 2, \dots \quad (2.24)$$

where Θ is the admissible uncertainty set. The disturbances w_t are the additive uncertainty, while the uncertain model parameters (A_t, B_t, D_t) are the multiplicative uncertainty. This is a very general description of the uncertainty, and most strategies consider a more precise description. Additionally, constraints defined by (2.2)

$$x_t \in \mathbb{X}, \quad u_t \in \mathbb{U},$$

are also considered.

2.2.1 Invariance and stability for Robust MPC

Feasibility of the predictions, invariance, the cost functions and stability are the most relevant aspects which are treated in RMPC differently from classical MPC.

Under the presence of uncertainty there is more than just a single predicted sequence, so then feasibility of the constraints in Mode 1 and the terminal constraint have to be guaranteed for all realizations of the uncertainty.

Regarding invariance, Definition 2.1 is modified to include the presence of uncertainty.

Definition 2.7 (Robust positively invariant set). *A set $\Omega \subseteq \mathbb{X}$ is said to be robust positively invariant for system (2.23) and constraints (2.2) under a control law given by $u = \kappa(x)$, if and only if, $Ax + B\kappa(x) + Dw \in \Omega$ and $\kappa(x) \in \mathbb{X}$ for all $(A, B, D, w) \in \Theta$ and for all $x \in \Omega$.*

Then, for any realization of the additive and model uncertainty, if $x \in \Omega$, it is guaranteed that x will remain in Ω for all future instants. Additionally, $x \in \Omega$ guarantees the satisfaction of the system constraints, and then their satisfaction is guaranteed for all future instants. The definition of the maximal robust positively invariant set is analogous to the case of the maximal positively invariant set in Section 2.1.2.

Regarding stability, the definition of Local uniform stability and Theorem 2.4 are unchanged except that the decrease property (2.16) has to hold for any realization of the uncertainty. This gives rise to what is referred to as robust stability. However, when additive disturbances are considered, it is not possible to drive the state to the origin, and therefore the usual notions of stability with respect to the origin cannot be used.

Alternative stability notions have then been considered in the literature. Some of the most relevant are Input-to-State Stability (ISS) [41] and Lyapunov functions that define stability with respect to a set rather than a point [12]. These notions will be presented in general for non-linear systems since the dynamics of system (2.23) under a control law given by $\kappa(x_t)$ become $x_{t+1} = A_t x_t + B_t \kappa(x_t) + D_t w_t$, which is in general non-linear since in general $\kappa(x_t)$ is non-linear.

Consider a time-invariant non-linear system with dynamics given by

$$x_{t+1} = f(x_t, \delta_t), \tag{2.25}$$

where δ_t is an uncertain parameter that lies in a compact set Δ , and $f(\cdot, \cdot)$ is continuous such that the origin is an equilibrium in the absence of disturbances (i.e. $f(0, 0) = 0$). Note that this definition includes systems with both additive and multiplicative uncertainty; $x_{t+1} = f(x_t, 0)$ corresponds to the nominal dynamics (for $w_t = 0$ and some nominal values of (A_t, B_t, D_t)), and δ_t defines w_t and the deviation of (A_t, B_t, D_t) from the nominal values.

Definition 2.8 (Input-to-state stability). *System (2.25) is Input-to-State Stable (ISS) with domain of attraction $\mathcal{S} \subseteq \mathbb{R}^n$, which contains the origin in its interior, if there exist a $\mathcal{KL}$ -function $\beta(\cdot, \cdot)$ and a $\mathcal{K}$ -function $\gamma(\cdot)$ such that, for all $x_0 \in \mathcal{S}$ and all admissible disturbances $\delta \in \Delta$, the evolution of (2.25) satisfies for all $t = 0, 1, 2, \dots$*

$$|x_t| \leq \beta(\|x_0\|, t) + \gamma(\sup\{\|\delta_k\| : k \in \mathbb{N}_{[0, t-1]}\}) \quad (2.26)$$

Lemma 2.8. *Let $\mathcal{S} \subseteq \mathbb{R}^n$ contain the origin in its interior and be a robust positively invariant set for system (2.25). Furthermore, let there exist $\mathcal{K}_\infty$ -functions $\alpha_1(\cdot)$, $\alpha_2(\cdot)$ and $\alpha_3(\cdot)$ and a function $V : \mathcal{S} \rightarrow \mathbb{R}_+$ that is Lipschitz continuous on $\mathcal{S}$ such that for all $x \in \mathcal{S}$,*

$$\alpha_1(\|x\|) \leq V(x) \leq \alpha_2(\|x\|) \quad (2.27a)$$

$$V(f(x, 0)) - V(x) \leq -\alpha_3(\|x\|) \quad (2.27b)$$

$V(\cdot)$ is a ISS-Lyapunov function and system (2.25) is ISS with domain of attraction $\mathcal{S}$ if either (i) $f(\cdot, \cdot)$ is Lipschitz continuous on $\mathcal{S} \times \Delta$, or (ii) $f(x, \delta) := g(x) + \delta$, where $g : \mathcal{S} \rightarrow \mathbb{R}^n$ is continuous on $\mathcal{S}$.

Remark 2.9. *ISS implies that: (i) system $x^+ = f(x, 0)$ is asymptotically stable with domain of attraction $\mathcal{S}$; (ii) all state trajectories of system (2.25) are bounded for bounded δ_t ; and (iii) as $t \rightarrow \infty$ all trajectories of system (2.25) go to the origin if*

$\delta_t \rightarrow 0$.

Definition 2.9 (Lyapunov function outside a set). *The continuous function $V(\cdot)$, $V(x) \geq 0$, is a Lyapunov function outside $\mathcal{S}$ for system (2.25) if $\mathcal{S} = \{x : V(x) = 0\}$, and for all $x \notin \mathcal{S}$ the inequality*

$$V(f(x, \delta)) - V(x) \leq -\phi(\|x\|),$$

holds for some κ -function $\phi(\cdot)$ and for all $w \in \mathbb{W}$.

Definition 2.10 (Uniform robust stability with respect to a set). *Let $\mathcal{S}$ be a neighborhood of the origin and let $d(x, X)$ be a distance function of x to the set X . System (2.25) is said to be uniformly robustly stable w.r.t. $\mathcal{S}$ if for all $\epsilon > 0$ there exists μ such that $d(x_0, \mathcal{S}) \leq \mu$ implies that $d(x_t, \mathcal{S}) \leq \epsilon$ for all $t = 0, 1, 2, \dots$ and all functions $w_t \in \mathbb{W}$.*

Definition 2.11 (Robust exponential stability). *System (2.25) is said to be Robustly Exponentially Stable w.r.t. $\mathcal{S}$ if and only if there exists $\mu > 0$ and $0 < \lambda < 1$ such that $\text{dist}(x_t, \mathcal{S}) \leq \mu \lambda^t \text{dist}(x_0, \mathcal{S})$ for all $t > 0$.*

Theorem 2.10. *Assume system (2.25) admits a Lyapunov function $V(\cdot)$ outside $\mathcal{S}$. Then it is uniformly robustly stable w.r.t. $\mathcal{S}$. If V is also upper and lower bounded by polynomial functions of the distance $d(x, \mathcal{S})$, then system (2.25) is uniformly robustly asymptotically (and exponentially stable) w.r.t. $\mathcal{S}$.*

As it has been mentioned earlier, uniform stability w.r.t. a set is particularly relevant to MPC strategies for systems with additive uncertainty since due to its presence the state cannot converge to zero. A particularly interesting set is the minimal robust positively invariant set.

Definition 2.12 (Minimal positively invariant set). *The set $\underline{\Omega} \subseteq \mathbb{X}$ is said to be the minimal positively invariant set for system (2.23) under the control law $u = \kappa(x)$ if*

and only if it is positively invariant and $\underline{\Omega} \subseteq \Omega$ for all robust positively invariant sets Ω .

For the case of additive uncertainty only, the minimal robust positively invariant set under the linear control law $u = Kx$ is given by [44]

$$\Omega_{\infty}(\mathbb{W}) = \bigoplus_{k=0}^{\infty} (A + BK)^k \mathbb{W}. \quad (2.28)$$

In this case, the system is robustly exponentially stable with respect to $\Omega_{\infty}(\mathbb{W})$.

Since the presence of uncertainty implies that there is more than just a single predicted sequence, the definition of a suitable cost function of the predicted sequences is not as obvious as in the previous section. Some of the typical alternatives are:

- Worst-case costs (whose minimization forms min-max problems).
- Costs that are zero inside the terminal set and greater than zero outside.
- Nominal costs. Here the presence of the uncertainty is disregarded and a nominal value for the model matrices is used. The costs can then take the form as those used in the previous section.

Worst case costs are normally used for systems with multiplicative uncertainty, and in these cases can yield asymptotic stability. For the additive uncertainty case, typical choices are nominal costs (that provide ISS stability) and costs that are zero in the terminal set and greater than zero outside (that yield stability with respect to that set).

The aspects introduced here are handled in a different manner by different strategies. More details are given in the review of the most relevant works in Section 2.2.3.

2.2.2 Open-loop versus feedback Robust MPC

Early developments of RMPC consisted of formulations that relied on open-loop optimal control problems [95], [64], where the decision variables are a sequence of control inputs $\mathbf{u}_t = \{u_{t|t}, \dots, u_{t+N-1|t}\}$ that guarantee the satisfaction of the constraints for any possible uncertainty realization, and minimize an appropriate criterion (usually worst case performance).

These strategies are conservative because they do not take into account the effect of the future uncertainties on future states, that are currently unknown by the controller, but which effect will be known when the future states are measured and that information is available to compute the inputs to be applied to the system. Feedback RMPC formulations consider closed-loop optimal control problems, as opposed to open-loop, and do take into account this information in the predictions by having as a decision variable a control policy $\pi(x_t)$, instead of a sequence of control inputs. A control policy is a sequence of control laws $\{\mu_0(x_t), \mu_1(x_{t+1|t}), \dots, \mu_{N-1}(x_{t+N-1|t})\}$ that are functions that depend on the state at the instant they would have to be applied (which in turn depend on the uncertainties that have affected the system).

When optimizing over control policies, future predicted inputs $u_{t+k|t}$ will depend on the future predicted states $x_{t+k|t}$, which in turn depend on the unknown current and future uncertainties. This may seem to present difficulties since the future predicted inputs $u_{t+k|t}$ will depend on unknown information, and are thus unknown at time t . This is not a problem for two reasons. First, only $u_{t|t}$ is actually applied to the system, and $u_{t|t}$ is known because it depends only on x_t which is known at time t . And second, instead of knowing the values of the future predicted $u_{t+k|t}$, in the feedback RMPC formulation one seeks to know that there exist control laws that will be feasible at time $t+k$ and use only information that will be available at the time they would be applied (i.e., so that they are causal), but there is no need to actually find these control laws. Clearly, $u_{t+k|t} = \mu_k(x_{t+k|t})$ depends on $x_{t+k|t}$, which will be known at

time $t + k$, so the predicted control laws are causal.

The use in the predicted control policies of the information that is unavailable at the current instant but will be available in the future offers various advantages, such as improvements in the performance or feasibility in cases where open-loop formulations would be infeasible. The most general solution of the closed-loop optimal control problem, in terms of considering a general class of control policies, is obtained via dynamic programming (DP) [9], which is computationally prohibitively complex. Consequently, researchers have focused on simplifications of the structure of the control policies to obtain simpler optimal control formulations in order to strike an appropriate compromise between sub-optimality (relative to the exact DP solution) and complexity of the online problems.

2.2.3 Review of Robust MPC with multiplicative uncertainty

2.2.3.1 Optimized state linear feedback

In [45] a robust feedback MPC scheme is proposed for linear, discrete time systems, with multiplicative uncertainty

$$x_{t+1} = A_t x_t + B_t u_t, \quad (2.29)$$

where $(A_t, B_t) \in \Theta = \text{conv} \{(A_i, B_i) : i = 1, \dots, \nu_m\}$, that are subject to state and input constraints as in (2.2).

The strategy of [45] avoids the exponential complexity of DP by optimizing at each instant t a gain K_t such that the predicted control policy is a feedback state controller

$$u_{t+k|t} = K_t x_{t+k|t}, \quad k = 0, 1, \dots \quad (2.30)$$

The optimization of K_t at each instant t aims to minimize the worst case of the

infinite horizon cost $J_\infty(x_t) = \sum_{k=0}^{\infty} x_{t+k|t}^T Q x_{t+k|t} + u_{t+k|t}^T R u_{t+k|t}$. However, since the optimization of the worst case would require consideration of all possible sequences of extreme realizations of the uncertain matrices A_t and B_t , this is performed by minimizing over K_t and V_t an attainable upper bound $V(x_t) = x_t^T P_t x_t$ of the cost $J_\infty(x_t)$.

So that $V(x_t)$ is an upper bound of $J_\infty(x_t)$, P_t has to satisfy the decrease condition

$$V((A_{t+k} + B_{t+k}K_t)x_{t+k|t}) - V(x_{t+k|t}) \leq -(x_{t+k|t}^T Q x_{t+k|t} + u_{t+k|t}^T R u_{t+k|t}), \quad (2.31)$$

for all $(A_{t+k}, B_{t+k}) \in \Theta$ and $P_t \geq 0$. This condition is non-linear due to the products $P_t K_t$, however it can be convexified by the change of variables $K_t = Y_t Q_t^{-1}$ and $Q_t = P_t^{-1}$. Condition (2.31) implies that the ellipsoid defined by $E_t = \{x : x^T Q_t^{-1} x \leq 1\}$ is robustly invariant. Thus, if one imposes that E_t is such that for all $x_t \in E_t$ then $x_t \in \mathbb{X}$ and $K_t x_t \in \mathbb{U}$, then the satisfaction of the system constraints will be guaranteed for the current and all future prediction instants if $x_t \in E_t$.

The conditions that E_t is such that for all $x \in E_t$ then $x \in \mathbb{X}$ and $K_t x \in \mathbb{U}$, $x_t \in E_t$ and (2.31) and can be enforced through the use of a number of linear matrix inequalities (LMIs), and hence need only be invoked at the vertices (A_i, B_i) of Θ . Then, the optimal control problem consisting in the minimization of $V(x_t)$ subject to all the previously mentioned conditions and can be cast as a semi-definite program (SDP). The invariance of E_t and the decrease property of (2.31) allow one to prove recursive feasibility and asymptotic stability.

2.2.3.2 Augmented dynamics for efficient RMPC

Since the strategy presented in the previous section requires one to solve an SDP with several LMIs, [51] proposes an augmented state formulation which allows one to robustly stabilize a system by imposing only one LMI. The strategy considers the

same type of uncertain systems as in (2.29)

$$x_{t+1} = A_t x_t + B_t u_t,$$

and mixed state and input constraints of the type

$$(x_t, u_t) \in \mathbb{Y} = \{(x, y) : Fx + Gu \leq \underline{1}\}, \quad (2.32)$$

with $F \in \mathbb{R}^{n_c \times n}$ and $G \in \mathbb{R}^{n_c \times n_u}$. The formulation is enhanced with the closed-loop paradigm [82], which consists of using additional terms $c_{t+k|t}$ such that the predicted inputs are given by

$$\begin{aligned} u_{t+k|t} &= Kx_{t+k|t} + c_{t+k|t}, \quad k = 0, \dots, N-1, \\ u_{t+k|t} &= Kx_{t+k|t}, \quad k \geq N, \end{aligned} \quad (2.33)$$

where $c_{t+k|t}, k = 0, \dots, N-1$ are the degrees of freedom (dof) of the inputs. The resulting system $x_{t+k+1|t} = \Phi_{t+k}x_{t+k|t} + B_{t+k}c_{t+k|t}$, where $\Phi_{t+k} = A_{t+k} + B_{t+k}K$ and $c_{t+k|t} = 0, k \geq N$, can also be described by the autonomous augmented state space model

$$z_{t+k+1|t} = \Psi_{t+k}z_{t+k|t}, \quad k = 0, 1, \dots, \quad \text{where } z_{t|t} = \begin{bmatrix} x_t \\ \mathbf{c}_t \end{bmatrix}, \quad \mathbf{c}_t = \begin{bmatrix} c_{t|t} \\ c_{t+1|t} \\ \vdots \\ c_{t+N-1|t} \end{bmatrix}, \quad (2.34)$$

and

$$\Psi_{t+k} = \begin{bmatrix} \Phi_{t+k} & \begin{bmatrix} B_{t+k} & 0 & \cdots & 0 \end{bmatrix} \\ 0 & M \end{bmatrix}, \quad M = \begin{bmatrix} 0 & I & 0 & \cdots \\ 0 & 0 & \ddots & 0 \\ \vdots & & \ddots & I \\ 0 & \cdots & 0 & 0 \end{bmatrix}, \quad (2.35)$$

where $I \in \mathbb{R}^{n_u \times n_u}$. Assuming that it is possible to find matrices Q_x, K such that $E_x = \{x : x^T Q_x x \leq 1\}$ is robustly invariant for the system (2.29) and constraints (2.32) under $u = Kx$, then it is also possible to find matrices Q_z, K such that $E_z = \{z : z^T Q_z z \leq 1\}$ is robustly invariant for the augmented closed-loop system defined by (2.34) and for constraints (2.32).

Consideration is given the projection of E_z onto the x -space, E_{xz} which is the ellipsoid that contains all x for which there exists at least one $\mathbf{c}$ such that $z = [x^T, \mathbf{c}^T]^T \in E_z$. It is proved that $E_x \subseteq E_{xz}$, and that E_{xz} grows in size with the prediction horizon N . The conditions for invariance and feasibility are given by

$$\begin{bmatrix} Q_z & Q_z \Psi_i^T \\ * & Q_z \end{bmatrix} \succeq 0, \quad i = 1, \dots, \nu_m, \quad \begin{bmatrix} Z & [F + GK, GM] \\ * & Q_z \end{bmatrix} \succeq 0, \quad Z_{ii} \leq 1, \quad i = 1, \dots, n_c, \quad (2.36)$$

where the symmetric matrix $Z \in \mathbb{R}^{n_c \times n_c}$ is a free variable with elements $Z_{p,q}$, $p, q = 1, \dots, n_c$ and $*$ denotes symmetric blocks. The goal is to maximize the volume of E_{xz} . This is achieved by solving $\max \log \det(TQ_z T^T)$ subject to (2.36), where T is such that $x = Tz$. However, if X_0 is the set of relevant initial conditions, it is also desired that E_{xz} contains X_0 .

Based on all these arguments the control algorithm consists of an offline and online stage. In the offline stage, K is computed ignoring constraints so as to optimize robust performance in some sense, which could be a worst case cost or nominal performance,

or K could be taken to be an $\mathcal{H}_2/\mathcal{H}_\infty$ controller. Then for this K , Q_z is chosen so as to maximize the volume of E_{xz} . If $X_0 \in E_{xz}$ the MPC scheme can be applied online. Otherwise, N is increased and the offline design is repeated. The online stage is repeated at each sampling instant and consists of performing the optimization:

$$\begin{aligned} & \underset{\mathbf{c}_t}{\text{minimize}} && \mathbf{c}_t^T W \mathbf{c}_t \\ & \text{subject to} && \mathbf{z}_t^T Q_z \mathbf{z}_t \leq 1 \end{aligned} \tag{2.37}$$

where $W \succ 0$, and then only the first element $c_{t|t}^*$ of the optimizing $\mathbf{c}_t^*$ is implemented. This optimization problem can be cast as an SDP since it involves a quadratic cost $J_t(\mathbf{c}_t) = \mathbf{c}_t^T W \mathbf{c}_t$ and condition $\mathbf{z}_t^T Q_z \mathbf{z}_t \leq 1$ which can be cast as an LMI by using Schur complement.

Recursive feasibility follows from the invariance for the augmented states. The strategy is proved to be robustly stabilizing towards the origin in the sense that $\mathbf{c}_t^T W \mathbf{c}_t \rightarrow 0$ and $x_t \rightarrow 0$ as $t \rightarrow \infty$ since feasibility of the receding candidate solution $\tilde{\mathbf{c}}_{t+1} = [c_{t+1|t}, \dots, c_{t+N-1|t}, 0]$ implies the decrease property $J_{t+1}(\tilde{\mathbf{c}}_{t+1}) - J_t(\mathbf{c}_t) \leq -c_{t|t}^T W_1 c_{t|t}$. It is also proved that this optimization problem can be equivalently solved by a single univariate algebraic problem, which is defined by the computation of the shortest distance of an ellipsoid from the origin; this is much more efficient than the problem presented in [45] and other previous formulations of robust MPC.

However, the strategy is still conservative for two reasons. First, because it guarantees constraint satisfaction through the use of ellipsoidal, rather than polytopic sets. Later works then use polytopic sets at the expense of increased computational cost [40], [43]. The second reason is that the control policy is only quasi closed-loop because the online optimization is open-loop for the degrees of freedom, and in spite of the presence of a feedback controller, it is designed with no special feedback treatment of future (unknown) uncertainties; it applies feedback from future states, but only with a fixed law.

2.2.3.3 Optimized dynamics RMPC

An augmented state formulation is also proposed in [19], but this time use is made of a feedback that depends of the future (unknown) uncertainties in the control policy. The strategy considers the same type of uncertain systems as in (2.29), and constraints as in (2.32). The control policy is given by

$$u_{t+k|t} = Kx_{t+k|t} + Lv_{t+k|t}, \quad k = 0, 1, \dots \quad (2.38)$$

where $v_{t+k|t}$ can be interpreted as a controller state, which satisfies the dynamics $v_{t+k+1|t} \in \text{conv}\{M_i v_{t+k|t}, i = 1, \dots, \nu_m\}$. More precisely, it takes the value associated with the realization of the uncertain model matrices, such that if

$$(A_t, B_t) = \sum_{i=1}^{\nu_m} \lambda_i (A_i, B_i), \quad (2.39)$$

for some set of scalars $\{\lambda_1, \dots, \lambda_{\nu_m}\}$, then

$$v_{t+k+1|t} = \sum_{i=1}^{\nu_m} \lambda_i M_i v_{t+k|t}. \quad (2.40)$$

The control policy depends on future realizations of the uncertainty which are unknown. However, as pointed out in Section 2.2.2, only the first element $u_{t|t}$ is actually implemented to the system, and $u_{t|t}$ only depends on x_t and $v_{t|t}$ which are known at time t .

The system and controller dynamics can be described by the augmented dynamics

$$z_{t+k+1|t} = \Psi_{t+k} z_{t+k|t}, \quad k = 0, 1, \dots, \quad \text{where } z_{t|t} = \begin{bmatrix} x_t \\ v_{t|t} \end{bmatrix}, \quad (2.41a)$$

and

$$\Psi_{t+k} \in \text{conv}\{\Psi_i, i = 1, \dots, \nu_m\}, \quad \Psi_i = \begin{bmatrix} \Phi_i & BL \\ 0 & M_i \end{bmatrix}, \quad \Phi_i = A_i + B_i K, \quad i = 1, \dots, \nu_m. \quad (2.41b)$$

K is assumed to be optimal in some sense (e.g. LQR for the nominal system), and then L and M_i are designed to optimize the polytopic dynamics of (2.41).

Invariance and feasibility of an ellipsoid $E_z = \{z : z^T P z \leq 1\}$ are equivalent to:

$$P - \Psi_i^T P \Psi_i \succeq 0 \quad (2.42a)$$

$$\begin{bmatrix} Z & \begin{bmatrix} F + GK & GL \\ P \end{bmatrix} \\ * & \end{bmatrix} \succeq 0, \quad Z_{p,p} \leq 1, \quad p = 1, \dots, n_c \quad (2.42b)$$

where $Z \in \mathbb{R}^{n_c \times n_c}$ is a free variable with elements $Z_{p,q}$, $p, q = 1, \dots, n_c$. To maximize the volume of the projection, E_{xz} , of E_z onto the x -subspace one needs to maximize $\log \det (T^T P^{-1} T)$ over P , L , M_i subject to (2.42) and (2.42b), where T is the transformation for which $x = Tz$ for all z . This optimization is non-convex because (2.42) is a bilinear matrix inequality (BMI) in P, M_i, L , however, it is shown that this problem can be convexified through the use of a transformation proposed in [83]. The maximum volume E_{xz} can be achieved by taking the dimension of v to be n_x , and this maximum volume is as large as the volume of the largest invariant ellipsoid obtained under any linear state feedback law. This strategy then is as general, in terms of domain of attraction, as the strategy of [45] as described in 2.2.3.1. However, the design of this robustly invariant set that guarantees feasibility of the system constraints is performed offline. It is also more general than the strategy of [51] as presented in Section 2.2.3.2, which is subsumed as a feasible solution of the optimized dynamics.

The online optimization aims at minimizing a worst case cost (min-max problem)

through minimizing an attainable quadratic upper bound of the minimax cost, given by $J(x_t, v_t) = x_t^T W_x x_t + v_t^T W_v v_t$, subject to $z_t = [x_t^T, v_{t|t}^T]^T \in E_z$, which is clearly less computationally involved than the strategy of [45] due to the reduced number of LMIs. Since E_z is robustly invariant under the control law of (2.38) and constraints (2.32) the predictions are feasible and the algorithm is recursively feasible. Stability follows from the fact that the cost satisfies a decrease condition which makes the cost a Lyapunov function.

As in [51], the size of the domain of attraction can be improved by considering a polytopic robustly invariant set for the online optimization instead of E_z , which however, comes at the cost of increased computational burden since in general the polytope will be defined by a greater number of inequalities.

2.2.4 Review of Robust MPC with additive uncertainty

2.2.4.1 Feedback RMPC strategies

The feedback RMPC strategies of [86], [42] consider time-invariant linear systems with bounded additive disturbances

$$x_{t+1} = Ax_t + Bu_t + w_t, \quad (2.43)$$

and state and input constraints as in (2.2)

$$u_t \in \mathbb{U}, x_t \in \mathbb{X},$$

where $\mathbb{X}, \mathbb{U}, \mathbb{W}$ are polytopic sets that contain the origin in their interiors. A dual-mode prediction based MPC strategy is considered, where the terminal set $\mathbb{X}_f$ is robust positively invariant for the system and constraints of (2.43) under $u = Kx$, which is the terminal control law that operates in Mode 2. In Mode 1 a predicted control policy is used to steer the state to $\mathbb{X}_f$.

The feedback in the predicted control policy is introduced by considering a different control sequence for every realization of the disturbances. This is implemented in practice by computing a different control sequence for every extreme realization of the disturbance, so the problem consists of finding the inputs $u_{t+k|t}^l$, $k = 0, \dots, N-1, l \in \mathcal{L}$, where $\mathcal{L}$ is the set that enumerates all extreme realizations of the disturbance and $u_{t+k|t}^l$ denotes the predicted input at time $t+k$ (made at time t) associated with the l -th disturbance realization. A causality constraint is included, which enforces that the control actions for different realizations of the disturbance if they result in the same state at some prediction instant. The online stage of the strategy consists of the minimization of worst case costs which can be cast as a linear program. However, it requires a number of constraints that grows exponentially with the prediction horizon N , due to the need to enumerate all the realizations of the disturbance for guaranteeing constraint satisfaction and the evaluation of the cost.

[58], [38] propose an approximate feedback strategy that tries to avoid the exponential increase in DP or the strategy of [86]. It considers linear time-invariant systems with additive disturbances as in equation (2.43) and mixed state and input constraints as in (2.32). A natural approach to reduce the complexity is to consider a predicted control policy where each predicted control law in Mode 1 is an affine function of the sequence of past states:

$$u_{t+k|t} = r_{t+k|t} + \sum_{j=0}^k M_{k,j} x_{t+j|t}, \quad k = 0, \dots, N-1. \quad (2.44)$$

However the mapping from $M_{k,j}$ and $r_{t+k|t}$ to $x_{t+k|t}$ is non-convex, thus yielding a non-convex optimization problem. Then [58], [38] propose an equivalent parameterization of (2.44) that yields a convex problem. In this new parameterization each predicted

control law in Mode 1 is an affine function of the sequence of past disturbances:

$$u_{t|t} = v_{t|t}, \quad u_{t+k|t} = v_{t+k|t} + \sum_{j=0}^{k-1} L_{k,j} w_{t+j}, \quad k = 1, \dots, N-1, \quad (2.45)$$

where $M_{k,j}$, $v_{t+k|t}$ are the degrees of freedom of the predicted control policy. It is also assumed that a fixed state feedback $u = Kx$ is deployed at Mode 2. [88] also propose a predicted control policy of the form of (2.44) but in a Stochastic MPC context, where the non-convexity is solved by using the Youla parameter, as will be shown in Section 2.3. The control policy of (2.45) depends on future unknown disturbances, but as noted earlier this is not problematic since these predicted inputs are never executed; all that is needed is the knowledge that an admissible control sequence that respects causality exists. Only the first input is actually implemented, but that one does not depend on any unknown disturbances.

The approach of [58] considers two cost functions (worst case or quadratic cost, both over a finite prediction horizon) that do not allow one to prove stability, whereas the approach of [38] does allow one to prove a type of stability, and is therefore examined next.

The online procedure consists of optimizing a performance criterion (a quadratic function of the predicted nominal states and inputs plus a quadratic terminal cost) such that the constraints (2.32) are satisfied for every possible realization of the disturbances. The constraints are enforced explicitly in Mode 1, and in Mode 2 these are enforced implicitly by requiring that the terminal state $x_{t+k|t}$ is inside the terminal set $\mathbb{X}_f$ which is robust positively invariant for system (2.43) and constraints (2.32) under $u = Kx$. The constraints in Mode 1 can be expressed as

$$F\mathbf{v}_t + \max_{\mathbf{w}_t \in \mathbb{W}^N} (FL + G)\mathbf{w}_t \leq c + Hx_t, \quad (2.46)$$

where Q, R, P are positive definite weighting matrices. F, G, c are matrices defined

to conveniently express the constraints as a function of $\mathbf{v}_t = [v_{t|t}^T, \dots, v_{t+N-1|t}^T]^T$, $\mathbf{w}_t = [w_t^T, \dots, w_{t+N-1}^T]^T$, and

$$\mathbf{L} = \begin{bmatrix} 0 & \cdots & \cdots & 0 \\ L_{1,0} & 0 & \cdots & 0 \\ \vdots & \ddots & \ddots & \vdots \\ L_{N-1,0} & \cdots & L_{N-1,N-2} & 0 \end{bmatrix}.$$

The max term in (2.46) can be eliminated by using duality as long as $\mathbb{W}$ is compact and convex. For instance, if $\mathbb{W}$ is polytopic, one can write $\mathbb{W}^N = \{\mathbf{w} : S\mathbf{w} \leq h\}$, and then

$$\begin{aligned} \max_{\mathbf{w}_t \in \mathbb{W}^N} (FL + G)_i \mathbf{w}_t &= \min_{z_i} h^T z_i \\ \text{subject to} \quad S^T z_i &= (FL + G)_i, \quad z_i \geq 0, \end{aligned} \quad (2.47)$$

where z_i represents the dual variables associated with the i -th row of the maximization in (2.46). With this, the online optimization of the RMPC strategy is given by

$$\begin{aligned} \underset{\mathbf{L}, \mathbf{v}_k, \mathbf{Z}}{\text{minimize}} \quad & \sum_{k=0}^{N-1} \bar{x}_{t+k|t}^T Q \bar{x}_{t+k|t} + v_{t+k|t}^T R v_{t+k|t} + \bar{x}_{t+N|t}^T P \bar{x}_{t+N|t} \\ \text{subject to} \quad & F\mathbf{v}_t + Z^T h \leq c + Hx_t, \\ & (FL + G) = Z^T S, Z \geq 0, \end{aligned} \quad (2.48)$$

where $\bar{x}_{t+k|t}$ denotes the nominal predictions (i.e. in the absence of disturbances) of the state. This optimization is tractable since the number of variables and constraints of this problem grows quadratically with the prediction horizon. Alternatively, if $\mathbb{W}$ is ellipsoidal, the problem can be solved as a second-order cone program.

Following from the structure of the predicted control policy the receded solution is feasible, and thus the strategy is recursively feasible. With the nominal cost function

of (2.48) and under appropriate conditions on Q, P , one can obtain a decrease condition in absence of disturbances, which allows one to guarantee ISS stability when disturbances are present. Other types of cost, such as expected value costs [37] and $\mathcal{H}_\infty$ [34] have also been used. The former also yields ISS stability, while the latter yields an l_2 performance bound. Further modifications of the cost along with an equivalent re-parameterization of the control policy, allow one to obtain convergence to the minimal invariance set $\Omega_\infty(\mathbb{W})$ [90], convergence in finite time to the terminal set, and, assuming that the disturbance is a stochastic process with a bounded distribution, convergence to $\Omega_\infty(\mathbb{W})$ with probability one [91].

The structure of the predicted control policy of (2.44) and (2.45) used in this strategy is more general than those of [45], [51], since these are subsumed as particular cases of the former.

2.2.4.2 Tube MPC strategies

Another family of strategies that have been successful in dealing with the additive uncertainty in RMPC is the so-called tube MPC (TMPC). The principle is that when uncertainty is present, the states move through a set rather than a single trajectory, and since the system under consideration is linear the predictions can be separated in a nominal portion, which is not affected by the disturbances, and an uncertain one. Predictions then include a nominal trajectory which defines the centre of the tubes and the cross sections of the tubes that define the zones through which the uncertain portion can move. The complexity of the strategies usually depends on the descriptions of the tubes, as will be seen next.

Tubes have been used in MPC for some time now, although the earlier work ([36], [84], [61], [56]) did not make explicit use of the term “tubes”. Among these, [61] devised a strategy for systems described by (2.43),

$$x_{t+1} = Ax_t + Bu_t + w_t,$$

and constraints as in (2.2)

$$u_t \in \mathbb{U}, x_t \in \mathbb{X}.$$

The strategy uses a re-parameterization of the predicted dynamics into a nominal (without disturbances) and an uncertain portion, and considers a control policy for Mode 1 given by the closed-loop paradigm [82], expressed as

$$u_{t+k|t} = \bar{u}_{t+k|t} + K(x_{t+k|t} - \bar{x}_{t+k|t}), \quad k = 0, \dots, N-1, \quad (2.49)$$

where $\bar{x}_{t+k|t}$ and $\bar{u}_{t+k|t}$ are associated to the nominal portions of the predictions. As usual, it is assumed that a state feedback $u = Kx$ will be used in Mode 2. If the uncertain portion of the predicted state is denoted by $e_{t+k|t}$, the dynamics of the uncertain and nominal portions in Mode 1 can be expressed by

$$e_{t+k+1|t} = (A + BK)e_{t+k|t} + w_{t+k}, \quad k = 0, \dots, N-1, \quad (2.50a)$$

$$\bar{x}_{t+k+1|t} = A\bar{x}_{t+k|t} + B\bar{u}_{t+k|t}, \quad k = 0, \dots, N-1, \quad (2.50b)$$

respectively. From this decomposition it is clear that the uncertain portion evolves according to inputs given by the static state feedback gain K , thus the predictions of the uncertain part will always stay in the minimal invariant set under $u = Kx$, $\Omega_\infty(\mathbb{W})$, provided that $e_{t|t} \in \Omega_\infty(\mathbb{W})$. Then, $x_{t+k|t}$ and $u_{t+k|t}$ are contained in sets given by $\bar{x}_{t+k|t} \oplus \Omega_\infty(\mathbb{W})$ and $\bar{u}_{t+k|t} \oplus K\Omega_\infty(\mathbb{W})$, with cross sections given by $\Omega_\infty(\mathbb{W})$ and $K\Omega_\infty(\mathbb{W})$ and centred over $\bar{x}_{t+k|t}$ and $\bar{u}_{t+k|t}$. The tubes are then the sequence of sets that contain the predictions. Subsequently, the uncertain portion is omitted from the optimization, and instead its effect is accounted for by a tightening of the constraints applied to the nominal states and inputs. The online optimization then includes these tightened constraints and the typical objective function for the nominal

predictions with horizon N and a terminal cost:

$$\begin{aligned}
 & \underset{\bar{\mathbf{u}}_t = \{\bar{u}_{t|t}, \dots, \bar{u}_{t+N-1|t}\}}{\text{minimize}} & \bar{J}_N(\bar{x}_t, \bar{\mathbf{u}}_t) &= \sum_{k=0}^{N-1} \ell(\bar{x}_{t+k|t}, \bar{u}_{t+k|t}) + V_f(\bar{x}_{t+N|t}) \\
 & \text{subject to} & \bar{x}_{t+k|t} &\in \bar{\mathbb{X}}, \quad k = 0, \dots, N-1, \\
 & & \bar{u}_{t+k|t} &\in \bar{\mathbb{U}}, \quad k = 0, \dots, N-1, \\
 & & \bar{x}_{t+N|t} &\in \bar{\mathbb{X}}_f
 \end{aligned} \tag{2.51}$$

where $\bar{\mathbb{X}} = \mathbb{X} - \Omega_\infty(\mathbb{W})$, $\bar{\mathbb{U}} = \mathbb{U} - K\Omega_\infty(\mathbb{W})$, $\bar{\mathbb{X}}_f = \mathbb{X}_f - \Omega_\infty(\mathbb{W})$, and $\ell(x, u)$ is convex. The structure of the predictions guarantees satisfaction of the system constraints and that if $\bar{x}_{t|t}$ is an optimization variable, rather than being fixed and equal to x_t , then the receded $\bar{u}_{t+k|t}$ will be feasible and guarantee that the optimal cost is decreasing. Thus the online application of the strategy is recursively feasible, and since optimal cost is made zero when $x_t \in \Omega_\infty(\mathbb{W})$ it can be seen as a Lyapunov function outside $\Omega_\infty(\mathbb{W})$ which along with the decrease property of the optimal cost yields exponential stability and convergence of the state to the set $\Omega_\infty(\mathbb{W})$ [63].

This strategy, since it considers tightened constraints and performs the online optimization only over the nominal sequences, is quite convenient in the sense that it requires only a minimal computational effort because it has the same complexity as the open-loop deterministic formulations. It is however conservative because the tubes have a fixed cross-section and hence they are usually larger than what could be accomplished. Further improvements have been proposed [28], [81], [53] where the tubes cross-sections are smaller and different for each prediction instant, thus making the constraint tightening less conservative.

In order to obtain tighter tubes, other strategies that compute the tubes online have been proposed [54], [74], [75]. For instance [54] suggest tubes $X_{t+k|t}$, $k = 0, \dots, N$ where the centre of the sequences $\bar{x}_{t+k|t}$ can be freely chosen (as opposed to determined by the nominal evolution of the system), the cross sections are polytopes

not necessarily invariant, with arbitrary fixed shapes given by $Z = \text{conv}\{v^1, \dots, v^P\}$.

The sequence α_k allows the sizes of the tubes to vary, such that

$$X_{t+k|t} = \bar{x}_{t+k|t} \oplus \alpha_k Z = \text{conv}\{x_k^1, \dots, x_k^P\}, \quad k = 1, \dots, N, \quad (2.52)$$

where the initial tube is the current state, $X_{t|t} = x_t$. The control tubes are also polytopic, and there is a control u_k^p associated to each x_k^p such that, $U_{t+k|t} = \text{conv}\{u_k^1, \dots, u_k^P\}$, $k = 1, \dots, N-1$ and $U_{t|t} = u_t$, and define a control policy $\pi(x_t) = \{\mu_0(x_{t|t}), \mu_1(x_{t+1|t}), \dots, \mu_{N-1}(x_{t+N-1|t})\}$, where each control law $\mu_k(x)$ defines the predicted input $u_{t+k|t}$ as

$$u_{t+k|t} = \mu_k(x_{t+k|t}) = \sum_{p=1}^P \lambda_k^p(x_{t+k|t}) u_k^p, \quad k = 1, \dots, N-1 \quad (2.53)$$

where $\lambda_k(x) = \{\lambda_k^1, \dots, \lambda_k^P\}$ is the unique least-squares decomposition of $\sum_{p=1}^P \lambda_k^p x_k^p = x$. The online optimization problem includes $\bar{x}_{t+k|t}$, α_k and $U_{t+k|t}$ as variables thus providing more freedom to choose the tubes.

A further improvement is proposed in [73]. It uses a new way of building the tubes that stems from the linear separability of the effects of the disturbances acting at different instants. The tubes are shown to be tighter and strong stability results are obtained. The strategy considers linear, time-invariant, discrete time systems as in (2.43) and constraints as in (2.2) where $\mathbb{X}$, $\mathbb{U}$ and $\mathbb{W}$ are convex polytopic sets that contain the origin in their interiors and are defined as

$$\mathbb{X} = \{x : Fx \leq \underline{1}\}, \quad \mathbb{U} = \{u : Gu \leq \underline{1}\}, \quad \mathbb{W} = \text{conv}\{\tilde{w}_i : i = 1, \dots, q\}, \quad (2.54)$$

where $F \in \mathbb{R}^{n_{cx} \times n}$, $G \in \mathbb{R}^{n_{cu} \times n_u}$. A separable prediction scheme for Mode 1 is assumed where the nominal dynamics and the contribution of each future disturbance are separated into different sequences. $\mathbf{x}_{(0,N)} = \{x_{(0,0)}, \dots, x_{(0,N)}\}$ and $\mathbf{u}_{(0,N-1)} =$

$\{u_{(0,0)}, \dots, u_{(0,N-1)}\}$ are the 0^{th} partial state and control sequences, which account for the nominal dynamics

$$x_{(0,k+1)} = Ax_{(0,k)} + Bu_{(0,k)}, \quad k = 0, \dots, N-1, \quad (2.55)$$

where $x_{(0,0)} = x_t$ and $\mathbf{x}_{(j,N)} = \{x_{(j,j)}, x_{(j,j+1)}, \dots, x_{(j,N)}\}$, $\mathbf{u}_{(j,N-1)} = \{u_{(j,j)}, u_{(j,j+1)}, \dots, u_{(j,N-1)}\}$ are the j^{th} partial state and control sequences describing the dynamical contribution of the disturbance acting at the $(j-1)^{th}$ prediction time, $w_{(j-1)}$. These dynamics are given by

$$\begin{aligned} x_{(j,j)} &= w_{(j-1)}, \\ x_{(j,k+1)} &= Ax_{(j,k)} + Bu_{(j,k)}, \quad k = j, \dots, N-1. \end{aligned} \quad (2.56)$$

For simplicity the notation of the type $x_{t+k|t}$ is not used for the k -step-ahead predictions made at time t , and instead $x_{(0,k)}, u_{(0,k)}$ and $x_{(j,k)}, u_{(j,k)}$ are the k -step-ahead predictions made at each time t associated with the 0^{th} partial sequences and the j^{th} partial sequences, while $x_{(k)}, u_{(k)}$ are the full k -step-ahead predictions made at each time t . Accordingly, the full predictions for $k = 0, \dots, N$ are given by

$$x_{(k)} = \sum_{j=0}^k x_{(j,k)}, \quad k = 0, \dots, N, \quad u_{(k)} = \sum_{j=0}^k u_{(j,k)}, \quad k = 0, \dots, N-1, \quad (2.57)$$

with

$$\begin{aligned} x_{(j,k)} \in X_{(j,k)} &= \begin{cases} x_{(0,k)} & \text{if } k = 0 \\ \text{conv}\{x_{(i,j,k)} : i = 1, \dots, q\} & \text{if } k = 1, \dots, N \end{cases} \\ u_{(j,k)} \in U_{(j,k)} &= \begin{cases} u_{(0,k)} & \text{if } k = 0 \\ \text{conv}\{u_{(i,j,k)} : i = 1, \dots, q\} & \text{if } k = 1, \dots, N-1. \end{cases} \end{aligned} \quad (2.58)$$

thus defining a triangular prediction structure. From (2.58), the j^{th} partial state and control sequences are defined by the j^{th} extreme partial state and control sequences, $\mathbf{x}_{(i,j,N)} = \{x_{(i,j,j)}, x_{(i,j,j+1)}, \dots, x_{(i,j,N)}\}$ and $\mathbf{u}_{(i,j,N-1)} = \{u_{(i,j,j)}, u_{(i,j,j+1)}, \dots, u_{(i,j,N-1)}\}$, respectively, which describe the effect of the extreme disturbances acting at the prediction instant $j - 1$. The dynamics of the j^{th} partial extreme state and control sequences are, for $i = 1, \dots, q$,

$$\begin{aligned} x_{(i,j,j)} &= \tilde{w}_i, & j &= 1, \dots, N, \\ x_{(i,j,k+1)} &= Ax_{(i,j,k)} + Bu_{(i,j,k)}, & j &= 1, \dots, N-1, \quad k = j, \dots, N-1. \end{aligned} \quad (2.59)$$

The partial sequences as defined by (2.55) and (2.59) imply that the predicted states and inputs will be contained in the tubes with cross-sections

$$X_{(k)} = \bigoplus_{j=0}^k X_{(j,k)}, \quad k = 0, \dots, N, \quad U_{(k)} = \bigoplus_{j=0}^k U_{(j,k)}, \quad k = 0, \dots, N-1, \quad (2.60)$$

where the partial tubes $X_{(j,N)} = \{X_{(j,j)}, X_{(j,j+1)}, \dots, X_{(j,N)}\}$ and $U_{(j,N-1)} = \{U_{(j,j)}, U_{(j,j+1)}, \dots, U_{(j,N-1)}\}$ contain the partial sequences $\mathbf{x}_{(j,N)}$ and $\mathbf{u}_{(j,N-1)}$, such that $X_{(0,k)} = x_{(0,k)}$, $U_{(0,k)} = u_{(0,k)}$, $X_{(j,k)} = \text{conv}\{x_{(i,j,k)} : i = 1, \dots, q\}$, and $U_{(j,k)} = \text{conv}\{u_{(i,j,k)} : i = 1, \dots, q\}$.

The separable prediction scheme induces a separable non-linear predicted control policy $\pi_{N-1} = \{\pi_0(\cdot), \dots, \pi_{N-1}(\cdot)\}$ given by,

$$\begin{aligned} \pi_k(x_{(k)}) &= \sum_{j=0}^k \pi_{(j,k)}(x_{(j,k)}), \quad \text{with } x_{(k)} = \sum_{j=0}^k x_{(j,k)} \\ \pi_{(0,k)}(x_{(0,k)}) &= u_{(0,k)}, \\ \pi_{(j,k)}(x_{(j,k)}) &= \sum_{i=1}^q \lambda_{(i,j)}(w_{(j-1)}) u_{(i,j,k)}, \quad j = 1, \dots, N-1, \quad k = j, \dots, N-1 \end{aligned} \quad (2.61)$$

where $\lambda_{(i,j)}(w_{(j-1)})$ are convex interpolation parameters that satisfy $\lambda_{(i,j)}(w_{(j-1)}) \geq 0$

for $i = 1, \dots, q$, $\sum_{i=1}^q \lambda_{(i,j)}(w_{(j-1)}) = 1$ and $\sum_{i=1}^q \lambda_{(i,j)}(w_{(j-1)}) \tilde{w}_i = w_{(j-1)}$. A well-defined and unique selection of such convex multipliers $\{\lambda_{(i,j)}, i = 1, \dots, q\}$ can be obtained employing, for instance, a least squares procedure.

This control policy is non-linear and is proved to be more general than a linear policy depending on the past states (or its equivalent formulation as a policy that is affine-in-the-disturbances). The control policy of (2.61) is more general than that of (2.33) (closed-loop paradigm [82], [84]) and (2.45) (affine-in-the-disturbances [38], [58]), which are subsumed as special cases of PTMPC.

Remark 2.11. *Since the disturbances can be expressed as $w_t = \sum_{i=1}^q \lambda_t \tilde{w}_i$, such that λ_t is contained in the simplex $\Lambda = \{\lambda \in \mathbb{R}^q | \lambda_i \geq 0, i = 1, \dots, q, \sum_{i=1}^q \lambda_i = 1\} \subset \mathbb{R}^q$, system (2.43) can be re-expressed as*

$$x_{t+1} = Ax_t + Bu_t + V\lambda, \quad (2.62)$$

where $V = [\tilde{w}_1, \dots, \tilde{w}_q]$, $\lambda = [\lambda_1, \dots, \lambda_q]^T \in \Lambda$ is a disturbance in the lifted space $\mathbb{R}^q$. Thus the terms $\lambda_{(i,j)}(w_{(j-1)})$ that appear in the predicted policy of (2.61) are actually the disturbances of system (2.62), so then (2.61) can be re-expressed as

$$\begin{aligned} \pi_k(x_{(k)}) &= \sum_{j=0}^k \pi_{(j,k)}(x_{(j,k)}), \quad \text{with } x_{(k)} = \sum_{j=0}^k x_{(j,k)}, \\ \pi_0(x_{(0)}) &= u_{(0,0)}, \\ \pi_k(x_{(k)}) &= u_{(0,k)} + \sum_{j=1}^k u_{(1 \rightarrow q,j,k)} \lambda_{(j)}, \quad k = 1, \dots, N-1, \end{aligned} \quad (2.63)$$

where $\lambda_{(j)} = [\lambda_{(1,j)}, \dots, \lambda_{(q,j)}]^T$ is the disturbance in the lifted space associated with $w_{(j-1)}$ and $u_{(1 \rightarrow q,j,k)} = [u_{(1,j,k)}, \dots, u_{(q,j,k)}]$. This is clearly an affine-in-the-disturbances policy as that of (2.45), where $v_{t+k|t} = u_{(0,k)}$, $k = 0, \dots, N-1$ and $L_{k,j} = u_{(1 \rightarrow q,j+1,k+1)}$, $k = 1, \dots, N-1$, $j = 0, \dots, k-1$, for the equivalent system of (2.62), where the

disturbances $\lambda_t \in \Lambda$ are in the lifted space $\mathbb{R}^q$.

Remark 2.12. It is said in Remark 2.11 that the disturbances $\lambda \in \Lambda$, $\Lambda \subset \mathbb{R}^q$, are in a lifted space because for a bounded polytope $\mathbb{W} \subset \mathbb{R}^n$ with non-empty interior, the smallest number of vertices is $q = n + 1$, while it can be larger in general (e.g. 2^n for hypercubes). Thus Λ is in a higher dimensional space than $\mathbb{W}$.

The conditions for constraint satisfaction are given by

$$\begin{aligned} X_{(k)} &\subseteq \mathbb{X}, \quad k = 0, \dots, N-1, \\ U_{(k)} &\subseteq \mathbb{U}, \quad k = 0, \dots, N-1, \\ X_{(N)} &\subseteq \mathbb{X}_f, \end{aligned} \tag{2.64}$$

where $\mathbb{X}_f = \{Hx \leq \underline{1}\}$, $H \in \mathbb{R}^{n_f \times n}$, is the maximal robust positively invariant set under the terminal control law $u = Kx$. These conditions can be guaranteed equivalently by constraints on the nominal sequences with slack variables that account for the worst (extreme) realizations of the uncertainty.

$$\begin{aligned} F_l x_{(0,0)} &\leq 1, \quad l = 1, \dots, n_{cx}, \quad G_l u_{(0,0)} \leq 1, \quad l = 1, \dots, n_{cu}, \\ F_l x_{(0,k)} + \sum_{j=1}^k f_{(l,j,k)} &\leq 1, \quad l = 1, \dots, n_{cx}, \quad k = 1, \dots, N-1 \\ G_l u_{(0,k)} + \sum_{j=1}^k g_{(l,j,k)} &\leq 1, \quad l = 1, \dots, n_{cu}, \quad k = 1, \dots, N-1 \\ H_l x_{(0,N)} + \sum_{j=1}^N h_{(l,j,N)} &\leq 1, \quad l = 1, \dots, n_f, \\ F_l x_{(i,j,k)} &\leq f_{(l,j,k)}, \quad l = 1, \dots, n_{cx}, \quad j = 1, \dots, N-1, \quad k = j, \dots, N-1, \quad i = 1, \dots, q, \\ G_l u_{(i,j,k)} &\leq g_{(l,j,k)}, \quad l = 1, \dots, n_{cu}, \quad j = 1, \dots, N-1, \quad k = j, \dots, N-1, \quad i = 1, \dots, q, \\ H_l x_{(i,j,N)} &\leq h_{(l,j,N)}, \quad l = 1, \dots, n_f, \quad j = 1, \dots, N, \quad i = 1, \dots, q, \end{aligned} \tag{2.65}$$

where F_l , G_l and H_l are the l^{th} rows of F , G , and H , respectively.

In order to evaluate the cost, a re-parameterization of variables is performed such that $x_{(j,k)} = z_{(j,k)} + e_{(j,k)}$ and $u_{(j,k)} = v_{(j,k)} + \psi_{(j,k)}$, where $e_{(j,k)} = (A + BK)^{k-j} e_{(j,j)}$ and $\psi_{(j,k)} = K e_{(j,k)}$ are associated to the evolution of the system according with the terminal control law $u = Kx$, and $z_{(j,k)}$ and $v_{(j,k)}$ are the required deviation, in order to satisfy constraints, from $e_{(j,k)}$ and $\psi_{(j,k)}$. The cost function contains penalizations of the terms $z_{(j,k)}$ and $v_{(j,k)}$, and is structured as the sum of the partial costs associated to each partial sequence, which in turn is structured as the sum of the costs associated with each extreme partial sequence, such that, if $\mathbf{d}_N$ contains all the decision variables of the formulation,

$$J(\mathbf{d}_N) = \sum_{j=0}^N J_{(j,N)}(\mathbf{d}_N), \quad (2.66)$$

where

$$\begin{aligned} J_{(0,N)}(\mathbf{d}_N) &= \sum_{k=0}^{N-1} \ell(z_{(0,k)}, v_{(0,k)}) + V_f(z_{(0,N)}), \\ J_{(j,N)}(\mathbf{d}_N) &= \sum_{i=1}^q J_{(i,j,N)}(\mathbf{d}_N), \quad j = 1, \dots, N-1, \quad J_{(N,N)}(\mathbf{d}_N) = \sum_{i=1}^q V_f(z_{i,N,N}), \\ J_{(i,j,N)}(\mathbf{d}_N) &= \sum_{k=j}^{N-1} \ell(z_{(i,j,k)}, v_{(i,j,k)}) + V_f(z_{(i,j,N)}), \quad i = 1, \dots, q, \quad j = 1, \dots, N-1 \\ \ell(z, v) &= |z|_{\mathcal{Q}} + |v|_{\mathcal{R}}, \quad V_f(z) = |z|_{\mathcal{P}}, \end{aligned}$$

$\mathcal{Q}, \mathcal{P}, \mathcal{R}$ are symmetric polytopic sets that contain the origin in their interiors and $V_f(\cdot)$ satisfies the Lyapunov like condition $V_f((A + BK)z) - V_f(x) \leq -\ell(z, Kz)$ for all $x \in \mathbb{R}^n$.

On account of the penalizations of the terms $z_{(j,k)}$ and $v_{(j,k)}$, the value function $J^*(x)$ will be zero if the entire evolution can be satisfied by $u = Kx$ and greater than zero otherwise. It follows that $J^*(x) > 0$ if $x \notin \mathbb{X}_f$ and that $J^*(x) = 0$ if $x \in \mathbb{X}_f$, in which case the control reduces to $u = Kx$. Since, additionally, the

tubes structure guarantees recursive feasibility, and the decrease property $J^*(x_{t+1}) - J(x_t) \leq -(|x_t|_{\mathcal{Q}} + |u_t|_{\mathcal{R}})$ can be obtained from the online application where the input is assigned to be $u_t = u_{(0,0)}^*$, it follows that $J^*(x)$ is a Lyapunov function outside of $\mathbb{X}_f$, thus proving exponential convergence of the state to $\mathbb{X}_f$. Furthermore, since in this set the control law reduces to $u = Kx$, the state converges exponentially to the minimal robust positively invariant set $\Omega_\infty(\mathbb{W})$.

PTMPC is proved to be equivalent to DP whenever $n = n_u = 1$ for any prediction horizon N , or for any n, n_u when $N = 1, 2$ [73], while the affine in the disturbances strategy [58], [38] is only equivalent to DP in the first case [73], [10].

The optimization using the cost from (2.66) can be cast as a linear program, where the triangular tubes structure means that the number of variables and constraints grows quadratically with the prediction horizon N . However, in order to account for the worst cases in the constraints and the cost one needs to consider sequences associated with the q vertices of $\mathbb{W}$ (more precisely, the size of the problem grows linearly with the product nq). Since q can grow exponentially with n (e.g. if $\mathbb{W}$ is a hypercube) the size of the problem can also grow exponentially with n . On the other hand, in the affine-in-the-disturbances strategies [38], [58] one can handle the worst case scenarios with a dual problem as shown in (2.47), which uses a description of disturbances of the type $\mathbb{W} = \{w : Sw \leq h\}$. Then the size of the optimization problem grows linearly with the product nn_s , where n_s is the number of rows of S . Generally n_s grows more slowly than the number of vertices of $\mathbb{W}$ as n increases (e.q. linearly if $\mathbb{W}$ is a hypercube), so then the affine-in-the-disturbances strategies are better suited for large n .

2.3 Stochastic Model Predictive Control

RMPC techniques only take into account the hard bounds of the uncertainties, and not their probabilistic distribution, which is problematic when soft probabilistic con-

straints are present. Additionally, the optimization problems in RMPC, where the costs and constraints are based on worst cases, can be extremely conservative because the uncertainties near the bounds can be extremely unlikely.

For these reasons, stochastic MPC (SMPC) has emerged as a family of techniques that incorporates the stochastic descriptions of the uncertainties for handling probabilistic constraints and probabilistic measures of performance.

This section first introduces the main aspects of SMPC for systems with additive uncertainty. This is followed by a review of the advances in SMPC that are most relevant for the development of Chapter 6.

2.3.1 Main aspects of Stochastic MPC

In order to introduce the main aspects of SMPC consider linear, discrete-time systems

$$x_{t+1} = Ax_t + Bu_t + Dw_t, \quad (2.67)$$

where the additive uncertainty w_t is an exogenous disturbance with unknown current and future values (as in RMPC) but such that some type of description of its stochastic nature is known, e.g. probabilistic distribution.

The main aspects to consider in SMPC are the stochastic characterization of the uncertainty, probabilistic constraints, and cost functions (as well as typical aspects such as recursive feasibility and stability).

Characterization of the uncertainty

The additive uncertainty has been characterized in several ways. The most common characterization (e.g. [4], [6], [87],[88], [23], [22], [20], [24]) considers w_t as an independent and identically distributed sequence (i.i.d., meaning that w_t is independent from w_k for all $t \neq k$, and that they have the same distribution), where each w_t has a normal distribution with mean 0 and covariance matrix Σ .

This characterization has a big weakness: it does not allow one to obtain recur-

sive feasibility of the online optimization problem while guaranteeing satisfaction of constraints due to the possibility of *infinite* size disturbances. Additionally, when recursive feasibility of the online optimization is not guaranteed, stability results do not hold. In practice, however, disturbances are usually bounded. Using a bounded characterization of the additive uncertainty, namely that of i.i.d. sequences where w_t has a bounded domain distribution, enables the possibility of obtaining recursive feasibility of the control strategy [49], [21], [50], [32].

Constraints

The most common type of constraints used in SMPC is that of probabilistic constraints, in which the probability of the satisfaction of some condition has to be greater than a predefined value. Examples of such constraints are:

- Individual probabilistic constraints

$$\Pr(y_t \leq 1) \geq p_t, \quad t = 0, 1, 2, \dots \quad (2.68)$$

- Joint probabilistic constraints

$$\Pr(y_t \leq 1, t = 0, 1, \dots, N) \geq p, \quad (2.69)$$

where $y_t = g(x_t, u_t, w_t) \in \mathbb{R}^{n_y}$ is an output function and $p_t, p \in (0, 1]$. Outputs with $n_y > 1$ can be considered, in which case the probability of the satisfaction of all constraints has to be greater than p . Alternatively, several constraints for several outputs with $n_y = 1$ can be considered such that the probability of the satisfaction of each constraint on its own has to be greater than p . Appropriate choices of $g(x, u, w)$ allow one to set up state, input and mixed state-input constraints. Both cases, additionally, cover the robust case by setting $p = 1$.

The individual probabilistic constraints [85], [92], [47], [49], [21] indicate that the

condition for each prediction time has to be satisfied with a probability greater or equal to the predefined value, while in the joint probabilistic constraints [87], [88], [89], [68], [16] the joint condition on all the predictions has to be satisfied with a probability greater or equal to the predefined value. If $p = p_t$, for all $t = 0, \dots, N$, then the joint probability constraints are clearly more stringent.

These constraints can be handled in different manners, as will be seen in the next section. Most strategies attempt to use the distributions of the uncertainty explicitly to find appropriate tightening parameters to be used in equivalent constraints for the nominal sequences [88], [89], [23], [21], [49]. A more recent trend however, consists of using a set of samples of realizations of the uncertainty on which to apply the original constraints, where among other conditions, the size of the sample is the main factor to ensure satisfaction of the probabilistic constraints [1], [14], [8], [16], [15] with a given confidence.

Although probabilistic constraints have been the most commonly used type, constraints based on expected values of relevant variables (or a function of them) have also been considered in SMPC [69], [70].

Recursive feasibility in SMPC is a bit harder to deal with than in RMPC. As it has been mentioned before, formulations that consider unbounded disturbances allow for the presence of *infinite* size states, and then feasibility of the probabilistic constraints for the predicted sequences at every sample time cannot be guaranteed. Even for the case of bounded uncertainty, satisfaction of the probabilistic constraints for the predicted sequences does not guarantee in general that at the next sample time these constraints will be satisfied, as usually happens with RMPC formulations. Therefore the constraints have to be tightened in a robustification process so as to guarantee recursive feasibility [49], [21]. This will be shown in more detail in the following section.

Cost function and stability

In order to account for the stochastic nature of the system, the control problem is to minimize, at each time k , the expected quadratic cost [4]

$$J(x_t, \mathbf{u}_t) = \sum_{k=0}^{\infty} \mathbb{E}_t (x_{t+k|t}^T Q x_{t+k|t} + u_{t+k|t}^T R u_{t+k|t}), \quad (2.70)$$

where $Q \succeq 0, R \succ 0$, and $\mathbb{E}_t$ is the expected value conditional on the known value of x_t , or $\mathbb{E}(\cdot|x_t)$.

A common stability notion used in SMPC (when multiplicative uncertainty, but no additive uncertainty, is considered) [22], [29], [22] is that of mean-square stability [52] (A system is said to be mean-square stable if $\lim_{t \rightarrow \infty} \mathbb{E}(\|x_t\|^2) = 0$). A direct consequence when a system is mean-square stable is that the performance of (2.70) is upper bounded, and then it can be concluded that $x_t \rightarrow 0$ with probability 1.

Under the presence of additive uncertainty however, the stage cost (2.70) does not converge to zero, and then the cost is infinite, thus precluding the possibility of using the notion of mean-square stability. This problem has been effectively fixed by modifying the stage cost such that the new overall cost is finite [23]

$$J(x_t, \mathbf{u}_t) = \sum_{k=0}^{\infty} \mathbb{E}_t (x_{t+k|t}^T Q x_{t+k|t} + u_{t+k|t}^T R u_{t+k|t} - l_{ss}), \quad (2.71)$$

$$l_{ss} = \lim_{k \rightarrow \infty} \mathbb{E}_t (x_{t+k|t}^T Q x_{t+k|t} + u_{t+k|t}^T R u_{t+k|t}).$$

Note that under stable prediction dynamics l_{ss} is a constant that does not depend on t . This modification allows one to obtain an analogous quadratic stability condition.

The repeated solution of the online optimization problems using the cost of (2.71) and implementing the input $u_t = u_{t|t}^*$, gives the decrease condition

$$J(x_{t+1}, \mathbf{u}_{t+1}^*) - J(x_t, \mathbf{u}_t^*) \leq -\mathbb{E}_t (x_t^T Q x_t + u_t^T R u_t - l_{ss}), \quad (2.72)$$

which, summing for $t = 0, \dots, \infty$ gives

$$\lim_{n \rightarrow \infty} \frac{1}{n} (J(x_0, \mathbf{u}_0^*) - J(x_n, \mathbf{u}_n^*)) \geq \lim_{n \rightarrow \infty} \frac{1}{n} \sum_{t=0}^{n-1} \mathbb{E}_0(x_t^T Q x_t + u_t^T R u_t) - l_{ss}.$$

This, used in conjunction with the fact that $J(x, \mathbf{u})$ is lower bounded [23], yields the quadratic stability condition

$$\lim_{n \rightarrow \infty} \frac{1}{n} \sum_{t=0}^{n-1} \mathbb{E}_0(x_t^T Q x_t + u_t^T R u_t) \leq l_{ss}, \quad (2.73)$$

which proves that $\lim_{t \rightarrow \infty} \mathbb{E}_0(x_t^T Q x_t + u_t^T R u_t) = l_{ss}$.

2.3.2 Review of Stochastic MPC

2.3.2.1 Probabilistic constraints via ellipsoidal bounding

A relevant early formulation for SMPC is presented in [88], [89]. The plant model is described in terms of the predicted relevant variables over a prediction horizon:

$$\begin{aligned} \mathbf{x}_t &= G_{xx}x_t + G_{xu}\mathbf{u}_t + G_{xw}\mathbf{w}_t, \\ \mathbf{y}_t &= G_{yx}x_t + G_{yu}\mathbf{u}_t + G_{yw}\mathbf{w}_t, \\ \mathbf{z}_t &= G_{zx}x_t + G_{zu}\mathbf{u}_t + G_{zw}\mathbf{w}_t, \end{aligned} \quad (2.74)$$

where the bold characters $\mathbf{x}_t$, $\mathbf{u}_t$, $\mathbf{y}_t$, $\mathbf{z}_t$ are the vectors with the stacked predictions over a prediction horizon of length N of the state, input, measured outputs and performance outputs, and the matrices G_{ab} denote the system transfer matrices from b to a . The disturbance is an i.i.d. sequence, where each w_t has a normal distribution with zero mean and variance equal to 1.

The online problem seeks to find reference (or nominal) inputs $\mathbf{u}_t^r$, that along with a known reference initial state x_t^r generate reference state and outputs $\mathbf{x}_t^r$, $\mathbf{y}_t^r$, $\mathbf{z}_t^r$, which

satisfy (2.74) with $\mathbf{w}_t = 0$

$$\mathbf{x}_t^r = G_{xx}x_t^r + G_{xu}\mathbf{u}_t^r, \quad \mathbf{y}_t^r = G_{yx}x_t^r + G_{yu}\mathbf{u}_t^r, \quad \mathbf{z}_t^r = G_{zx}x_t^r + G_{zu}\mathbf{u}_t^r. \quad (2.75)$$

The control policy is a sequence of feedback functions where the reference inputs are corrected based on the difference of the actual and the reference outputs such that:

$$\mathbf{u}_t = \mathbf{u}_t^r + T(\mathbf{y}_t - \mathbf{y}_t^r) \quad (2.76)$$

where T is a matrix with appropriate dimensions and a lower-triangular structure that ensures causality; it is noted that if $y_t = x_t$ this control policy is the same as in (2.44). However, the non-linear term $T(I - G_{yu}T)^{-1}$ appears when expressing this control policy and the predictions of (2.74) as a function of the degrees of freedom of the formulation, $\mathbf{u}^r$ and T , and the known actual and reference initial states x_t and x_t^r . This causes the optimization problem to be non-linear and non-convex. However, this non-linearity can be removed through the use of the Youla parameter, $Q = T(I - G_{yu}T)^{-1}$ which converts the control policy and the predictions of (2.74) into a linear function of $\mathbf{u}_t^r$ and Q . The control policy now can be expressed as

$$\mathbf{u}_t = QG_{yx}(x_t - x_t^r) + \mathbf{u}_t^r + QG_{yw}\mathbf{w}_t. \quad (2.77)$$

It can be proved that QG_{yw} is also lower-triangular, thus ensuring causality. It follows that, as shown in Section 2.2.4, this control policy is affine-in-the-disturbances, and is equal to that of (2.45) when $y_t = x_t$ with $\mathbf{v}_k = QG_{yx}(x_t - x_t^r) + \mathbf{u}_t^r$ and $QG_{yw} = \mathbf{L}$.

The system constraints are defined as joint probabilistic constraints (over a prediction horizon) for the performance outputs $\mathbf{z}_k$, which have to be contained in the polytope $\mathcal{P} = \{\mathbf{z} : F\mathbf{z} \leq h\}$, $F \in \mathbb{R}^{n_c \times n_z}$, with a probability larger than some

user-chosen level α :

$$\mathbb{P}\{\mathbf{z}_t \in \mathcal{P}\} \geq \alpha. \quad (2.78)$$

In order to ensure this constraint, the following ellipsoid is defined

$$\mathcal{E}_r = \{\mathbf{z} : \mathbf{z}^T Z^{-1} \mathbf{z} \leq r^2\} \quad (2.79)$$

where Z is the covariance matrix of $\mathbf{z}_t$, and the scaling r is chosen to be the smallest value such that the satisfaction of $\hat{\mathbf{z}}_t + \mathcal{E}_r \subseteq \mathcal{P}$, where $\hat{\mathbf{z}}_t = \mathbb{E}(\mathbf{z}_t)$, guarantees the satisfaction of $\mathbb{P}\{\mathbf{z} \in \mathcal{P}\} \geq \alpha$. The condition $\hat{\mathbf{z}}_t + \mathcal{E}_r \subseteq \mathcal{P}$ defines a tube that contains the entire vector $\mathbf{z}_t$, where the cross-section associated with each prediction instant is given by the projection of $\hat{\mathbf{z}}_t + \mathcal{E}_r \subseteq \mathcal{P}$ to the relevant coordinates. The condition $\hat{\mathbf{z}}_t + \varepsilon_r \subseteq \mathcal{P}$ can be imposed by:

$$F_l \mathbf{z}_t = F_l \hat{\mathbf{z}}_t + r \sqrt{F_l Z F_l^T} \leq g_l, \quad l = 1, \dots, n_c, \quad (2.80)$$

where F_l denotes the l -th row of F , and the covariance matrix Z can be expressed as a function of the Youla parameter Q , such that $Z = S(Q)S(Q)^T$ where $S(Q)$ is a linear expression of Q . Additionally, the expected value of the predictions are related with the predicted references by

$$\hat{\mathbf{x}}_t = G_{xx} x_t^r + [G_{xx} + G_{yu} Q G_{yx}](\hat{x}_t - x_t^r) + G_{xu} \mathbf{u}_t^r, \quad \hat{\mathbf{u}}_t = \mathbf{u}_t^r + Q G_{yx}(\hat{x}_t - x_t^r), \quad (2.81)$$

where $\hat{x}_t, x_t^r$ are given values. In the absence of historical data, these are simply related by $\hat{x}_t = x_t^r$, and then $\hat{\mathbf{x}}_t = \mathbf{x}_t^r$ and $\hat{\mathbf{u}}_t = \mathbf{u}_t^r$.

Finally, the optimization problem consists of the minimization over $Q, \mathbf{u}_t^r$ of a function of $\mathbf{z}^r, f(\mathbf{z}^r)$, subject to the reference dynamics of (2.75), to conditions (2.80), (2.81) and $Z = S(Q)S(Q)^T$. Under appropriate conditions the optimization problem can be expressed as a second-order cone program that can be solved efficiently.

The formulation discussed here is introduced by [88], while [89] extends this work to include bounded disturbances in addition to the stochastic disturbances. These works have merit in that they propose a strategy for the handling of stochastic constraints based on the tightening of the constraints. However, the handling of the constraints is conservative because the entire tube is given by a single ellipsoid (as defined by matrix Z) which does not give enough freedom to each cross-section. Additionally, it is only a finite horizon formulation and therefore cannot guarantee recursive feasibility and stability of the closed-loop application.

2.3.2.2 Tight constraints based on explicit use of probabilistic distributions

The unnecessary conservativeness of the strategy of [88], [89] is overcome by [49], wherein the explicit distributions of the disturbances and their effect on the state evolution are used to build tightening parameters that are as small as possible (given the control policy and description of uncertainty) to guarantee the satisfaction of the probabilistic constraints and recursive feasibility. Consideration is given to linear systems with stochastic additive disturbances as in (2.67)

$$x_{t+1} = Ax_t + Bu_t + Dw_t,$$

which are subject to individual probabilistic constraints, so for each instant t the state is subject to

$$\mathbb{P}\{fx_t \leq h\} \geq p, \tag{2.82}$$

where $p \in (0, 1]$, $f \in \mathbb{R}^{1 \times n}$. Condition (2.82) is presented as a single state constraint for simplicity, but the following analysis can be easily generalized for multiple constraints. The disturbance w_t is an i.i.d. sequence, the components of w_t , $w_{t,i}$, have zero mean, are independent, and have arbitrary but known distributions over a finite support.

The predicted control policy is given by the closed-loop paradigm [82], so that the state and input predictions are given by:

$$x_{t+k|t} = z_{t+k|t} + e_{t+k|t}, \quad (2.83a)$$

$$u_{t+k|t} = Kx_{t+k|t} + c_{t+k|t}, \quad (2.83b)$$

$$z_{t+k+1|t} = \Phi z_{t+k|t} + Bc_{t+k|t}, \quad (2.83c)$$

$$e_{t+k+1|t} = \Phi e_{t+k|t} + Dw_{t+k}, \quad (2.83d)$$

where $\Phi = (A + BK)$ is strictly stable. Here $x_{t+k|t}$, $u_{t+k|t}$, $z_{t+k|t}$, $e_{t+k|t}$ are the predicted state, input, nominal state and uncertainty of the state, with $e_{t|t} = 0$ and $z_{t|t} = x_{t|t} = x_t$. Also $c_{t+k|t}$, $k = 0, \dots, N-1$ are the degrees of freedom in the inputs and in the online optimization, and $c_{t+k|t} = 0$, $k \geq N$, for some Mode 1 horizon N .

It is proved in [49] that for a vector of degrees of freedom $\mathbf{c}_t = [c_{t|t}^T, \dots, c_{t+N-1|t}^T]^T$, the predictions of system (2.67) made at time t satisfy the probabilistic constraints (2.82) if and only if:

$$g^T H_k \mathbf{c}_t + g^T \Phi^k x_t + \gamma_k \leq h, \quad k = 1, 2, \dots \quad (2.84)$$

is satisfied, where $H_k = [\Phi^{k-1}B, \dots, B, 0, \dots, 0]$ and $g^T H_k \mathbf{c}_t + g^T \Phi^k x_t = g^T z_{t+k|t}$ are the nominal components. γ_k are the constraint tightening parameters associated with the stochastic components in the constraint term, $g^T e_{t+k|t} = g^T (\Phi^{k-1}Dw_t + \dots + Dw_{t+k-1})$, and are defined as the minimum value that satisfies:

$$\mathbb{P}\{g^T (\Phi^{k-1}Dw_t + \dots + Dw_{t+k-1}) \leq \gamma_k\} = p, \quad k = 1, 2, 3, \dots \quad (2.85)$$

This result only guarantees the feasibility of the probabilistic constraints for the predicted sequence, but does not ensure that these probabilistic constraints will be satisfied in the future. Thus it fails to guarantee recursive feasibility. In order to guar-

antee the recursive feasibility of a probabilistic constraint for the prediction instant $t + k$ computed at time t , i.e. for the prediction instant $t + k|t$, it must be ensured that the predictions $(\cdot)_{t+k|t+1}, (\cdot)_{t+k|t+2}, \dots, (\cdot)_{t+k|t+k-1}$, $k = 2, 3, \dots$, also satisfy the probabilistic constraint. Since the predictions computed at times $t + 1, \dots, t + k - 1$ depend on the states at those instants and in turn these depend on the future disturbances $w_t, \dots, w_{t+k-2}$, one must account for the worst cases associated with these future disturbances in order to check feasibility for the predictions $(\cdot)_{t+k|t+1}, (\cdot)_{t+k|t+2}, \dots, (\cdot)_{t+k|t+k-1}$, $k = 2, 3, \dots$. This means that in the constraints for predictions at $t + k|t$, the effects of $w_t, \dots, w_{t+k-2}$ must be robustified, while the effect of w_{t+k-1} is still accounted for in a probabilistic manner.

Along these lines, in order to satisfy the probabilistic constraints and the recursive feasibility, (2.85) must be replaced by:

$$g^T H_k \mathbf{c}_t + g^T \Phi^k x_t + \leq h - (\beta_k + \gamma_1), \quad k = 1, 2, \dots \quad (2.86)$$

where β_k accounts for the worst case of the effect of $w_t, \dots, w_{t+k-2}$, while γ still accounts for the probabilistic portion of the constraints and is associated with w_{t+k-1} . Under these considerations, noting that the stochastic components in the constraint term at time $t + k$ are given by $f e_{t+k|t} = f(\Phi^{k-1} D w_t + \dots + D w_{t+k-1})$ and that the worst case of $f \Phi^k D w_t$ is $a_k = |g^T \Phi^k D| \alpha$, β_k is given by

$$\beta_1 = 0, \quad \beta_k = a_1 + \dots + a_{k-1}, \quad k = 2, 3, \dots \quad (2.87)$$

These bounds, to provide a real recursive effect, should be applied for every $k = 0, 1, \dots$, but in practice this cannot be done because it would require an infinite number of constraints. Using the arguments of [35] it can be proved that there exists a $\hat{N} \geq 0$ such that (2.86) needs to be invoked only for $k = N, \dots, N + \hat{N} - 1$. This defines a polytopic set $\mathbb{Z}_{\hat{N}}$, that allows one to replace the infinite number of

constraints in (2.86) by $z_{t+N|t} \in \mathbb{Z}_{\hat{N}}$.

The usual objective function used in MPC as in (2.70) is unbounded due to the presence of the additive uncertainty, which causes the limit of the stage cost to be non-zero. Then, as proposed in [23], the modified cost of (2.71) is used. It is proved that this objective function can be expressed as a quadratic expression of the initial state and the degrees of freedom $\mathbf{c}_t = [c_{t|t}^T, \dots, c_{t+N|t}^T]^T$:

$$\tilde{J}(x_t, \mathbf{c}_t) = \begin{bmatrix} x_t^T & \mathbf{c}_t^T & 1 \end{bmatrix} \Omega \begin{bmatrix} x_t^T & \mathbf{c}_t^T & 1 \end{bmatrix}^T. \quad (2.88)$$

The stochastic MPC strategy then consists of minimizing the same objective function as in [21] given by (2.88), subject to constraints (2.86) and $z_{t+N|t} \in \mathbb{Z}_{\infty}$, and then implementing the first predicted input given by $u_t = Kx_t + c_{t|t}^*$. Under appropriate assumptions, it can be guaranteed that the closed-loop application of this strategy yields the Lyapunov inequality of (2.72), which consequently implies the quadratic stability criterion of (2.73).

The proposed strategy, in spite of using tightening parameters that are as small as possible, can still be improved, for instance by changing the predicted control policy. The strategy is only quasi closed-loop because, in spite of the fact that there is feedback in the control policy of the closed-loop paradigm, the online optimal control problem is open-loop in the degrees of freedom $\mathbf{c}_t$ and there is no special feedback treatment of future (unknown) uncertainties.

It is suggested in [50] that more general control policies can effectively be used to find smaller tightening parameters, which provides increased control authority. The authors consider a control policy given by

$$u_{t+k|t} = Kx_{t+k|t} + c_{t+k|t} + v_{t+k|t}, \quad k = 0, 1, 2, \dots \quad (2.89)$$

where $v_{t+k|t}$ are chosen to eliminate the effect of the future disturbances that are

known at the prediction instant $t+k|t$ in the constraints. For this purpose, note that the predictions for the state constraint terms $fx_{t+k|t}$ are given by

$$\begin{aligned} fx_{t+1|t} &= f(\Phi x_t + Bc_t + Dw_t) \\ fx_{t+2|t} &= f(\Phi^2 x_t + \Phi Bc_t + Bc_{t+1|t} + Bv_{t+1|t} + \Phi Dw_t + Dw_{t+1}) \\ &\vdots \end{aligned} \quad (2.90)$$

Since $v_{t+k|t}$ is being applied at the same time that w_{t+k} is being applied, nothing can be done to eliminate the effect of w_{t+k} . However, the effect of $w_t, w_{t+1}, \dots, w_{t+k-1}$ is already known, so their effect can be eliminated by choosing $v_{t+k|t}$ such that:

$$\begin{bmatrix} fB & 0 & 0 & \cdots \\ f\Phi B & fB & 0 & \cdots \\ f\Phi^2 B & f\Phi B & fB & \cdots \\ \vdots & \vdots & \vdots & \ddots \end{bmatrix} \begin{bmatrix} v_{t+1|t} \\ v_{t+2|t} \\ v_{t+3|t} \\ \vdots \end{bmatrix} + \begin{bmatrix} f\Phi & 0 & 0 & \cdots \\ f\Phi^2 & f\Phi & 0 & \cdots \\ f\Phi^3 & f\Phi^2 & f\Phi & \cdots \\ \vdots & \vdots & \vdots & \ddots \end{bmatrix} \begin{bmatrix} Dw_t \\ Dw_{t+1} \\ Dw_{t+2} \\ \vdots \end{bmatrix} = 0 \quad (2.91)$$

which can be achieved by setting $v_{t+k|t}$, for some sequence $\{L_0, L_1, L_2, \dots\}$, to be given by

$$\begin{bmatrix} v_{t|t} \\ v_{t+1|t} \\ v_{t+2|t} \\ v_{t+3|t} \\ \vdots \end{bmatrix} = \underbrace{\begin{bmatrix} 0 & 0 & 0 & \cdots \\ L_0 & 0 & 0 & \cdots \\ L_1 & L_0 & 0 & \cdots \\ L_2 & L_1 & L_0 & \cdots \\ \vdots & \vdots & \vdots & \ddots \end{bmatrix}}_{\mathbf{L}} \begin{bmatrix} Dw_t \\ Dw_{t+1} \\ Dw_{t+2} \\ Dw_{t+3} \\ \vdots \end{bmatrix}. \quad (2.92)$$

This implies that (2.89) can be re-written as

$$u_{t+k|t} = Kx_{t+k|t} + c_{t+k|t} + v_{t+k|t}, \quad k = 0, 1, \dots, \quad (2.93)$$

where

$$v_{t|t} = 0, \quad v_{t+k|t} = \sum_{j=0}^{k-1} L_j D w_{t+k-j-1}, \quad k = 1, 2, \dots, \quad (2.94)$$

which is affine-in-the-disturbances like the policy of (2.45) [58], [38], but with two differences: the compensation given by $v_{t+k|t}$ goes into Mode 2 for all prediction horizons, and the matrix L has a striped structure.

The elimination of the effect in the constraints of the future disturbances that are known at the prediction instant $t + k|t$ means that $\beta_k = 0$, and then γ_1 will be the only tightening in the constraints of (2.86).

This exact elimination can only be performed when there is only one constraint present. If there are more, the above suggests that the gains L_k can be chosen sequentially to achieve the minimization of some norm of the effect of the disturbances in the constraint terms. This work has proved that the tightening parameters β_k can be improved, but no further analysis regarding the stability or feasibility of the receding horizon implementation has been carried out. The development of an SMPC strategy with improved tightening parameters with recursive feasibility and stability guarantees is presented in Chapter 6.

Chapter 3

Offline tube design for efficient implementation of Parameterized Tube Model Predictive Control

Recently introduced parameterized tube model predictive control (PTMPC) [73] is a control technique for linear time-invariant systems subject to additive uncertainty. It supersedes, in terms of feasibility, the RMPC technique that uses an affine-in-the-disturbances control policy. In common with the affine-in-the-disturbances technique, PTMPC online implementation reduces to a standard convex optimization for which the number of variables and constraints grow quadratically with respect to the prediction horizon N . This chapter considers the reduction of the rate of growth of the optimization size to linear in N by allowing part of the computation to be performed offline. It is shown that the suggested offline simplifications preserve the desirable control theoretic properties of PTMPC, and it is illustrated that the reduction in computation can be gained at a modest cost in terms of size of the domain of attraction. Furthermore, the reduction of the online computation burden is such that the prediction horizon of the proposed strategy can be increased in order to get larger

domains of attraction than PTMPC while still having online problems with less (or equal) computational effort.

3.1 Introduction

As discussed in Section 2.2, RMPC is an area of considerable practical importance which also presents significant challenges. Ideal RMPC uses DP [9] or feedback optimization [86] which are impracticable on account of computational complexity, and thus it has been necessary to seek a good compromise between sub-optimality and complexity of the online optimization problems. Some of the approaches developed consider quasi closed-loop policies that are computationally efficient, with a number of variables and constraints of the online optimization that grow linearly with N , but can be conservative (i.e., rigid tube MPC [36], [61], [28], [63], [81], [53]). The degree of optimality can be improved by using the affine-in-the-disturbances control policy [58], [38], but the number of decision variables and constraints of the online optimization grow quadratically with N . These RMPC methods were recently superseded in terms of feasibility by PTMPC [73], which is shown to be equivalent to DP in a broader range of cases than the strategies with the affine-in-the-disturbances policy [10]. PTMPC uses a separable prediction structure and a separable non-linear control policy. The key is the use of $1 + qN$ (where q is the number of vertices of the admissible disturbance set) suitably defined partial control sequences in order to generate controlled and dynamically consistent uncertain state and control predictions. This approach results in a triangular separable prediction structure that induces $N + 1$ partial state and control tubes. The 0^{th} (or nominal) partial tube pair accounts for the dynamical contribution of the initial prediction system state while each of the partial tubes captures the dynamical contribution associated to a future disturbance acting on the system during the Mode 1. However, the triangular structure results in a number of variables and constraints of the online optimization that also grow

quadratically with N .

The aim of the work presented in this chapter is to reduce the computational complexity of PTMPC synthesis yet degrading (sub-)optimality as little as possible while preserving the desirable system theoretic properties (namely recursive feasibility and exponential convergence to some minimal robust invariant set). To this end, attention is restricted to offline design of the partial tube pairs that account for future disturbances by considering the qN partial extreme sequences that generate from the 1st until the N^{th} partial tube pairs, while the nominal control sequences inducing the 0th partial tube pair (which are in fact sequences of points, i.e., singleton sets) are left to be optimized online. This simplification yields an online optimization problem on the nominal sequences with tightened constraints (as in [36], [61], [28], [63], [81], [53]), and reduces the online computational complexity from quadratic to linear. Yet, due to the non-linearity of the underlying control policy (parts of which are computed offline), the overall approach preserves a number of desirable system theoretic properties associated with PTMPC. Clearly there is a cost to pay in terms of a reduction in the size of the domain of attraction, but as will be shown by simulations, this reduction can be small. Furthermore, the reduction of the online computation burden is such that the prediction horizon of the proposed strategy can be increased in order to get larger domains of attraction than PTMPC while still having online problems with less (or equal) computational effort.

In the interest of generality, the simplified PTMPC approach is presented considering mixed state and control constraints; it is noted that the presentation in [73] considers separated state and control constraints.

Section 3.2 provides the setting and re-visits the PTMPC background. The simplifications of the PTMPC synthesis based on offline computations are discussed in Section 3.3. Finally, Sections 3.4 and 3.5 present simulation results and conclusions.

3.2 System description and PTMPC background

Consider a linear, time-invariant, discrete time system with additive disturbances

$$x_{t+1} = Ax_t + Bu_t + w_t, \quad (3.1)$$

where $x_t \in \mathbb{R}^n$, $w_t \in \mathbb{W} \subseteq \mathbb{R}^n$, and $u_t \in \mathbb{R}^{n_u}$ are the state, disturbance and control input. w_t is unknown at time t but is known to belong to a polytopic set given by

$$\mathbb{W} = \text{conv}\{\tilde{w}_i : i = 1, \dots, q\}. \quad (3.2)$$

The system is subject to mixed state and input constraints

$$(x_t, u_t) \in \mathbb{Y}, \quad (3.3)$$

where $\mathbb{Y}$ is a polytopic set that contains the origin in its interior and is given by

$$\mathbb{Y} = \{(x, u) : Fx + Gu \leq \underline{1}\}, \quad (3.4)$$

where $F \in \mathbb{R}^{n_c \times n}$ and $G \in \mathbb{R}^{n_c \times n_u}$.

Assumption 3.1.

- (i) (A, B) is stabilizable.
- (ii) The matrix K is such that $\Phi = A + BK$ is strictly Hurwitz and the minimal robust positively invariant set for system (3.1) under $u = Kx$, $\Omega_\infty(\mathbb{W})$, satisfies $(I, K)\Omega_\infty(\mathbb{W}) \subseteq \text{interior}(\mathbb{Y})$.
- (iii) The terminal constraint set $\mathbb{X}_f$ is the maximal robust positively invariant set for system (3.1) under $u = Kx$ subject to constraints (3.3), and is a polytopic set that

contains the origin in its interior and is given by

$$\mathbb{X}_f = \{x : Hx \leq \underline{1}\}, \quad (3.5)$$

where $H \in \mathbb{R}^{n_f \times n}$.

Recall from Section 2.2.4.2 that a separable prediction scheme is assumed for Mode 1, where the nominal dynamics and the contribution of each future disturbance are separated into different sequences. The 0^{th} partial state and control sequences are denoted $\mathbf{x}_{(0,N)} = \{x_{(0,0)}, \dots, x_{(0,N)}\}$ and $\mathbf{u}_{(0,N-1)} = \{u_{(0,0)}, \dots, u_{(0,N-1)}\}$, and account for the nominal dynamics

$$\begin{aligned} x_{(0,0)} &= x_t, \\ x_{(0,k+1)} &= Ax_{(0,k)} + Bu_{(0,k)}, \quad k = 0, \dots, N-1, \end{aligned} \quad (3.6)$$

and $\mathbf{x}_{(j,N)} = \{x_{(j,j)}, x_{(j,j+1)}, \dots, x_{(j,N)}\}$, $\mathbf{u}_{(j,N-1)} = \{u_{(j,j)}, u_{(j,j+1)}, \dots, u_{(j,N-1)}\}$, are the j^{th} partial state and control sequences describing the dynamical contribution of the disturbance $w_{(j-1)}$ which acts on the system at the $(j-1)^{th}$ prediction instant. These dynamics are given by

$$\begin{aligned} x_{(j,j)} &= w_{(j-1)}, \\ x_{(j,k+1)} &= Ax_{(j,k)} + Bu_{(j,k)}, \quad k = j, \dots, N-1. \end{aligned} \quad (3.7)$$

Recall that for simplicity the notation of the type $x_{t+k|t}$ is not used for the k -step-ahead predictions made at time t , and instead $(x_{(0,k)}, u_{(0,k)})$ and $(x_{(j,k)}, u_{(j,k)})$ are the k -step-ahead predictions made at each time t associated with the 0^{th} partial sequences and the j^{th} partial sequences, while $(x_{(k)}, u_{(k)})$ are the full k -step-ahead predictions

made at each time t . Accordingly, the full predictions for $k = 0, \dots, N$ are given by

$$x_{(k)} = \sum_{j=0}^k x_{(j,k)}, \quad k = 0, \dots, N, \quad u_{(k)} = \sum_{j=0}^k u_{(j,k)}, \quad k = 0, \dots, N-1, \quad (3.8)$$

with

$$\begin{aligned} x_{(j,k)} \in X_{(j,k)} &= \begin{cases} x_{(0,k)} & \text{if } k = 0 \\ \text{conv}\{x_{(i,j,k)} : i = 1, \dots, q\} & \text{if } k = 1, \dots, N \end{cases} \\ u_{(j,k)} \in U_{(j,k)} &= \begin{cases} u_{(0,k)} & \text{if } k = 0 \\ \text{conv}\{u_{(i,j,k)} : i = 1, \dots, q\} & \text{if } k = 1, \dots, N-1. \end{cases} \end{aligned} \quad (3.9)$$

Eq.(3.8) defines a triangular prediction scheme for Mode 1 as shown in Figure 3.1 and implies that

$$\begin{aligned} x_{(k)} \in X_{(k)} &= \bigoplus_{j=1}^k X_{(j,k)}, \quad k = 0, \dots, N, \\ u_{(k)} \in U_{(k)} &= \bigoplus_{j=1}^k U_{(j,k)}, \quad k = 0, \dots, N-1, \end{aligned} \quad (3.10)$$

where $X_{(k)}$ and $U_{(k)}$ are the state and control cross-sections of the state tube $\mathbf{X}_N = \{X_{(0)}, \dots, X_{(N)}\}$ and control tube $\mathbf{U}_{N-1} = \{U_{(0)}, \dots, U_{(N-1)}\}$, respectively, and can be obtained as the column-wise (Minkowski) sums in Figure 3.1. From (3.9), the j^{th} partial state and control sequences are defined in terms of the j^{th} extreme partial state and control sequences, $\mathbf{x}_{(i,j,N)} = \{x_{(i,j,j)}, x_{(i,j,j+1)}, \dots, x_{(i,j,N)}\}$ and $\mathbf{u}_{(i,j,N-1)} = \{u_{(i,j,j)}, u_{(i,j,j+1)}, \dots, u_{(i,j,N-1)}\}$, respectively, which describe the effect of the extreme disturbances acting at the prediction instant $j-1$. The dynamics of the j^{th} partial

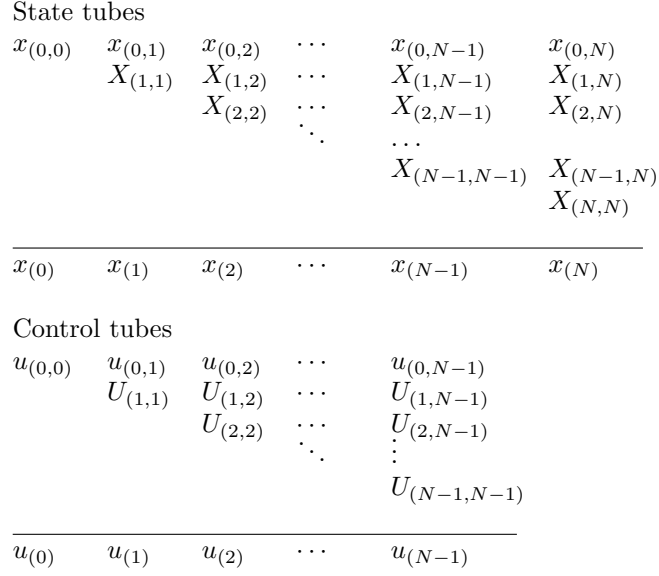


Figure 3.1: Triangular prediction scheme of PTMPC

extreme state and control sequences are given by

$$\begin{aligned}
 x_{(i,j,j)} &= \tilde{w}_i, & j &= 1, \dots, N, \ i = 1, \dots, q \\
 x_{(i,j,k+1)} &= Ax_{(i,j,k)} + Bu_{(i,j,k)}, & j &= 1, \dots, N-1, \ k = j, \dots, N-1, \ i = 1, \dots, q.
 \end{aligned} \tag{3.11}$$

The partial sequences as defined by (3.6) and (3.11) imply that the predicted $(x_{(k)}, u_{(k)})$ will be contained in the overall parameterized mixed state and control tubes (overall PMSCT)

$$Y_{(k)} = \bigoplus_{j=0}^k Y_{(j,k)}, \quad k = 0, \dots, N-1, \tag{3.12}$$

where the partial parameterized mixed state and control tubes (partial PMSCT) $\mathbf{Y}_{(j,N-1)} = \{Y_{(j,j)}, Y_{(j,j+1)}, \dots, Y_{(j,N-1)}\}$ contain the predicted $(x_{(j,k)}, u_{(j,k)})$, where $Y_{(0,k)} = (x_{(0,k)}, u_{(0,k)})$, $Y_{(j,k)} = \text{conv}\{(x_{(i,j,k)}, u_{(i,j,k)}) : i = 1, \dots, q\}$. The partial sequences as defined by (3.6) and (3.11) also imply that the terminal state will be contained in the overall parameterized terminal state tube (overall PTST) cross sec-

tion

$$X_{(N)} = \bigoplus_{j=0}^N X_{(j,N)}, \quad (3.13)$$

where the terminal partial state cross sections $X_{(j,N)}$ are given by $X_{(0,N)} = x_{(0,N)}$ and $X_{(j,N)} = \text{conv}\{x_{(i,j,N)} : i = 1, \dots, q\}$.

While the separable prediction scheme is used for the analysis of the predictions in Mode 1, a terminal control law is assumed in Mode 2 of the form $u_{(k)} = Kx_{(k)}$, $k \geq N$. Under these considerations, the conditions for constraint satisfaction are

$$\begin{aligned} Y_{(k)} &\subseteq \mathbb{Y}, \quad k = 0, \dots, N-1, \\ X_{(N)} &\subseteq \mathbb{X}_f. \end{aligned} \quad (3.14)$$

In the online problem these conditions can be guaranteed equivalently by linear constraints on the 0^{th} partial state and control sequences by using slack variables that account for the worst (extreme) realizations of the extreme j^{th} partial state and control sequences.

$$\begin{aligned} F_l x_{(0,0)} + G_l u_{(0,0)} &\leq 1, \quad l = 1, \dots, n_c, \\ F_l x_{(0,k)} + G_l u_{(0,k)} + \sum_{j=1}^k f_{(l,j,k)} &\leq 1, \quad l = 1, \dots, n_c, \quad k = 1, \dots, N-1 \\ H_l x_{(0,N)} + \sum_{j=1}^N h_{(l,j,N)} &\leq 1, \quad l = 1, \dots, n_f, \end{aligned} \quad (3.15)$$

and for $i = 1, \dots, q$, $k = j, \dots, N-1$,

$$\begin{aligned} F_l x_{(i,j,k)} + G_l u_{(i,j,k)} &\leq f_{(l,j,k)}, \quad l = 1, \dots, n_c, \quad j = 1, \dots, N-1, \\ H_l x_{(i,j,N)} &\leq h_{(l,j,N)}, \quad l = 1, \dots, n_f, \quad j = 1, \dots, N, \end{aligned}$$

where F_l , G_l and H_l are the l^{th} rows of F , G and H , respectively.

The separable prediction scheme induces a separable non-linear predicted control

policy $\pi_{N-1} = \{\mu_0(\cdot), \dots, \mu_{N-1}(\cdot)\}$ given by,

$$\begin{aligned}\mu_k(x_{(k)}) &= \sum_{j=0}^k \mu_{(j,k)}(x_{(j,k)}), \text{ with } x_{(k)} = \sum_{j=0}^k x_{(j,k)} \\ \mu_{(0,k)}(x_{(0,k)}) &= u_{(0,k)}, \\ \mu_{(j,k)}(x_{(j,k)}) &= \sum_{i=1}^q \lambda_{(i,j)}(w_{(j-1)}) u_{(i,j,k)}, \quad j = 1, \dots, N-1, \quad k = j, \dots, N-1\end{aligned}\tag{3.16}$$

where $\lambda_{(i,j)}(w_{(j-1)})$ are convex interpolation parameters that satisfy $\lambda_{(i,j)} \geq 0$ for $i = 1, \dots, q$, $\sum_{i=1}^q \lambda_{(i,j)}(w_{(j-1)}) = 1$ and $\sum_{i=1}^q \lambda_{(i,j)}(w_{(j-1)}) \tilde{w}_i = w_{(j-1)}$. A well-defined and unique selection of such convex multipliers $\{\lambda_{(i,j)}, i = 1, \dots, q\}$ can be obtained employing, for instance, a least squares procedure.

Although the current and future disturbances $w_{(j-1)}$, $j = 1, 2, \dots$, are unknown at the current instant, they will become known at prediction time j . The predicted control policy accounts for this fact in that the predicted inputs $u_{(k)} = \mu_k(x_{(k)})$ depend on the interpolation parameters $\lambda_{(i,j)}$, $j = 1, \dots, k$ that are associated with all the disturbances $w_{(j-1)}$ that are currently unknown, but would be known at prediction time k . For the application of the MPC strategy however, it is only needed to know that the predicted control laws for $k = 1, \dots, N-1$ exist in order to guarantee feasibility of the predictions; it is not necessary to find them since only the first one, $u_{(0)} = \mu_0(x_{(0)})$, is actually used and applied to the system. The existence is ensured by finding the predictions associated with the extreme realizations of the disturbances, and as a result, it is not necessary to find the interpolation parameters $\lambda_{(i,j)}$.

An equivalent re-parameterization, which will be useful for the definition of the objective function and the stability analysis, is considered such that the 0^{th} partial sequences are re-written as,

$$\begin{aligned}x_{(0,k)} &= z_{(0,k)} + \Phi^k(x_t - z_{(0,0)}), \quad k = 0, \dots, N, \\ u_{(0,k)} &= v_{(0,k)} + K\Phi^k(x_t - z_{(0,0)}), \quad k = 0, \dots, N-1,\end{aligned}\tag{3.17}$$

where $\mathbf{z}_{(0,N)} = \{z_{(0,0)}, z_{(0,1)}, \dots, z_{(0,N)}\}$ and $\mathbf{v}_{(0,N-1)} = \{v_{(0,0)}, v_{(0,1)}, \dots, v_{(0,N-1)}\}$ satisfy

$$z_{(0,k+1)} = Az_{(0,k)} + Bv_{(0,k)}, \quad k = 0, \dots, N-1, \quad (3.18)$$

$$H_l(x_t - z_{(0,0)}) \leq 1, \quad l = 1, \dots, n_f. \quad (3.19)$$

$\mathbf{z}_{(0,N)}$ and $\mathbf{v}_{(0,N-1)}$ represent the deviation of $\mathbf{x}_{(0,N)}$ and $\mathbf{u}_{(0,N-1)}$ from an evolution generated by the state feedback $u = Kx$ (which would be feasible if the state is inside $\mathbb{X}_f$) where $u_{(0,k)} = Kx_{(0,k)}$ and $x_{(0,k)} = \Phi^k x_{(0,0)}$. Likewise, the j^{th} partial extreme sequences follow the re-parameterization given by

$$\begin{aligned} x_{(i,j,k)} &= z_{(i,j,k)} + \Phi^{k-j} \tilde{w}_i, \quad j = 1, \dots, N, \quad k = j, \dots, N, \quad i = 1, \dots, q, \\ u_{(i,j,k)} &= v_{(i,j,k)} + K\Phi^{k-j} \tilde{w}_i, \quad j = 1, \dots, N-1, \quad k = j, \dots, N-1, \quad i = 1, \dots, q, \end{aligned} \quad (3.20)$$

where $\mathbf{z}_{(i,j,N)} = \{z_{(i,j,j)}, z_{(i,j,j+1)}, \dots, z_{(i,j,N)}\}$ and $\mathbf{v}_{(i,j,N-1)} = \{v_{(i,j,j)}, v_{(i,j,j+1)}, \dots, v_{(i,j,N-1)}\}$ satisfy

$$\begin{aligned} z_{(i,j,j)} &= 0, \quad j = 1, \dots, N, \quad i = 1, \dots, q, \\ z_{(i,j,k+1)} &= Az_{(i,j,k)} + Bv_{(i,j,k)}, \quad j = 1, \dots, N-1, \quad k = j, \dots, N-1, \quad i = 1, \dots, q. \end{aligned} \quad (3.21)$$

Likewise, $\mathbf{z}_{(i,j,N)}$ and $\mathbf{v}_{(i,j,N-1)}$ represent the deviation of $\mathbf{x}_{(i,j,N)}$ and $\mathbf{u}_{(i,j,N-1)}$ from the evolution generated by $u = Kx$, i.e. $u_{(i,j,k)} = Kx_{(i,j,k)}$ and $x_{(i,j,k)} = \Phi^{k-j} x_{(i,j,j)}$.

The cost function penalizes the terms $z_{(0,k)}$, $v_{(0,k)}$, $z_{(i,j,k)}$ and $v_{(i,j,k)}$, and is structured as the sum of the costs associated to each partial extreme sequence, such that, if $\mathbf{d}_N$ contains all the decision variables of the formulation,

$$V_N(\mathbf{d}_N) = \sum_{j=0}^N V_{(j,N)}(\mathbf{d}_N), \quad (3.22)$$

where

$$\begin{aligned}
 V_{(0,N)}(\mathbf{d}_N) &= \sum_{k=0}^{N-1} \ell(z_{(0,k)}, v_{(0,k)}) + V_f(z_{(0,N)}), \\
 V_{(j,N)}(\mathbf{d}_N) &= \sum_{i=1}^q V_{(i,j,N)}(\mathbf{d}_N), \quad j = 1, \dots, N-1, \quad V_{(N,N)}(\mathbf{d}_N) = \sum_{i=1}^q V_f(z_{(i,N,N)}), \\
 V_{(i,j,N)}(\mathbf{d}_N) &= \sum_{k=j}^{N-1} \ell(z_{(i,j,k)}, v_{(i,j,k)}) + V_f(z_{(i,j,N)}), \quad i = 1, \dots, q, \quad j = 1, \dots, N-1, \\
 \ell(z, v) &= |z|_{\mathcal{Q}} + |v|_{\mathcal{R}}, \quad V_f(z) = |z|_{\mathcal{P}},
 \end{aligned}$$

$\mathcal{Q}, \mathcal{P}, \mathcal{R}$ are symmetric polytopic sets that contain the origin in their interiors and $V_f(\cdot)$ satisfies the Lyapunov condition $V_f(\Phi z) - V_f(x) \leq -\ell(z, Kz)$ for all $x \in \mathbb{R}^n$.

The online optimal control problem seeks to minimize $V_N(\mathbf{d}_N)$. Since the cost penalizes the terms $z_{(0,k)}$, $v_{(0,k)}$, $z_{(i,j,k)}$, and $v_{(i,j,k)}$, it clearly attempts to penalize the deviations from the predicted state and input sequences generated by $u = Kx$. As a result, if the constraints become inactive, which will happen when $x_t \in \mathbb{X}_f$, the evolution of the system generated by $u = Kx$ is feasible, in which case $z_{(0,k)} = 0$, $v_{(0,k)} = 0$, $z_{(i,j,k)} = 0$ and $v_{(i,j,k)} = 0$ are a feasible solution. Therefore the value function $V_N^0(x)$ will be zero, and one will recover the state feedback $u = Kx$ from the optimal control problem. On the other hand, if $x \notin \mathbb{X}_f$ then $z_{(0,k)} = 0$, $v_{(0,k)} = 0$, $z_{(i,j,k)} = 0$ and $v_{(i,j,k)} = 0$ is not feasible and then $V_N^0(x) > 0$. Since the tube structure guarantees recursive feasibility and the decrease condition $V_N^0(x_{t+1}) - V_N^0(x_t) \leq -(|x_t|_{\mathcal{Q}} + |u_t|_{\mathcal{R}})$ can be obtained from the online application, it follows that $V_N^0(x)$ is a Lyapunov function outside of $\mathbb{X}_f$, which leads to the exponential convergence of the state to $\mathbb{X}_f$. Furthermore, since in this set the control law becomes $u = Kx$, the state converges exponentially to the associated minimal robust invariant set $\Omega_\infty(\mathbb{W})$.

The minimization of (3.22), which is a sum of gauge functions of polytopic sets, can be performed equivalently (which as pointed out in [73] is customary in convex analysis) by minimizing the sum of a collection of slack decision variables that, via

additional linear constraints, upper bound tightly the optimal value of the underlying gauge functions. Thus the online optimal control problem can be solved by minimization of a linear objective subject to linear constraints, which is an LP.

The most important advantage of PTMPC [73] is that, since the predicted control laws $\mu_k(x_{(k)})$ depend on the interpolation parameters $\lambda_{(i,j)}(\cdot)$, $i = 1, \dots, q$, that are non-linear functions of $w_{(j-1)}$ (for instance, when obtained by a least-squares procedure, they are piecewise affine functions of $w_{(j-1)}$), the associated control policy is formed from separable non-linear control laws of the past disturbances. Thus the control policy of (3.14) is more general than others such as the closed-loop paradigm [82], [84] and the affine-in-the-disturbances approach [38], [58], which are subsumed as special cases of PTMPC.

As already indicated, the main objective here is to reduce the quadratic growth with N of the number of decision variables and constraints required for the online implementation of PTMPC. For this reason, the focus is on an intuitive idea based on an offline optimization of the j^{th} partial extreme state and control sequences satisfying (3.11) and an online optimization of the 0^{th} partial state and control sequences satisfying (3.6).

3.3 PTMPC: An offline Simplification

The simplified PTMPC strategy, which is based on the offline computation of the j^{th} partial extreme state and control partial sequences that satisfy (3.11), and the online optimization of the 0^{th} state and control sequences, is described in this section.

As will be seen below, once the j^{th} partial extreme state and control sequences are computed offline, they are available to be used of the online optimization for the first instant $t = 0$, but need to be updated for the future instants $t = 1, \dots, N$. This update can be carried out according to offline computations. The online optimization, due to these updates, will be time varying, and then it is necessary to make a temporal index

explicit. Let $\hat{\mathbf{x}}_{((i,j,N),t)} = \{ \hat{x}_{((i,j,j),t)}, \hat{x}_{((i,j,j+1),t)}, \dots, \hat{x}_{((i,j,N),t)} \}$ and $\hat{\mathbf{u}}_{((i,j,N-1),t)} = \{ \hat{u}_{((i,j,j),t)}, \hat{u}_{((i,j,j+1),t)}, \dots, \hat{u}_{((i,j,N-1),t)} \}$, be the j^{th} partial extreme state and control sequences to be used at time t . Assume in this section that the offline computations of the j^{th} partial extreme state and control partial sequences to be used at time $t = 0$, $\hat{\mathbf{x}}_{((i,j,N),0)}$ and $\hat{\mathbf{u}}_{((i,j,N-1),0)}$, has already been performed and are ready to be used by the controller.

Assumption 3.2. $N \geq 2$.

Remark 3.1. It will be assumed that $N \geq 2$ for the developments of Section 3.3 until the end of Section 3.3.1. The reason behind this is that if $N = 1$, then there are no disturbance compensation terms that can be designed offline, since only the first input $u_{(0,0)}$ is free. From Section 3.3.2 onwards one can also choose $N = 1$ under the assumption that the offline design has been skipped (in this case only the slacks $\hat{h}_{((i,j,N),t)}$ are needed, and can be found as will be shown in (3.31)).

Assumption 3.3. In common with (3.11), $\hat{\mathbf{x}}_{((i,j,N),t)}$ and $\hat{\mathbf{u}}_{((i,j,N-1),t)}$ satisfy for $i = 1, \dots, q$, $t = 0, 1, \dots$

$$\begin{aligned} \hat{x}_{((i,j,j),t)} &= \tilde{w}_i, & j &= 1, \dots, N, \\ \hat{x}_{((i,j,k+1),t)} &= A\hat{x}_{((i,j,k),t)} + B\hat{u}_{((i,j,k),t)}, & j &= 1, \dots, N-1, \quad k = j, \dots, N-1. \end{aligned} \quad (3.23)$$

The overall PMSCT $\hat{Y}_{(k)}$, for $k = 1, \dots, N-1$, and the overall PTST cross-section $\hat{X}_{(N)}$ satisfy for $t = 0$

$$\hat{Y}_{(k)} = \bigoplus_{j=1}^k \hat{Y}_{((j,k),t)} \subseteq \mathbb{Y}, \quad \text{and} \quad \hat{X}_{(N)} = \bigoplus_{j=1}^N \hat{X}_{((j,N),t)} \subseteq \mathbb{X}_f, \quad (3.24)$$

where the corresponding partial cross-sections are given by $\hat{Y}_{((j,k),t)} = \text{conv}\{(\hat{x}_{((i,j,k),t)}, \hat{u}_{((i,j,k),t)}) : i = 1, \dots, q\}$ and $\hat{X}_{((j,N),t)} = \text{conv}\{\hat{x}_{((i,j,N),t)} : i = 1, \dots, q\}$.

The role of (3.24) in this assumption is to guarantee that the domain of attraction

is not empty. Under the expressed conditions at least the origin, i.e. $x_t = 0$, is feasible when solving the online optimization problem.

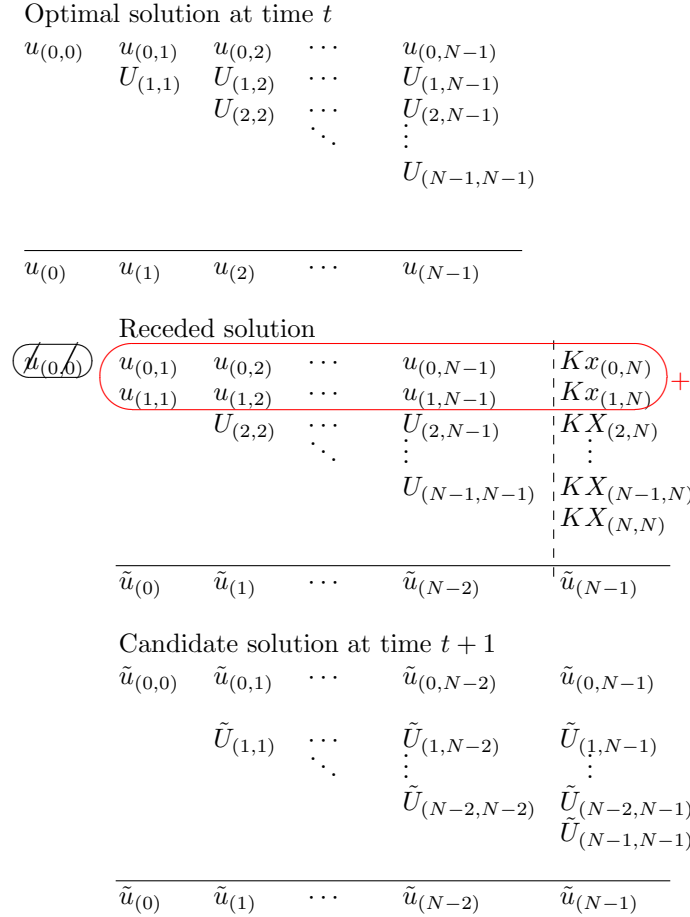


Figure 3.2: Receded candidate solutions that yield recursive feasibility and stability in PTMPC

The invariance and stability analyses reported in [73] are based on the definition of a candidate solution for the online optimization problem of time $t + 1$, which is defined using the optimal solution for the online optimization problem of time t . The construction of such a candidate solution is based on the *receding horizon* strategy, in what will be called a *receding* operation, and gives a *receded* solution. This procedure is illustrated in Figure 3.2 for the inputs and is described next. Since the new problem deals with predictions starting from time $t + 1$, the predictions associated with $k = 0$ have to be discarded from the solution at time t , and each new input $\tilde{u}_{(k)}$ of time

$t + 1$ is associated with the input $u_{(k+1)}$ of time t . u_N is not defined by the separable prediction scheme, but it is known that in Mode 2 the control law is given by $u = Kx$, so then $\tilde{u}_{(N-1)}$ is defined according to that state feedback. Finally, the effect of the disturbance w_t will be available in x_{t+1} which implies that the 0^{th} sequences at time $t + 1$ have to account for the effect of w_t . Such effect is contained in the 1^{st} partial tubes $\mathbf{Y}_{(1,N-1)}$ of time t ; since at time $t + 1$ the effect of w_t is known, the sequence of sets $\mathbf{Y}_{(1,N-1)}$ collapses to a sequence of elements which is determined by the value of w_t . Then the candidate 0^{th} sequences of time $t + 1$ will be composed by the 0^{th} sequences plus a sequence of elements that are contained in the 1^{st} partial tubes of time t .

With these considerations, the precise construction of the candidate solution of time $t + 1$, identified by $(\tilde{\cdot})$, from the optimal solutions at time t , denoted by $(\cdot)^0(x_t)$, is given by

$$\begin{aligned}\tilde{x}_{(0,k)} &= x_{(0,k+1)}^0(x_t) + \sum_{i=1}^q \lambda_{(i,1)}^*(w_t) x_{(i,1,k+1)}^0(x_t), \quad k = 0, \dots, N-1, \\ \tilde{x}_{(0,N)} &= \Phi \tilde{x}_{(0,N-1)},\end{aligned}\tag{3.25}$$

$$\tilde{u}_{(0,k)} = u_{(0,k+1)}^0(x_t) + \sum_{i=1}^q \lambda_{(i,1)}^*(w_t) u_{(i,1,k+1)}^0(x_t), \quad k = 0, \dots, N-2,$$

$$\tilde{u}_{(0,N-1)} = K \tilde{x}_{(0,N-1)},$$

$$\begin{aligned}\tilde{x}_{(i,j,k)} &= x_{(i,j+1,k+1)}^0(x_t), \quad j = 1, \dots, N-1, \quad k = j, \dots, N-1, \quad i = 1, \dots, q \\ \tilde{x}_{(i,j,N)} &= \Phi x_{(i,j+1,N)}^0(x_t), \quad j = 1, \dots, N-1, \quad i = 1, \dots, q, \\ \tilde{x}_{(i,N,N)} &= \tilde{w}_i, \quad i = 1, \dots, q, \\ \tilde{u}_{(i,j,k)} &= u_{(i,j+1,k+1)}^0(x_t), \quad j = 1, \dots, N-2, \quad k = 1, \dots, N-2, \quad i = 1, \dots, q, \\ \tilde{u}_{(i,j,N-1)} &= K \tilde{x}_{(i,j,N-1)}, \quad j = 1, \dots, N-1, \quad i = 1, \dots, q,\end{aligned}\tag{3.26}$$

where $\lambda_{(i,1)}^*(w_t)$ are the least squares (or any other convenient selection criterion)

convex interpolation parameters that satisfy $w_t = \sum_{i=1}^q \lambda_{(i,1)}^*(w_t) \tilde{w}_i$.

The online update of j^{th} partial extreme state and control sequences is defined so as to replicate the *receding* operation that is used in the construction of the candidate solutions of time $t+1$ based on the solutions of time t , as shown in (3.25) and (3.26). This is done in order to enable invariance and stability analyses. Thus the j^{th} partial extreme state and control sequences to be used at instant t , $\hat{\mathbf{x}}_{((i,j,N),t)}$ and $\hat{\mathbf{u}}_{((i,j,N-1),t)}$, are updated offline by

$$\begin{aligned} \hat{x}_{((i,j,k),t+1)} &= \hat{x}_{((i,j+1,k+1),t)}, \quad j = 1, \dots, N-1, \quad k = j, \dots, N-1, \quad i = 1, \dots, q, \\ \hat{x}_{((i,j,N),t+1)} &= \Phi \hat{x}_{((i,j+1,N),t)}, \quad j = 1, \dots, N-1, \quad i = 1, \dots, q, \\ \hat{x}_{((i,N,N),t+1)} &= \tilde{w}_i, \quad i = 1, \dots, q, \end{aligned} \tag{3.27}$$

and

$$\begin{aligned} \hat{u}_{((i,j,k),t+1)} &= \hat{u}_{((i,j+1,k+1),t)}, \quad j = 1, \dots, N-2, \quad k = j, \dots, N-2, \quad i = 1, \dots, q, \\ \hat{u}_{((i,j,N-1),t+1)} &= K \hat{x}_{((i,j+1,N),t)}, \quad j = 1, \dots, N-2, \quad i = 1, \dots, q, \\ \hat{u}_{((i,N-1,N-1),t+1)} &= K \hat{x}_{((i,N,N),t)}, \quad i = 1, \dots, q. \end{aligned} \tag{3.28}$$

The update equations (3.27) and (3.28) define a tube update scheme by which $\hat{Y}_{((j,k),t+1)} = \hat{Y}_{((j+1,k+1),t)}$, for $j = 1, \dots, N-2$, $k = j, \dots, N-2$, such that $\hat{U}_{((j,k),t+1)} = \hat{U}_{((j+1,k+1),t)}$, for $j = 1, \dots, N-2$, $k = j, \dots, N-2$, and $\hat{X}_{((j,k),t+1)} = \hat{X}_{((j+1,k+1),t)}$ for $j = 1, \dots, N-1$, $k = j, \dots, N-1$, where $\hat{X}_{((j,k),t)} = \text{conv}\{\hat{x}_{((i,j,k),t)} : i = 1, \dots, q\}$ and $\hat{U}_{((j,k),t)} = \text{conv}\{\hat{u}_{((i,j,k),t)} : i = 1, \dots, q\}$. The update scheme is illustrated in Figure 3.3 for the inputs.

Provided that the j^{th} partial extreme state and control sequences are known at each instant t , the constraints of (3.14) become time-varying conditions on the 0^{th} partial state and control sequences $\mathbf{x}_{(0,N)} = \{x_{(0,0)}, x_{(0,1)}, \dots, x_{(0,N)}\}$ and $\mathbf{u}_{(0,N-1)} =$

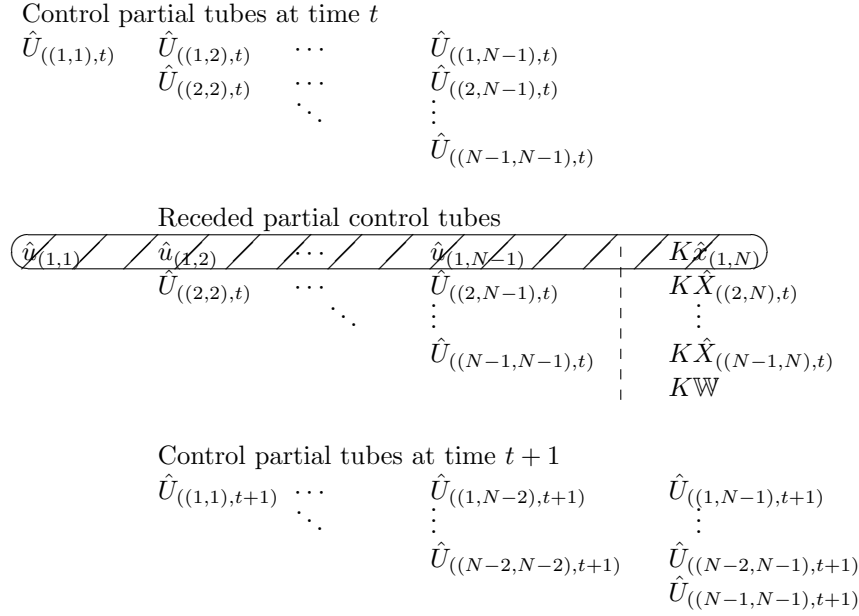


Figure 3.3: Tubes update in the simplified PTMPC

$\{u_{(0,0)}, u_{(0,1)}, \dots, u_{(0,N-1)}\}$ to be found at instant t

$$\begin{aligned}
 Y_{(0,0)} &\subseteq \mathbb{Y}, \\
 Y_{(0,k)} \bigoplus_{j=1}^k \hat{Y}_{((j,k),t)} &\subseteq \mathbb{Y}, \quad k = 1, \dots, N-1, \\
 X_{(0,N)} \bigoplus_{j=1}^k \hat{X}_{((j,N),t)} &\subseteq \mathbb{X}_f.
 \end{aligned} \tag{3.29}$$

where $Y_{(0,k)} = (x_{(0,k)}, u_{(0,k)})$, $k = 0, \dots, N-1$ and $X_{(0,N)} = x_{(0,N)}$. (3.29) can be cast as linear constraints on $\mathbf{x}_{(0,N)}$ and $\mathbf{u}_{(0,N-1)}$

$$\begin{aligned}
 F_l x_{(0,0)} + G_l u_{(0,0)} &\leq 1, & l = 1, \dots, n_c, \\
 F_l x_{(0,k)} + G_l u_{(0,k)} + \sum_{j=1}^k \hat{f}_{((l,j,k),t)} &\leq 1, & l = 1, \dots, n_c, \quad k = 1, \dots, N-1, \\
 H_l x_{(0,N)} + \sum_{j=1}^N \hat{h}_{((l,j,N),t)} &\leq 1, & l = 1, \dots, n_f,
 \end{aligned} \tag{3.30}$$

where the associated slack variables $\hat{f}_{((l,j,k),t)}$ and $\hat{h}_{((l,j,N),t)}$ can be evaluated offline for $t = 0, 1, \dots$ according to

$$\begin{aligned} \hat{f}_{((l,j,k),t)} &= \max_{i=1,\dots,q} F_l \hat{x}_{((i,j,k),t)} + G_l \hat{u}_{((i,j,k),t)}, \quad l = 1, \dots, n_c, \quad j = 1, \dots, N-1, \\ &\quad k = j, \dots, N-1, \\ \hat{h}_{((l,j,N),t)} &= \max_{i=1,\dots,q} H_l + \hat{x}_{((i,j,N),t)}, \quad l = 1, \dots, n_f, \quad j = 1, \dots, N. \end{aligned} \quad (3.31)$$

Remark 3.2. The update equations (3.27) and (3.28) are such that the updated $\hat{\mathbf{x}}_{((i,j,N),t)}$ and $\hat{\mathbf{u}}_{((i,j,N-1),t)}$ satisfy (3.23) for $t = 0, 1, 2, \dots$. It can also be seen that for $t \geq N$, $\hat{\mathbf{x}}_{((i,j,N),t)}$ and $\hat{\mathbf{u}}_{((i,j,N-1),t)}$ satisfy

$$\begin{aligned} \hat{x}_{((i,j,k),t)} &= \Phi^{k-j} \tilde{w}_i, \quad i = 1, \dots, q, \quad j = 1, \dots, N, \quad k = 1, \dots, N \\ \hat{u}_{((i,j,k),t)} &= K \Phi^{k-j} \tilde{w}_i, \quad i = 1, \dots, q, \quad j = 1, \dots, N-1, \quad k = j, \dots, N-1. \end{aligned} \quad (3.32)$$

which do not depend on t . This implies that the slack variables $\hat{f}_{((l,j,k),t)}$ and $\hat{h}_{((l,j,N),t)}$ do not depend on t , for $t \geq N$. So, while constraints (3.29) and (3.30) are time-varying for $t = 0, \dots, N-1$, these become time-invariant for $t \geq N$.

Remark 3.3. The fact that the optimization problems are time-varying implies that there is an increased memory requirement when compared with usual constraint tightening approaches that are time-invariant [28], [81], [53]. More precisely, it requires $N+1$ sets $n_c(N-1) + n_f$ of tightening parameters which map into $N+1$ sets of the $n_c(N-1) + n_f$ tightened inequality constraints of (3.30). If the overall decision variable is $\mathbf{d} \in \mathbb{R}^{n_d}$, the inequalities of (3.30) that involve tightening parameters can be expressed as $E\mathbf{d} \leq \mathbf{r}_t$, where $E \in \mathbb{R}^{(n_c(N-1)+n_f) \times n_d}$ and only $\mathbf{r}_t \in \mathbb{R}^{n_c(N-1)+n_f}$ is affected by the different tightening parameters. Thus one only needs to store $N+1$ different RHS vectors, $\mathbf{r}_t \in \mathbb{R}^{n_c(N-1)+n_f}$ for $t = 0, \dots, N$, while only 1 RHS vector is needed for the time-invariant constraints tightening approaches. Note that there is no need

to store each $\hat{f}_{((l,j,k),t)}$ and $\hat{h}_{((l,j,N),t)}$ since what one actually uses is $1 - \sum_{j=1}^k \hat{f}_{((l,j,k),t)}$ and $1 - \sum_{j=1}^N \hat{h}_{((l,j,N),t)}$, which are contained in r_t .

Remark 3.4. *The update of the slack variables implies that at every time step there will be more terms that are associated with the control law $u = Kx$ in the control policy. Thus the feasibility regions of the optimizations tend to decrease until $t = N$ when this region, as well as the optimization problem, becomes time-invariant. This means that this controller is less robust than others with constant regions of feasibility (which nevertheless, are not fully protected) against unmodelled disturbances that might render the problem infeasible. If this were to happen, one might try to solve the problem again re-setting the optimization problem to $t = 0$ so as to have an increased feasibility region, thus improving the robustness of the method. But of course, this would come at the cost of increased computational burden.*

3.3.1 Offline design of initial partial tube cross sections

Here the offline design of the j^{th} partial extreme state and control sequences, $\hat{\mathbf{x}}_{((i,j,N),0)} = \{\hat{x}_{((i,j,j),0)}, \hat{x}_{((i,j,j+1),0)}, \dots, \hat{x}_{((i,j,N),0)}\}$ and $\hat{\mathbf{u}}_{((i,j,N-1),0)} = \{\hat{u}_{((i,j,j),0)}, \hat{u}_{((i,j,j+1),0)}, \dots, \hat{u}_{((i,j,N-1),0)}\}$, that satisfy (3.23) and (3.24) from Assumption 3.3, is described.

It is clear from (3.30) that the slack variables effect a constraint tightening on the nominal dynamics that allow one to guarantee constraint satisfaction in the presence of uncertainty. Less tight constraints imply larger domains of attraction, and thus one should seek the minimization of the slacks. A sensible procedure for selecting the slacks is to design these sequentially: starting from the first prediction instant given by $k = 1$, proceeding into the second, given by $k = 2$, and so on; this is both computationally convenient and is consistent with the expectation that constraints are more stringent during the immediate transients and less so towards the steady state of predictions.

One more requirement is imposed in addition to constraints (3.23) and (3.24):

that suitable sums of the optimized scalars $\sum_{j=1}^k \hat{f}_{((l,j,k),0)}$ and $\sum_{j=1}^N \hat{h}_{((l,j,k),0)}$ are not greater than those that would have been obtained by using the state linear feedback $u = Kx$ (i.e., obtained by setting each $\hat{u}_{((i,j,k),0)}$ to be equal to $K\hat{x}_{((i,j,k),0)}$) and (3.31). This is a safeguard against the possibility that the sequential slacks design, though convenient, may result in excessively large slacks far in the horizon at the expense of getting small slacks at the beginning of the horizon. This additional constraint can be expressed as

$$\begin{aligned} \sum_{j=1}^k \hat{f}_{((l,j,k),0)} &\leq \sum_{j=1}^k \bar{f}_{(l,j,k)}, \quad l = 1, \dots, n_c, \quad k = 1, \dots, N-1, \\ \sum_{j=1}^N \hat{h}_{((l,j,N),0)} &\leq \sum_{j=1}^k \bar{h}_{(l,j,N)}, \quad l = 1, \dots, n_f, \end{aligned} \quad (3.33)$$

$$\begin{aligned} \hat{f}_{((l,j,k),0)} &\geq F_l \hat{x}_{((i,j,k),0)} + G_l \hat{u}_{((i,j,k),0)}, \quad l = 1, \dots, n_c, \quad j = 1, \dots, N-1, \\ &\quad k = j, \dots, N-1, \end{aligned} \quad (3.34)$$

$$\hat{h}_{((l,j,N),0)} \geq F_l \hat{x}_{((i,j,N),0)} + G_l \hat{u}_{((i,j,N),0)}, \quad l = 1, \dots, n_f, \quad j = 1, \dots, N,$$

where the constants $\bar{f}_{((l,j,k),0)}$ and $\bar{h}_{((l,j,N),0)}$ are obtained by evaluating (3.23) and (3.31) with $\hat{u}_{((i,j,k),0)} = K\hat{x}_{((i,j,k),0)}$. Note that if $\hat{f}_{((l,j,k),0)}$ and $\hat{h}_{((l,j,N),0)}$ are defined by (3.31), condition (3.33) would not be linear. So in order to solve an LP in the offline design, (3.31) is not used to define $\hat{f}_{((l,j,k),0)}$ and $\hat{h}_{((l,j,N),0)}$ and instead the conditions of (3.34) are considered.

Note that satisfaction of conditions (3.33) and (3.34) imply satisfaction of (3.24) since Assumption 3.1 implies that the dynamics under the state feedback $u = Kx$ satisfy constraints (3.3), from which

$$\begin{aligned} \sum_{j=1}^k \bar{f}_{(l,j,k)} &\leq 1, \quad l = 1, \dots, n_c, \quad k = 1, \dots, N-1, \\ \sum_{j=1}^k \bar{h}_{(l,j,N)} &\leq 1, \quad l = 1, \dots, n_c. \end{aligned}$$

and therefore there is no need for (3.24) in the offline optimization.

Let $\hat{\mathbf{x}}_{(N,0)}$ and $\hat{\mathbf{u}}_{(N-1,0)}$ denote the collections of the associated j^{th} partial extreme state and control sequences, $\hat{\mathbf{x}}_{((i,j,N),0)}$ and $\hat{\mathbf{u}}_{((i,j,N-1),0)}$. Let also $\hat{\mathbf{f}}_{(N-1,0)}$ and $\hat{\mathbf{h}}_{(N,0)}$ denote the collections of the slack variables $\hat{f}_{((i,j,k),0)}$ and $\hat{h}_{((i,j,N),0)}$. Finally, let $\hat{\boldsymbol{\eta}}_0$ denote the collection of $\hat{\mathbf{x}}_{(N,0)}$, $\hat{\mathbf{u}}_{(N-1,0)}$, $\hat{\mathbf{f}}_{(N-1,0)}$ and $\hat{\mathbf{h}}_{(N,0)}$.

Based on the considerations expressed above, the corresponding sequential optimization of the values of $\hat{f}_{((l,j,k),0)}$ and $\hat{h}_{((l,j,k),0)}$ is presented in Algorithm 3.5.

Algorithm 3.5.

1. First, solve the LP problem

$$\begin{aligned} & \underset{(\hat{\boldsymbol{\eta}}_0, \gamma)}{\text{minimize}} \quad \gamma \\ & \text{subject to} \quad \hat{f}_{((l,1,1),0)} \leq \gamma, \quad l = 1, \dots, n_c, \\ & \quad \quad \quad (3.23), (3.33) \text{ and } (3.34) \end{aligned} \tag{3.35}$$

Set $\hat{f}_{((l,1,1),0)}^1 = \hat{f}_{((l,1,1),0)}^*$ for $l = 1, \dots, n_c$, where the super-script $*$ represents the optimal solution of the optimization.

2. For $s = 2, \dots, N - 1$, solve the LP problems

$$\begin{aligned} & \underset{(\hat{\boldsymbol{\eta}}_0, \gamma)}{\text{minimize}} \quad \gamma \\ & \text{subject to} \quad \sum_{j=1}^s \hat{f}_{((l,j,s),0)} \leq \gamma, \quad l = 1, \dots, n_c, \\ & \quad \quad \quad \sum_{j=1}^{k-1} \hat{f}_{((l,j,k-1),0)} = \sum_{j=1}^{k-1} \hat{f}_{((l,j,k-1),0)}^{k-1}, \quad l = 1, \dots, n_c, \quad k = 2, \dots, s \\ & \quad \quad \quad (3.23), (3.33) \text{ and } (3.34) \end{aligned} \tag{3.36}$$

Set $\hat{f}_{((l,j,s),0)}^s = \hat{f}_{((l,j,s),0)}^*$ for $l = 1, \dots, n_c$, $j = 1, \dots, s$.

3. Solve the LP problem

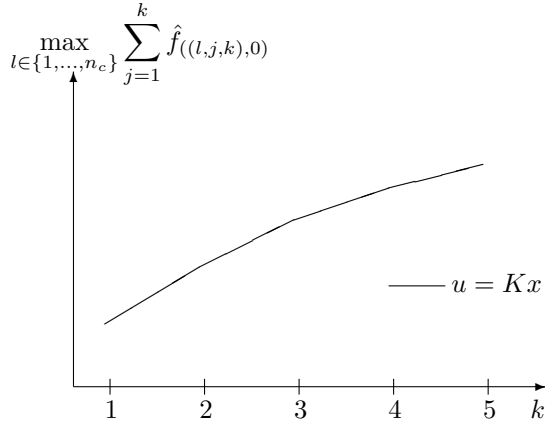
$$\begin{aligned}
 & \underset{(\hat{\mathbf{h}}_0, \gamma)}{\text{minimize}} \quad \gamma \\
 & \text{subject to} \quad \sum_{j=1}^N \hat{h}_{((l,j,N),0)} \leq \gamma, \quad l = 1, \dots, n_f, \\
 & \quad \quad \quad \sum_{j=1}^{k-1} \hat{f}_{((l,j,k-1),0)} = \sum_{j=1}^{k-1} \hat{f}_{((l,j,k-1),0)}^{k-1}, \quad l = 1, \dots, n_c, \quad k = 2, \dots, N \\
 & \quad \quad \quad (3.23), (3.33) \text{ and } (3.34)
 \end{aligned} \tag{3.37}$$

Finally, the desired collections of the optimized j^{th} partial extreme state and control sequences $\hat{\mathbf{x}}_{(N,0)}$ and $\hat{\mathbf{u}}_{(N-1,0)}$ are obtained from the solution of (3.37), and the associated scalar variables $\hat{\mathbf{f}}_{(N-1,0)}$ and $\hat{\mathbf{h}}_{(N,0)}$, are obtained from (3.31).

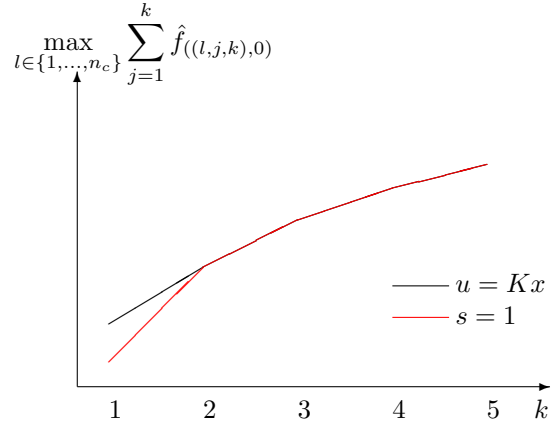
A graphical representation of Algorithm 3.5 is depicted in Figure 3.4, which shows how the values of the worst case of $\sum_{j=1}^s \hat{f}_{((l,j,s),0)}$ are sequentially optimized for $s = 1$, $s = 2$ and so on. As imposed by (3.33), the values of $\sum_{j=1}^k \hat{f}_{((l,j,k),0)}$ are never greater than those associated with the state feedback $u = Kx$, $\sum_{j=1}^k \bar{f}_{(l,j,k)}$.

Remark 3.6. A relevant aspect of having a sequential procedure is that once the slacks associated to a certain prediction instant are optimized, these should not be made worse in the future. This situation would clearly render the sequential procedure useless. An example is depicted in Figure 3.5. Here the slacks for prediction instants $k = 1, 2$ have already been optimized at steps $s = 1, 2$, but these are made larger when optimizing the slacks associated to prediction instant $k = 3$ at step $s = 3$. This situation is avoided by the constraints $\sum_{j=1}^{k-1} \hat{f}_{((l,j,k-1),0)} = \sum_{j=1}^{k-1} \hat{f}_{((l,j,k-1),0)}^{k-1}$, $l = 1, \dots, n_c$, $k = 2, \dots, s$, which guarantee that once $\sum_{j=1}^s \hat{f}_{((l,j,s),0)}$ is optimized, it remains the same for future iterations.

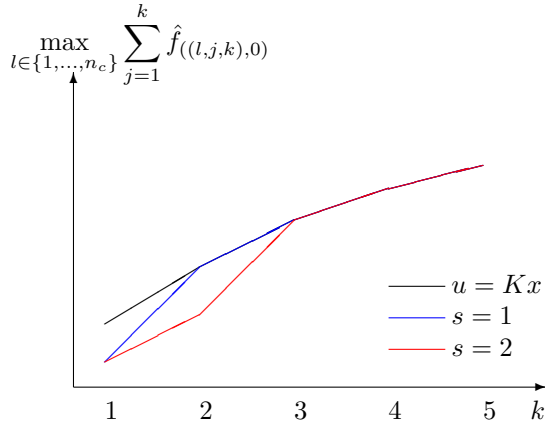
Remark 3.7. Although the associated scalar variables $\hat{\mathbf{f}}_{(N,0)}$ and $\hat{\mathbf{h}}_{(N,0)}$ could in practice also be obtained from the solution of (3.37), it is actually preferred to obtain these



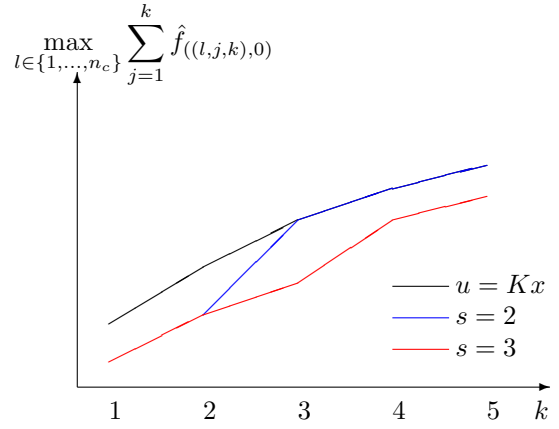
(a) Worst admissible values of slacks



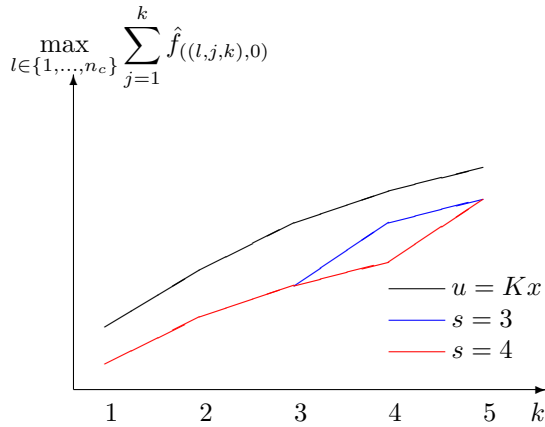
(b) Algorithm 3.5 - Step 1



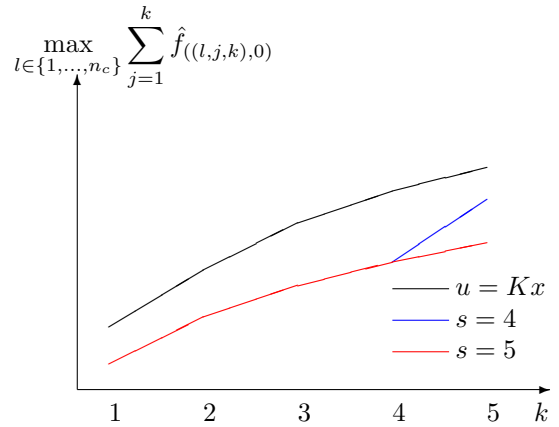
(c) Algorithm 3.5 - Step 2, $s = 2$



(d) Algorithm 3.5 - Step 2, $s = 3$



(e) Algorithm 3.5 - Step 2, $s = 4$



(f) Algorithm 3.5 - Step 2, $s = 5$

Figure 3.4: Graphical representation of Algorithm 3.5

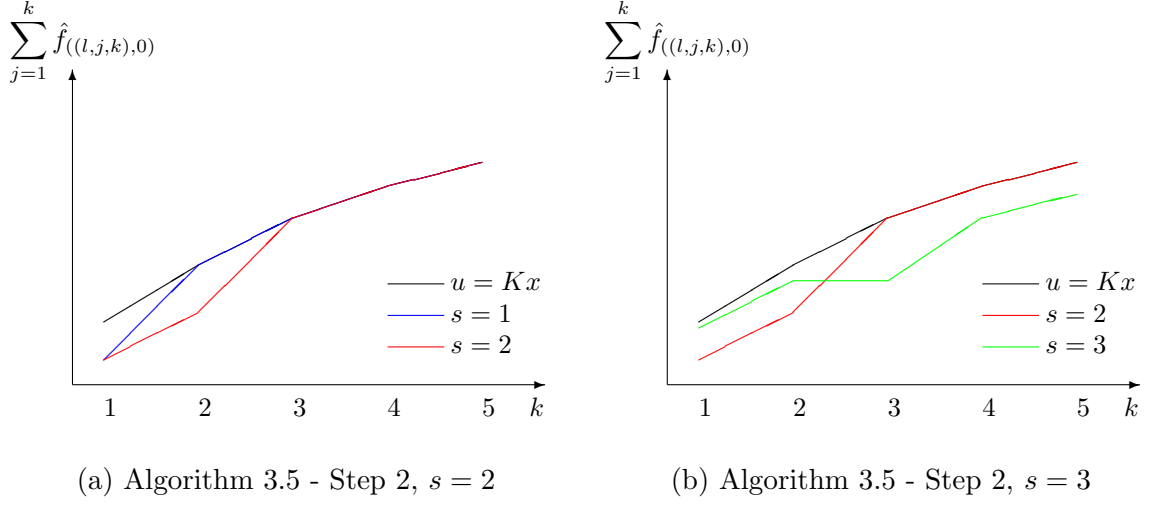


Figure 3.5: Undesirable situation in Algorithm 3.5 avoided by the use of the constraint $\sum_{j=1}^{k-1} \hat{f}_{((l,j,k-1),0)} = \sum_{j=1}^{k-1} \hat{f}_{((l,j,k-1),0)}^{k-1}$, $l = 1, \dots, n_c$, $k = 2, \dots, s$

by evaluating (3.31) with the solutions $\hat{\mathbf{x}}_{(N,0)}$ and $\hat{\mathbf{u}}_{(N,0)}$ of (3.37). The reason is as follows. $\hat{\mathbf{x}}_{(N,0)}$, $\hat{\mathbf{u}}_{(N,0)}$, $\hat{\mathbf{f}}_{(N,0)}$ and $\hat{\mathbf{h}}_{(N,0)}$ are obtained from solving LPs that attempt to minimize the worst case over $l = 1, \dots, n_c$ of the sum of the slacks for s -step ahead predictions, which means that only the worst case for l will be guaranteed to satisfy (3.31); the optimal solutions of LPs are not necessarily unique, and it follows from (3.34) that the slacks $\hat{f}_{((l,j,k),0)}$ and $\hat{h}_{((l,j,N),0)}$ that are not associated to this worst case can take values in between this worst case and what would be obtained by (3.31). As a result, the values of $\hat{\mathbf{f}}_{(N-1,0)}$ and $\hat{\mathbf{h}}_{(N,0)}$ are not necessarily as small as they could be, and then better (smaller) values of $\hat{\mathbf{f}}_{(N-1,0)}$ and $\hat{\mathbf{h}}_{(N,0)}$ can be obtained by evaluating (3.31) with the solutions $\hat{\mathbf{f}}_{(N-1,0)}$ and $\hat{\mathbf{h}}_{(N,0)}$ of (3.37).

3.3.2 Simplified PTMPC: Online implementation and relevant properties

The preceding analysis allows for the online optimization of the 0^{th} partial state and control sequences satisfying (3.6), which is described next.

Assumption 3.4. $N \geq 1$.

Remark 3.8. As stated in Remark 3.1, the offline design is only used for $N \geq 2$. From here on $N = 1$ can also be chosen, but under the assumption that the offline design has been skipped. In this case only the slacks $\hat{h}_{(l,j,N),t}$ are needed, and can be found as in (3.31). Using $N = 1$ however, is clearly of little interest since only the input at time t , $u_{(0,0),t}$, is optimized in the online problem.

As in [73], $\mathbf{x}_{(0,N)}$ and $\mathbf{u}_{(0,N-1)}$ are re-parameterized via

$$\begin{aligned} x_{(0,k)} &= z_{(0,k)} + \Phi^k(x_t - z_{(0,0)}), \quad k = 0, \dots, N, \\ u_{(0,k)} &= v_{(0,k)} + K\Phi^k(x_t - z_{(0,0)}), \quad k = 0, \dots, N-1, \end{aligned} \quad (3.38)$$

where $\mathbf{z}_{(0,N)} = \{z_{(0,0)}, z_{(0,1)}, \dots, z_{(0,N)}\}$ and $\mathbf{v}_{(0,N-1)} = \{v_{(0,0)}, v_{(0,1)}, \dots, v_{(0,N-1)}\}$ comprise the collection of decision variables $\mathbf{d}_{(N,t)}$ of the predicted control policy, and satisfy the dynamics

$$z_{(0,k+1)} = Az_{(0,k)} + Bv_{(0,k)}, \quad k = 0, \dots, N-1. \quad (3.39)$$

$\mathbf{z}_{(0,N)}$ and $\mathbf{v}_{(0,N-1)}$ must be such that (3.29) is satisfied, which, and as in (3.30), can be cast as linear constraints on $\mathbf{z}_{(0,N)}$ and $\mathbf{v}_{(0,N-1)}$, for $t = 0, 1, \dots$

$$F_l(z_{(0,0)} + (x_t - z_{(0,0)})) + G_l(v_{(0,0)} + K(x_t - z_{(0,0)})) \leq 1, \quad l = 1, \dots, n_c, \quad (3.40a)$$

$$\begin{aligned} &F_l(z_{(0,k)} + \Phi^k(x_t - z_{(0,0)})) + \\ &G_l(v_{(0,k)} + K\Phi^k(x_t - z_{(0,0)})) + \sum_{j=1}^k \hat{f}_{(l,j,k),t} \leq 1, \quad l = 1, \dots, n_c, \quad k = 1, \dots, N-1 \end{aligned} \quad (3.40b)$$

$$\begin{aligned} &H_l(z_{(0,N)} + \Phi^N(x_t - z_{(0,0)})) + \\ &G_l(v_{(0,N)} + K\Phi^N(x_t - z_{(0,0)})) + \sum_{j=1}^N \hat{h}_{(l,j,N),t} \leq 1, \quad l = 1, \dots, n_f, \end{aligned} \quad (3.40c)$$

Additionally, in order to guarantee stability it is required that,

$$H_l(x_t - z_{(0,0)}) \leq 1. \quad (3.41)$$

This constraint ensures that $\mathbf{z}_{(0,N)}$ and $\mathbf{v}_{(0,N-1)}$ are the necessary deviations of $\mathbf{x}_{(0,N)}$ and $\mathbf{u}_{(0,N-1)}$ from the dynamics generated by the state feedback $u = Kx$ for an initial state $x_t - z_{(0,0)}$ that is inside $\mathbb{X}_f$, that is necessary in order to satisfy (3.29).

The predicted cost to be optimized online, unlike that of [73], only considers the 0^{th} partial sequences (since the j^{th} partial sequences are fixed, as a function of t , from the offline computations), and is defined as

$$V_N(\mathbf{d}_{(N,t)}) = \sum_{k=0}^{N-1} \ell(z_{(0,k)}, v_{(0,k)}) + V_f(z_{(0,N)}), \quad (3.42)$$

where $\mathbf{d}_{(N,t)}$ is the vector of all the decision variables in the online optimization, $\mathbf{z}_{(0,N)}$ and $\mathbf{v}_{(0,N-1)}$.

Assumption 3.5. *The stage and terminal cost functions $\ell(\cdot, \cdot)$ and $V_f(\cdot)$ are defined for all $z \in \mathbb{R}^n$ and $v \in \mathbb{R}^{n_u}$ by*

$$\ell(z, v) = |z|_{\mathcal{Q}} + |v|_{\mathcal{R}}, \text{ and } V_f = |z|_{\mathcal{P}} \quad (3.43)$$

where $\mathcal{Q}$, $\mathcal{R}$ and $\mathcal{P}$ are symmetric polytopic sets that contain the origin in their interiors, and are defined by

$$\begin{aligned} \mathcal{Q} &= \{z : Qz \leq \mathbf{1}\}, \\ \mathcal{R} &= \{v : Rv \leq \mathbf{1}\}, \\ \mathcal{P} &= \{z : Pz \leq \mathbf{1}\}, \end{aligned} \quad (3.44)$$

where $Q \in \mathbb{R}^{s_1 \times n}$, $R \in \mathbb{R}^{s_2 \times n_u}$ and $P \in \mathbb{R}^{s_3 \times n}$ satisfy for all $z \in \mathcal{P}$

$$|\Phi z|_{\mathcal{P}} - |z|_{\mathcal{P}} \leq -(|z|_{\mathcal{Q}} + |Kz|_{\mathcal{R}}) \quad (3.45)$$

For purposes of analysis only however, consideration is given to a cost function like that of (3.22) so that the partial costs associated with the j^{th} partial extreme state and control sequences, albeit fixed (as a function of time t), are included.

$$V_{(N,t)}(\mathbf{d}_{(N,t)}) = V_{((0,N),t)}(\mathbf{d}_{(N,t)}) + \sum_{j=1}^N V_{((j,N),t)}(\hat{\boldsymbol{\eta}}_t), \quad (3.46)$$

where

$$\begin{aligned} V_{(0,N)}(\mathbf{d}_{(N,t)}) &= \sum_{k=0}^{N-1} \ell(z_{(0,k)}, v_{(0,k)}) + V_f(z_{(0,N)}), \\ V_{(j,N)}(\hat{\boldsymbol{\eta}}_t) &= \sum_{i=1}^q V_{(i,j,N)}(\hat{\boldsymbol{\eta}}_t), \quad j = 1, \dots, N-1, \\ V_{(N,N)}(\hat{\boldsymbol{\eta}}_t) &= \sum_{i=1}^q V_f(\hat{x}_{((i,N,N),t)} - \tilde{w}_i) = 0, \\ V_{(i,j,N)}(\hat{\boldsymbol{\eta}}_t) &= \sum_{k=j}^{N-1} \ell(\hat{x}_{((i,j,k),t)} - \Phi^{k-j} \tilde{w}_i, \hat{u}_{((i,j,k),t)} - K \Phi^{k-j} \tilde{w}_i) + \\ &\quad V_f(\hat{x}_{((i,N,N),t)} - \Phi^{N-j} \tilde{w}_i), \quad i = 1, \dots, q, j = 1, \dots, N-1 \end{aligned}$$

where $\hat{\boldsymbol{\eta}}_t$ is the collection of $\hat{\mathbf{x}}_{((i,j,N),t)}$, $\hat{\mathbf{u}}_{((i,j,N-1),t)}$, $\hat{\mathbf{f}}_{((i,j,N-1),t)}$ and $\hat{\mathbf{h}}_{((i,j,N),t)}$. Obviously, the only difference between the costs of (3.42) and (3.46) lies in the fixed (as a function of time t) terms $\sum_{j=1}^N V_{((j,N),t)}(\hat{\boldsymbol{\eta}}_t)$. Thus the optimization of (3.42) is equivalent to that of (3.46) since the optimal arguments obtained with both cost functions are the same.

The set $\mathcal{D}_{(N,t)}(x)$ of admissible decision variables $\mathbf{d}_{(N,t)}$ is given by

$$\mathcal{D}_{(N,t)}(x) = \{\mathbf{d}_{(N,t)} : \text{conditions (3.39), (3.40) and (3.41) hold}\}. \quad (3.47)$$

The feasibility domain at instant t , $\mathcal{X}_{(N,t)}$, is given by

$$\mathcal{X}_{(N,t)} = \{x : \mathcal{D}_{(N,t)}(x) \neq \emptyset\}. \quad (3.48)$$

$\mathcal{X}_N = \mathcal{X}_{(N,0)}$ is the domain of attraction of the simplified PTMPC strategy.

With the above constructions in mind, the online procedure of the simplified PTMPC reduces to solving the time varying optimal control problem $\mathbb{P}_{(N,t)}(x_t)$

$$\begin{aligned} V_{(N,t)}^0(x_t) &= \min_{\mathbf{d}_{(N,t)}} \{V_{(N,t)}(\mathbf{d}_{(N,t)}) : \mathbf{d}_{(N,t)} \in \mathcal{D}_{(N,t)}(x_t)\}, \\ \mathbf{d}_{(N,t)}^* &= \arg \min_{\mathbf{d}_{(N,t)}} \{V_{(N,t)}(\mathbf{d}_{(N,t)}) : \mathbf{d}_{(N,t)} \in \mathcal{D}_{(N,t)}(x_t)\}. \end{aligned} \quad (3.49)$$

Note that, as will be seen below, $\mathbb{P}(x_t)_{(N,t)}$ is an LP, the solution of which may not be unique, and in general $\mathbf{d}_{(N,t)}^*(x_t)$ is a set-valued function of x_t .

The simplified PTMPC law is a time-varying function $\kappa_{(N,t)}^0(x_t)$, which for $t = 0, 1, \dots$ and $x_t \in \mathcal{X}_{(N,t)}$ is given by

$$\kappa_{(N,t)}^0(x_t) = v_{((0,0),t)}^0(x_t) + K(x - z_{((0,0),t)}^0(x_t)), \quad (3.50)$$

where $v_{((0,0),t)}^0(x_t)$ and $z_{((0,0),t)}^0(x_t)$ denote a continuous selection from $v_{((0,0),t)}^*(x_t)$ and $z_{((0,0),t)}^*(x_t)$ since the latter two are not necessarily singled valued (as they are part of $\mathbf{d}_{(N,t)}^*(x_t)$). Note that from this point on, the dependence on t will be made explicit for the optimal solutions (and their selections) in the form of $(\cdot)_{((0,k),t)}^*$ (and $(\cdot)_{((0,k),t)}^0$), in order to highlight that they are solutions of a time-varying optimal control problem. The control input at time t is then given by $u_t := \kappa_{(N,t)}^0(x_t)$.

As shown in [73], $\mathbb{P}_{(N,t)}(x_t)$ can be re-cast as an LP. Since the minimization of $V_N(\mathbf{d}_{(N,t)})$ (which is equivalent to the minimization of $V_{(N,t)}(\mathbf{d}_{(N,t)})$) is a sum of gauge functions of polytopic sets, it can be performed equivalently by minimizing the sum of the slack variables $\boldsymbol{\alpha}_{(0,N)} = \{\alpha_{(0,0)}, \alpha_{(0,1)}, \dots, \alpha_{(0,N)}\}$, $\boldsymbol{\beta}_{(0,N-1)} =$

$\{\beta_{(0,0)}, \beta_{(0,1)}, \dots, \beta_{(0,N-1)}\}$, with the additional linear inequalities,

$$\begin{aligned} \mathcal{Q}_l^T z_{(0,k)} &\leq \alpha_{(0,k)}, \quad k = 0, \dots, N-1, \quad l = 1, \dots, s_1, \\ \mathcal{R}_l^T v_{(0,k)} &\leq \beta_{(0,k)}, \quad k = 0, \dots, N-1, \quad l = 1, \dots, s_2, \\ \mathcal{P}_l^T z_{(0,N)} &\leq \alpha_{(0,N)}, \quad l = 1, \dots, s_3. \end{aligned} \tag{3.51}$$

The equivalent LP form of $\mathbb{P}_{(N,t)}(x_t)$ is then given by

$$\begin{aligned} &\underset{\mathbf{d}_{(O,t)}}{\text{minimize}} && \sum_{k=0}^{N-1} \alpha_{(0,k)} + \beta_{(0,k)} + \alpha_{(0,N)} \\ &\text{subject to} && \mathcal{Q}_l^T z_{(0,k)} \leq \alpha_{(0,k)}, \quad k = 0, \dots, N-1, \quad l = 1, \dots, s_1, \\ & && \mathcal{R}_l^T v_{(0,k)} \leq \beta_{(0,k)}, \quad k = 0, \dots, N-1, \quad l = 1, \dots, s_2, \\ & && \mathcal{P}_l^T z_{(0,N)} \leq \alpha_{(0,N)}, \quad l = 1, \dots, s_3, \\ & && z_{(0,k+1)} = Az_{(0,k)} + Bv_{(0,k)}, \quad k = 0, \dots, N-1, \\ & && F_l(z_{(0,0)} + (x_t - z_{(0,0)})) + G_l(v_{(0,0)} + K(x_t - z_{(0,0)})) \leq 1, \\ & && \quad \quad \quad l = 1, \dots, n_c, \\ & && F_l(z_{(0,k)} + \Phi^k(x - z_{(0,0)})) + G_l(v_{(0,k)} + K\Phi^k(x - z_{(0,0)})) + \\ & && \sum_{j=1}^k \hat{f}_{((l,j,k),t)} \leq 1, \quad l = 1, \dots, n_c, \quad k = 0, \dots, N-1, \\ & && H_l(z_{(0,N)} + \Phi^N(x - z_{(0,0)})) + G_l(v_{(0,N)} + K\Phi^N(x - z_{(0,0)})) + \\ & && \sum_{j=1}^N \hat{h}_{((l,j,N),t)} \leq 1, \quad l = 1, \dots, n_f, \\ & && H_l(x_t - z_{(0,0)}) \leq 1, \quad l = 0, \dots, n_f. \end{aligned} \tag{3.52}$$

where the overall decision variable $\mathbf{d}_{(O,t)}$ comprises $\mathbf{z}_{(0,N)}$, $\mathbf{v}_{(0,N-1)}$, $\boldsymbol{\alpha}_{(0,N)}$ and $\boldsymbol{\beta}_{(0,N-1)}$.

It can be seen now that this is an LP with $N_{var} = N(n + n_u + 2) + n$ variables and $N_{cons} = N(s_1 + Ns_2 + n_c + n) + s_3 + 2n_f$ constraints (of which Nn are equality constraints and the rest are inequality constraints).

Remark 3.9. In view of the observations of Remark 3.2 and the construction of

this section, the sequence of time varying optimal control problems $\mathbb{P}_{(N,t)}$, $t = 0, 1, \dots$ converges in at most N time steps to the time-invariant optimal control problem $\mathbb{P}_{(N,N)}$. More precisely, the construction guarantees that, for $t = 0, 1, \dots$

$$\begin{aligned} \mathcal{X}_{(N,t+N)} &= \mathcal{X}_{(N,N)}, \\ \forall x \in \mathcal{X}_{(N,t+N)}, \quad V_{(N,t+N)}^0(x) &= V_{(N,N)}^0(x) \text{ and } \mathbf{d}_{(N,t+N)}^0(x) = \mathbf{d}_{(N,N)}^0(x) \end{aligned} \quad (3.53)$$

As a result, the time-varying simplified PTMPC law converges to a time-invariant simplified PTMPC law $\kappa_{(N,N)}^0(\cdot)$.

Some system theoretic properties of the simplified PTMPC strategy are analysed next.

Theorem 3.10. *The online application of $\kappa_{(N,t)}^0(x_t)$ to system (3.1) is recursively feasible and guarantees the satisfaction of (3.3) if $x_t \in \mathcal{X}_{(N,t)}$.*

Proof. For all $x_t \in \mathcal{X}_{(N,t)}$, $\mathbf{d}_{(N,t)}^0(x_t)$ is an optimal solution of $\mathbb{P}_{(N,t)}(x_t)$. Thus (3.30) is satisfied, where

$$\begin{aligned} z_{((0,k),t)}^0(x_t) &= z_{((0,k),t)}^0(x_t) + \Phi^k(x_t - z_{((0,0),t)}^0(x_t)), \\ u_{((0,k),t)}^0(x_t) &= v_{((0,k),t)}^0(x_t) + K\Phi^k(x_t - z_{((0,0),t)}^0(x_t)). \end{aligned}$$

We want to prove that the optimization is feasible at the next instant, i.e., that $x_{t+1} \in \mathcal{X}_{(N,t)}$. To this end let the candidate 0^{th} partial state and control sequences at time $t + 1$, $\tilde{\mathbf{z}}_{((0,N),t+1)}(x_{t+1}) = \{\tilde{z}_{((0,0),t+1)}(x_{t+1}), \tilde{z}_{((0,1),t+1)}(x_{t+1}), \dots, \tilde{z}_{((0,N),t+1)}(x_{t+1})\}$ and $\tilde{\mathbf{v}}_{((0,N-1),t+1)}(x_{t+1}) = \{\tilde{v}_{((0,0),t+1)}(x_{t+1}), \tilde{v}_{((0,1),t+1)}(x_{t+1}), \dots, \tilde{v}_{((0,N),t+1)}(x_{t+1})\}$, be given by the receded 0^{th} partial state and control sequences obtained at time t and of the 1^{st} partial state and control tubes obtained at time t which, since at time $t + 1$ the effect of w_t is known, have collapsed to a sequence of singletons. Then, for

$$k = 0, \dots, N - 1$$

$$\tilde{z}_{((0,k),t+1)}(x_{t+1}) = z_{((0,k+1),t)}^0(x_t) + \sum_{i=1}^q \lambda_{(i,1)}^*(w_t)(\hat{x}_{((i,1,k+1),t)} - \Phi^k \tilde{w}_i), \quad (3.54a)$$

$$\tilde{z}_{((0,N),t+1)}(x_{t+1}) = \Phi \tilde{z}_{((0,N-1),t+1)}(x_{t+1}), \quad (3.54b)$$

and, for $k = 0, \dots, N - 2$

$$\tilde{v}_{((0,k),t+1)}(x_{t+1}) = v_{((0,k+1),t)}^0(x_t) + \sum_{i=1}^q \lambda_{(i,1)}^*(w_t)(\hat{u}_{((i,1,k+1),t)} - K \Phi^k \tilde{w}_i), \quad (3.55a)$$

$$\tilde{v}_{((0,N-1),t+1)}(x_{t+1}) = K \tilde{z}_{((0,N-1),t+1)}(x_{t+1}), \quad (3.55b)$$

where $\lambda_{(i,1)}^*(w_t)$ are the least squares (or any other convenient selection criterion) convex interpolation parameters that satisfy $w_t = \sum_{i=1}^q \lambda_{(i,1)}^*(w_t) \tilde{w}_i$. These define the candidate 0^{th} partial state and control sequences at time $t + 1$

$$\tilde{x}_{((0,k),t+1)}(x_{t+1}) = \tilde{z}_{((0,k),t+1)}(x_{t+1}) + \Phi^k(x_{t+1} - \tilde{z}_{((0,0),t+1)}(x_{t+1})), \quad k = 0, \dots, N \quad (3.56a)$$

$$\tilde{u}_{((0,k),t+1)}(x_{t+1}) = \tilde{v}_{((0,k),t+1)}(x_{t+1}) + K \Phi^k(x_{t+1} - \tilde{z}_{((0,0),t+1)}(x_{t+1})), \quad k = 0, \dots, N - 1, \quad (3.56b)$$

thus defining the candidate 0^{th} mixed state and control partial cross sections $\tilde{Y}_{((0,k),t+1)} = (\tilde{x}_{((0,k),t+1)}(x_{t+1}), \tilde{u}_{((0,k),t+1)}(x_{t+1}))$, $k = 0, \dots, N - 1$, and the candidate terminal partial cross sections $\tilde{X}_{((0,N),t+1)} = \tilde{x}_{((0,N),t+1)}(x_{t+1})$.

Feasibility of (3.30) at time t (and the equivalence of (3.30) and (3.29)), the definition of the slacks in (3.31), the update dynamics of (3.27) and (3.28), the invariance of $\mathbb{X}_f$ under $u = Kx$ (which is the terminal control law), and (3.54), (3.55) and (3.56),

imply that

$$\begin{aligned}
 \tilde{Y}_{((0,0),t+1)} &\subseteq \mathbb{Y}, \\
 \tilde{Y}_{((0,k),t+1)} \bigoplus_{j=1}^k \hat{Y}_{((j,k),t+1)} &\subseteq \mathbb{Y}, \quad k = 0, \dots, N-1 \\
 \tilde{X}_{((0,N),t+1)} \bigoplus_{j=1}^N \hat{X}_{((j,N),t+1)} &\subseteq \mathbb{X}_f.
 \end{aligned} \tag{3.57}$$

This implies that the candidate decision variable $\tilde{\mathbf{d}}_{(N,t+1)}(x_{t+1})$, that contains the sequences $\tilde{\mathbf{z}}_{((0,N),t+1)}(x_{t+1})$ and $\tilde{\mathbf{v}}_{((0,N-1),t+1)}(x_{t+1})$, is a feasible solution of $\mathbb{P}_{(N,t+1)}(x_{t+1})$.

Thus $x_{t+1} \in \tilde{\mathcal{X}}_{N,t+1}$, which proves the desired result.

Satisfaction of (3.3) follows from use of (3.40) for $k = 0$. $\square$

Proposition 3.11. *For the optimal control problem $\mathbb{P}_{(N,t)}(x)$:*

(i) *There exists a piecewise affine and continuous function $\mathbf{d}_{(N,t)}^0(\cdot) : \mathcal{X}_{(N,t)} \rightarrow \mathbb{R}^{N_{\mathbf{d}}}$ such that for all $x \in \mathcal{X}_{(N,t)}$, $\mathbf{d}_{(N,t)}^0(x) \in \mathbf{d}_{(N,t)}^*(x)$ ($\mathbf{d}_{(N,t)}^0(x)$ being a selection of $\mathbf{d}_{(N,t)}^*(x)$).*

(ii) *For all $x \in \mathbb{X}_f$, $t \geq N$, $\mathbf{d}_{(N,t)}^*(x)$ is singleton given by $z_{((0,k),t)}^*(x) = 0$ and $v_{((0,k),t)}^*(x) = 0$.*

(iii) *For all $x \in \mathbb{X}_f$, $t \geq N$, $\kappa_{(N,t)}^0(x) = Kx$.*

The value function $V_{(N,t)}^0(x)$, has the following properties:

(iv) *$V_{(N,t)}^0(\cdot)$ is a convex, continuous and piecewise affine function.*

(v) *$V_{(N,t)}^0(\cdot) > 0$, for all $x \in \mathcal{X}_{(N,t)} \setminus \mathbb{X}_f$, for $t = 0, 1, \dots$*

(vi) *$V_{(N,t)}^0(\cdot) = 0$, for all $x \in \mathbb{X}_f$, $t \geq N$.*

Proof. (i) The optimal control problem $\mathbb{P}_{(N,t)}(\cdot)$ is a linear parametric program in $\mathcal{X}_{(N,t)}$. Then by standard results in parametric optimization problems [5] the stated property holds.

(ii) For $t \geq N$ the j^{th} partial extreme state and control sequences $\hat{\mathbf{x}}_{((i,j,N),t)}$, $\hat{\mathbf{x}}_{((i,j,N-1),t)}$ and the associated slacks $\hat{\mathbf{f}}_{((i,j,N-1),t)}$, $\hat{\mathbf{h}}_{((i,j,N),t)}$ are obtained by setting $\hat{u}_{((i,j,k),t)} = K\hat{x}_{((i,j,k),t)}$. Thus if $x \in \mathbb{X}_f$, the solution for the 0^{th} partial state and control

sequences given by the evolution with the linear controller, such that $u_{(0,k)} = Kx_{(0,k)}$, is feasible and optimal so that $z_{((0,k),t)}^*(x) = 0$, $v_{((0,k),t)}^*(x) = 0$ (and $V_{(N,t)}^0(x) = 0$).

(iii) This follows trivially from (ii).

(iv) The function $V_{(N,t)}^0(x)$ is the value function of a parametric linear program so that it is a convex, piecewise affine and continuous function in x (by standard results in parametric optimization problems [5]).

(v) $x - z_{(0,0)} \in \mathbb{X}_f$ is a constraint in the optimization problem. Then, if $x \notin \mathbb{X}_f$, it is necessary that $z_{((0,0),t)}^* \neq 0$ in order to satisfy that constraint, hence $V_{(N,t)}^0(x) > 0$.

(vi) For $t \geq N$ the j^{th} partial extreme state and control sequences $\hat{\mathbf{x}}_{((i,j,N),t)}$, $\hat{\mathbf{x}}_{((i,j,N-1),t)}$ (and the associated slacks $\hat{\mathbf{f}}_{((i,j,N-1),t)}$, $\hat{\mathbf{h}}_{((i,j,N),t)}$) are given by $\hat{u}_{((i,j,k),t)} = K\hat{x}_{((i,j,k),t)}$ so that $V_{(N,t)}^0(x) = V_{(N,N)}^0(x) = \sum_{k=0}^{N-1} |z_{((0,k),t)}^0(x)|_{\mathcal{Q}} + |v_{((0,k),t)}^0(x)|_{\mathcal{R}} + |z_{((0,N),t)}^0(x)|_{\mathcal{P}}$, but from (ii), $z_{((0,k),t)}^*(x_t) = 0$ and $v_{((0,k),t)}^*(x) = 0$, which yields the desired result. $\square$

Proposition 3.12. *For all x_0 in $\mathcal{X}_{(N,0)}$ and for $t = 0, 1, \dots$ the value function $V_{(N,t)}^0(x_t)$ satisfies*

$$V_{(N,t+1)}^0(x_{t+1}) \leq V_{(N,t)}^0(x_t) - (|z_{((0,0),t)}^0(x_t)|_{\mathcal{Q}} + |v_{((0,0),t)}^0(x_t)|_{\mathcal{R}}) \quad (3.58)$$

Proof. Consider the optimal solution of $\mathbb{P}_{(N,t)}(x_t)$, $\mathbf{d}_{(N,t)}^0(x_t)$, that contains the optimal sequences $\mathbf{z}_{((0,N),t)}^0(x_t)$ and $\mathbf{v}_{((0,N-1),t)}^0(x_t)$, and the candidate decision variable $\tilde{\mathbf{d}}_{(N,t+1)}(x_{t+1})$ containing the sequences $\tilde{\mathbf{z}}_{((0,N),t+1)}(x_{t+1})$ and $\tilde{\mathbf{z}}_{((0,N),t+1)}(x_{t+1})$ defined in (3.54) and (3.55).

From (3.23) and the Lyapunov condition of (3.45), the following inequality holds for $V_{((i,j,N),t)}(\cdot)$,

$$V_{((i,j,N),t+1)}(\hat{\boldsymbol{\eta}}_{t+1}) \leq V_{((i,j+1,N),t)}(\hat{\boldsymbol{\eta}}_t), \quad (3.59)$$

which implies

$$V_{((j,N),t+1)}(\hat{\boldsymbol{\eta}}_{t+1}) \leq V_{((j+1,N),t)}(\hat{\boldsymbol{\eta}}_t). \quad (3.60)$$

In a similar way, using also the convexity and sub-additivity of the associated gauge functions and the Lyapunov condition of (3.45), $\tilde{\mathbf{d}}_{(N,t+1)}(x_{t+1})$ and $\mathbf{d}_{(N,t)}^0(x_t)$ it follows that

$$\begin{aligned} V_{((0,N),t+1)}(\tilde{\mathbf{d}}_{(N,t+1)}(x_{t+1})) &\leq V_{((1,N),t)}(\hat{\boldsymbol{\eta}}_t) + V_{((0,N),t)}(\mathbf{d}_{(N,t)}^0(x_t)) \\ &\quad - \ell(z_{((0,0),t)}^0(x_t), v_{((0,0),t)}^0(x_t)). \end{aligned} \quad (3.61)$$

(3.60) and (3.61), in conjunction with the fact that $V_{((N,N),t+1)}(\hat{\boldsymbol{\eta}}_{t+1}) = 0$, yield

$$V_{(N,t+1)}^0(x_{t+1}) \leq V_{(N,t)}^0(x_t) - (|z_{((0,0),t)}^0(x_t)|_{\mathcal{Q}} + |v_{((0,0),t)}^0(x_t)|_{\mathcal{R}})$$

which is the desired result. $\square$

Theorem 3.13. *Consider system (3.1) under the control law $\kappa_{(N,t)}^0(x_t)$. There exist $a_N \in [0, 1)$ and $b_N \in (0, \infty)$ such that the value function $V_{(N,t)}^0(x_t)$ satisfies:*

- (i) $V_{(N,t)}^0(x_t) \leq a_N^{t-N} V_{(N,N)}^0(x_N)$, $t \geq N$,
- (ii) $\text{dist}(\mathcal{Q}, x_t, \mathbb{X}_f) \leq b_N a_N^{t-N} \text{dist}(\mathcal{Q}, x_N, \mathbb{X}_f)$, $t \geq N$.

Proof. (i) By construction it is clear that $|z_{((0,0),t)}^0(x_t)|_{\mathcal{Q}} + |v_{((0,0),t)}^0(x_t)|_{\mathcal{R}} < V_N^0(x_t)$. In addition, $V_{(N,t)}^0(x_t) = 0$ and $|z_{((0,0),t)}^0(x_t)|_{\mathcal{Q}} + |v_{((0,0),t)}^0(x_t)|_{\mathcal{R}} = 0$ for $x_t \in \mathbb{X}_f$, $t = N, N+1, \dots$; and $V_{(N,t)}^0(x_t) > 0$ and $|z_{((0,0),t)}^0(x_t)|_{\mathcal{Q}} + |v_{((0,0),t)}^0(x_t)|_{\mathcal{R}} > 0$ for $x_t \in \mathcal{X}_{(N,t)} \setminus \mathbb{X}_f$. Then, based on the continuity of $V_{(N,N)}^0(\cdot)$ and $|z_{((0,0),N)}^0(\cdot)|_{\mathcal{Q}} + |v_{((0,0),N)}^0(\cdot)|_{\mathcal{R}}$, the fact that $V_{(N,t)}^0(\cdot) = V_{(N,N)}^0(\cdot)$ and $\mathbf{d}_{(N,t)}^0(\cdot) = \mathbf{d}_{(N,N)}^0(\cdot)$ for $t \geq N$, and that $\mathbb{X}_f$ and $\mathcal{X}_N$ are compact, there exists a scalar $c_1 \in [1, \infty)$ such that $V_{(N,N)}^0(x_t) \leq c_1(|z_{((0,0),t)}^0(x_t)|_{\mathcal{Q}} + |v_{((0,0),t)}^0(x_t)|_{\mathcal{R}})$, for all $x \in \mathcal{X}_N$. Combining this with (3.58) it follows that for $t \geq N$, $V_{(N,t+1)}^0(x_{t+1}) \leq a_N V_{(N,t)}^0(x_t)$, with $a_N = (c_1 - 1)/c_1 < 1$, which yields the desired property.

(ii) Similarly to (i), it is possible to obtain $|z_{((0,0),t)}^0(x_t)|_{\mathcal{R}} \leq c_1 ((c_1 - 1)/c_1)^{t-N} |z_{((0,0),N)}^0(x_N)|_{\mathcal{R}}$, $t \geq N$. Also, $\text{dist}(\mathcal{Q}, x, \mathbb{X}_f) = \text{dist}(\mathcal{Q}, x - z + z, \mathbb{X}_f)$, and hence, since $x_t - z_{((0,0),t)}^0(x_t) \in \mathbb{X}_f$, it holds that for all $x \in \mathcal{X}_N$, $\text{dist}(\mathcal{Q}, x_t, \mathbb{X}_f) = \text{dist}(\mathcal{Q}, z_{((0,0),t)}^0(x_t) +$

$(x_t - z_{((0,0),t)}^0(x_t), \mathbb{X}_f) \leq \text{dist}(\mathcal{Q}, z_{((0,0),t)}^0(x_t), \mathbb{X}_f) \leq |z_{((0,0),t)}^0(x_t)|_{\mathcal{Q}}$. This, in addition with continuity of $z_{((0,0),N)}^0(x_t)$ and $\text{dist}(\mathcal{Q}, x_t, \mathbb{X}_f)$, the fact that $\mathbf{d}_{(N,t)}^0(\cdot) = \mathbf{d}_{(N,N)}^0(\cdot)$ for $t \geq N$, and compactness of $\mathbb{X}_f$ and $\mathcal{X}_N$, imply that there exists $c_2 \in (0, 1)$ such that $c_2 |z_{((0,0),t)}^0(x_t)|_{\mathcal{Q}} \leq \text{dist}(\mathcal{Q}, x_t, \mathbb{X}_f) \leq |z_{((0,0),t)}^0(x_t)|_{\mathcal{Q}}$, $t \geq N$. Combining this with $|z_{((0,0),t)}^0(x_t)|_{\mathcal{Q}} \leq c_2 ((c_2 - 1)/c_2)^{t-N} |z_{((0,0),t)}^0(x_N)|_{\mathcal{Q}}$, it can be obtained for $t \geq N$ that $\text{dist}(\mathcal{Q}, x_t, \mathbb{X}_f) \leq c_1/c_2 ((c_1 - 1)/c_1)^{t-N} \text{dist}(\mathcal{Q}, x_N, \mathbb{X}_f)$, which is the desired result with $b_N = c_1/c_2$. $\square$

Remark 3.14. Under the control law of $u = Kx$, system (3.1) is robustly exponentially stable with respect to the minimal robust positively invariant set $\Omega_{\infty}(\mathbb{W})$, with the domain of attraction being the maximal robust positively invariant set [3], [44]. Then, combining arguments presented in [73], Proposition 3.11(iii), Theorem 3.13 and classical results [3], [44] imply directly that system (3.1) under the control law of $u = \kappa^0(x)$ is robustly exponentially stable with respect to $\Omega_{\infty}(\mathbb{W})$ where the domain of attraction is $\mathcal{X}_N$.

3.4 Simulation Results

The benefits of the proposed strategy are illustrated with statistics for 50 randomly generated systems controlled by PTMPC, and the proposed strategy, OPTMPC, using different prediction horizons N .

The systems are such that $n = 2$, $n_u = 1$, with A, B being random matrices where each component is selected from an independent uniform distribution between -1.25 and 1.25, but only systems where the spectral radius of A is greater than 0.95 are used. F, G are chosen to represent the constraints $-\underline{1} \leq \tilde{F}x_t \leq \underline{1}$, $\tilde{F} \in \mathbb{R}^{2 \times 2}$, and $-1 \leq u_t \leq 1$, so that $F = [\tilde{F}^T, -\tilde{F}^T, [0, 0]^T, [0, 0]^T]^T$ and $G = [0, 0, 0, 0, -1, 1]^T$, with $\tilde{F}$ being a random matrix where each component is selected from an independent uniform distribution between -0.1 and 0.1. K is chosen to be the LQ optimal gain for the cost matrices $Q = I_{2 \times 2}$ and $R = 0.1$. $\mathbb{W} = \text{conv}(\{[0.1, 0.1]^T, [0.1, -0.1]^T, [-0.1, 0.1]^T,$

$[-0.1, -0.1]^T\}$).

The matrices in the costs to be used in the proposed strategy and PTMPC are given by $\mathcal{Q} = \{x : (1, 0)x \leq 1, (0, 1)x \leq 1, (-1, 0)x \leq 1, (0, -1)x \leq 1\}$, $\mathcal{R} = \{u : u \leq 1, -u \leq 1\}$ and $\mathcal{P} = [\mathcal{P}_1^T \quad -\mathcal{P}_1^T]$, where $\mathcal{P}_1^T \in \mathbb{R}^{s \times n}$, $s = 2$ is computed as in [55] in order to satisfy the Lyapunov condition of (3.45). The simulations have been performed on Matlab 7.11.0 (R2010b) with the Gurobi solver on a machine with Windows 7 and an Intel Core 2 Duo E8400 processor at 3GHz.

The domains of attraction using OPTMPC and PTMPC are computed for each of these systems. The areas that would be obtained with DP, which are the maximum that can be obtained for any control policies, are also included. Table 3.1 shows the mean values (over all the random systems) of the areas of the domains of attraction obtained with different prediction horizons. In order to make the values comparable, these areas are reported relative to the area of $\mathbb{X}_f$, such that $A_N = \text{Area}(\mathcal{X}_N)/\text{Area}(\mathbb{X}_f)$. The table also reports information relevant to the computational implementation, such as the mean values for the number of variables, equalities and inequalities in the online optimization and average computational times for OPTMPC and PTMPC. More precisely, the reported values correspond to the average time it takes the solver to perform the online optimization for a given prediction horizon, where the average is computed over a set initial conditions evenly distributed at the edge of the domain of attraction of OPTMPC and over all random systems.

First note from Table 3.1 that the areas obtained by PTMPC and DP are the same for the considered examples. This is consistent with the conjecture that PTMPC is equivalent to DP in terms of domain of attraction, as introduced in [73].

It can be deduced from Table 3.1 that for small prediction horizons the deterioration in the domain of attraction of OPTMPC is small, but it increases with the prediction horizon. This is expected since there are more variables fixed for greater N . For $N = 1$, OPTMPC and PTMPC are equivalent, and therefore the areas of the

Table 3.1: Domain of Attraction and computational aspects

N		1	2	3	4	5	6	7	8	9	10
A_N	OPTMPC	1.482	3.049	4.499	5.356	6.004	6.688	7.456	8.312	9.348	10.548
	PTMPC	1.482	3.139	4.718	5.710	6.558	7.458	8.530	9.742	11.203	12.953
	DP	1.482	3.139	4.718	5.710	6.558	7.458	8.530	9.742	11.203	12.953
vars	OPTMPC	8	13	18	23	28	33	38	43	48	53
	PTMPC	24.9	72.7	146.6	246.4	372.3	524.1	702.0	905.9	1136	1392
ineqs	OPTMPC	25.7	37.7	49.7	61.7	73.7	85.7	97.7	109.7	121.7	133.7
	PTMPC	63.6	161.4	307.3	501.1	743.0	1033	1371	1757	2190	2672
eqs	OPTMPC	2	4	6	8	10	12	14	16	18	20
	PTMPC	10	28	54	88	130	180	238	304	378	460
time [ms]	OPTMPC	1.70	1.89	2.14	2.36	2.45	2.55	2.66	2.78	2.90	2.97
	PTMPC	1.70	2.60	3.94	6.55	12.17	20.46	36.28	56.69	100.0	147.4

domains of attraction are the same. For $N = 5$, OPTMPC the areas of the domains of attraction are in average 93.65% the size of those obtained with PTMPC, and for $N = 10$ this number goes down to 85.50%.

This loss in domain of attraction is compensated by the fact that the reduction in computational times is increasing with the prediction horizon, which is driven by the fact that the number of constraints and variables grows quadratically with N for PTMPC, but only linearly for OPTMPC. For the case $N = 1$, the computational times are about the same for both methods. For $N = 5$, OPTMPC is on average 5.0 times faster than PTMPC, and for $N = 10$ OPTMPC is on average 49.5 times faster than PTMPC. This fosters the possibility of using considerably longer prediction horizons in OPTMPC. For instance, one might choose to use $N = 10$ OPTMPC, which will give a domain of attraction larger than would be obtained with PTMPC with $N = 8$ but with an online optimization 19 times faster on average.

All the above implies that for many systems, as illustrated by the average case, one can use OPTMPC with large horizons in order to obtain larger domains of attractions at a reduced computational cost (when compared with PTMPC).

This is obviously useful only as long as the obtained domains of attraction are within an acceptable size, or include the desired regions of the state space. For instance, one can find systems for which the offline design will impose such conserva-

tiveness that even for very large N , the domain of attraction of OPTMPC will never be as large as what can be obtained by PTMPC for even a short N . Such an example is provided next for a system that is defined by

$$A = \begin{bmatrix} 1.19 & 0.95 \\ -0.89 & 0.96 \end{bmatrix}, B = \begin{bmatrix} 0.66 \\ 0.74 \end{bmatrix}, F = \begin{bmatrix} 0.083 & 0.114 & -0.083 & -0.114 & 0 & 0 \\ -0.035 & -0.070 & 0.035 & 0.070 & 0 & 0 \end{bmatrix}^T,$$

$G = [0 \ 0 \ 0 \ 0 \ 1.44 \ -1.44]^T$ and $\mathbb{W} = \text{conv}(\{[0.1, 0.1]^T, [0.1, -0.1]^T, [-0.1, 0.1]^T, [-0.1, -0.1]^T\})$. The domains of attraction of this system controlled by PTMPC and OPTMPC are presented in Table 3.2.

Table 3.2: Domain of Attraction and computational aspects

N		1	2	3	4	5	6	7	8	9	10
A_N	OPTMPC	1.302	1.446	1.626	1.830	1.880	1.935	1.958	1.977	1.990	2.002
	PTMPC	1.302	1.747	2.174	2.508	2.707	2.850	2.975	3.073	3.137	3.178
	DP	1.302	1.747	2.174	2.508	2.707	2.850	2.975	3.073	3.137	3.178

Clearly in such cases the benefit of OPTMPC is limited due to the small domains of attraction that one can obtain with it.

The examples considered in this section have $n = 2$. As n increases, the advantages of OPTMPC over PTMPC in terms of online computational efficiency will also increase. As it is discussed in Section 2.2.4.2, the size of the online optimization in PTMPC increases linearly with the product nq , where the number of vertices of $\mathbb{W}$, q , is at least $n + 1$, but can even grow exponentially with n , so the size of the optimization problem grows at least quadratically or even exponentially with n . On the other hand, since the extreme partial sequences are computed offline the size of the online optimization problem in OPTMPC only grows linearly with n . And although the offline computations of OPTMPC have a comparable size to that of PTMPC (only the nominal sequences and the slacks associated with the cost are absent), the fact that these are offline computations without the stringent time constraints that can be present in real applications, makes them much easier to be handled as n increases.

For these reasons, OPTMPC is better suited than PTMPC for larger n .

3.5 Conclusions

The suggested offline simplifications reduce online computational complexity, particularly the growth of the number of variables and constraints of the online computations, from quadratic to linear. This comes at the cost of degrading to a reasonable extent the size of the domain of attraction as compared with PTMPC and DP, when both methods use the same N as OPTMPC.

This degradation is reasonably small in the sense that, for some systems, the prediction horizon of OPTMPC can be increased in order to get larger domains of attraction than PTMPC (which uses a smaller prediction horizon), while incurring a computational demand smaller than that required by PTMPC.

Since the extreme partial sequences are computed offline, the size of the online optimization problem in OPTMPC only grows linearly with n , while it can be exponential, and in the best case quadratic, for PTMPC. Thus OPTMPC is better suited than PTMPC for larger n .

However, the conservativeness imposed by the offline design in some cases can be large enough so that, no matter how large N is chosen, OPTMPC cannot have a domain of attraction as large as PTMPC even for short N . The use of OPTMPC for these systems does not confer any benefit other than the reduction in computational demand. This means that in the sense of both domain of attraction and speed of computation, neither method is better. For some systems, and a given range of acceptable domains of attraction, OPTMPC will be able to give better domains of attraction with faster computations than PTMPC. But for other systems, or when larger domains of attraction are desired, OPTMPC will not be able to accomplish that.

The nature of the offline design means that the online optimization problems need

to be time varying in order to guarantee recursive feasibility, but at execution time N the problem will become time-invariant. The update of the slack variables implies that at every time step there will be more terms that are associated with the control law $u = Kx$ in the control policy. Thus the feasibility domains of the optimizations tend to decrease until the N th instant when the feasibility domain, like the optimization problem, becomes time-invariant.

The time-varying property of the online optimizations means that there is an increased memory requirement and the approach is less robust against unmodelled disturbances, both when compared with other approaches with time-invariant constraint tightening. With the control policy considered here (part of which is computed offline), the use of the tightening parameters without the update might result in infeasibility. Thus, a potential line of research is the investigation of conditions on the tightening parameters so as to obtain guarantees of recursive feasibility and stability. These extra conditions, however, might impose extra conservativeness thus reducing the domains of attraction, but will solve the problems stated above.

Chapter 4

Striped Parameterized Tube Model Predictive Control

As was discussed in the previous chapter, the recently introduced parameterized tube model predictive control (PTMPC) [73] uses a triangular separable prediction strategy that induces a nominal sequence associated with the dynamics of the current initial state and N partial tube pairs, where each partial tube is associated with the effect of a particular future disturbance acting during Mode 1. In Mode 2, the terminal control law can be taken to be a static state feedback. PTMPC online implementation reduces to a standard convex optimization, and has a structure that causes the number of decision variables and constraints to rise quadratically with N . In common with the strategy of Chapter 3, this chapter proposes a modification of PTMPC with the aim of reducing the online computational complexity. The modification consists in letting all the partial tubes associated with different disturbances acting on the system to be given by the same partial sequences and allowing their compensatory effect to extend beyond the horizon N into Mode 2. With this modification, the prediction structure is of a triangular striped nature and extends over an infinite horizon. As illustrated by simulation examples, the resulting scheme can outperform PTMPC in terms of the

size of the region of attraction and affords a reduction in computation which allows one to use a longer horizon (for the same or reduced computational demand).

4.1 Introduction

As discussed in Sections 2.2 and 3.2, ideal RMPC uses DP [9] or feedback optimization [86] which are computationally intractable, and thus researchers have looked for sensible compromises between sub-optimality and complexity of the online problems. Attempts based on quasi closed-loop policies are computationally efficient, with a number of variables and constraints of the online optimization that grow linearly with N , but can be conservative (i.e. rigid tube MPC [63], [61]). RMPC strategies with the affine-in-the-disturbances control policy [58], [38] have an improved degree of sub-optimality, but this was later superseded by PTMPC [73]. PTMPC is based on a separable prediction strategy and a separable non-linear control policy that results in a triangular prediction structure of $N + 1$ partial tubes. The 0^{th} (or nominal) partial state and control tube accounts for the dynamical contribution of the current initial state while the remainder, the j^{th} partial tube pairs, capture the dynamical contribution due to disturbances. However, the number of variables and constraints of the online optimization in RMPC with the affine-in-the-disturbances control policy and PTMPC grow quadratically with N .

The simplified PTMPC strategy presented in Chapter 3, OPTMPC, is based on performing offline rather than online some of the computations of the predicted control policy. OPTMPC manages to reduce the computational complexity of PTMPC synthesis with little degradation on average of the degree of sub-optimality, in terms of size of the domain of attraction, while preserving its desirable system theoretic properties. Based on this, for some systems it will be possible to increase the prediction horizon N for the OPTMPC strategy which may allow one to obtain greater domains of attractions at lower computational efforts than PTMPC. This however, will not

happen every time. In some cases, the conservativeness imposed by the offline design will be such that regardless how big the prediction horizon is chosen, OPTMPC will not have a domain of attraction as large as PTMPC using even a short N . This will clearly be a problem when large domains of attraction are needed.

An alternative approach to reduce the computational complexity of PTMPC is presented in this chapter.

Unlike earlier work, this chapter presents an RMPC formulation where the degrees of freedom of the predicted control policy that provide disturbance compensation affect (directly) the inputs over an infinite prediction horizon, thus resulting in a terminal control law more general than the conventional static state feedback. This is achieved by a modification of the PTMPC structure, such that each j^{th} partial tube pair is defined by the same partial extreme sequences rather than having specific partial extreme sequences for each partial tube pair. This gives rise to a triangular striped prediction structure where the disturbance compensation terms are allowed to extend over an infinite prediction horizon. On account of this the constraints are less tight in the infinite horizon while the number of variables and constraints grows linearly with N , rather than quadratically as for PTMPC. Simulations demonstrate that this strategy can, for a comparable number of decision variables and constraints, lead to larger domains of attraction.

The chapter is structured as follows. In Section 4.2 the system description is given and the modification of the separable prediction scheme of [73], which consists on the imposition of a striped structure, is introduced. The control strategy of this chapter together with the approach for constraint satisfaction are presented in Section 4.3. Section 4.4 presents the online implementation and theoretical properties of two variations of the proposed strategy: one that gives increased control freedom and provides ISS stability; and another with a reduced control freedom but a stronger stability result. Section 4.5 illustrates the benefits of the proposed strategy with

simulation examples, and final conclusions are drawn in Section 4.6.

4.2 System description and separable prediction scheme

Consider a linear, time-invariant, discrete-time system with additive disturbances

$$x_{t+1} = Ax_t + Bu_t + w_t, \quad (4.1)$$

where $x_t \in \mathbb{R}^n$, $w_t \in \mathbb{W} \subseteq \mathbb{R}^n$, and $u_t \in \mathbb{R}^{n_u}$ are the state, disturbance and control input. The admissible disturbance set is a polytopic set that contains the origin and is given by

$$\mathbb{W} = \text{conv}\{\tilde{w}_i : i = 1, \dots, q\}. \quad (4.2)$$

The system is subject to mixed state and input constraints

$$(x_t, u_t) \in \mathbb{Y}, \quad (4.3)$$

where $\mathbb{Y}$ is a polytopic set that contains the origin in its interior and is given by

$$\mathbb{Y} = \{(x, u) : Fx + Gu \leq \underline{1}\}, \quad (4.4)$$

where $F \in \mathbb{R}^{n_c \times n}$ and $G \in \mathbb{R}^{n_c \times n_u}$.

Assumption 4.1.

(i) (A, B) is stabilizable.

(ii) The matrix K is such that $\Phi = A + BK$ is strictly Hurwitz and the minimal robust positively invariant set for system (4.1) under $u = Kx$, $\Omega_\infty(\mathbb{W})$, satisfies $(I, K)\Omega_\infty(\mathbb{W}) \subseteq \text{interior}(\mathbb{Y})$.

Assume a separable prediction scheme such as that of [73], where the nominal

dynamics and the contribution of each future disturbance are separated into different sequences. This though is done in a different way to that of [73] in the sense that the sequences cover an infinite prediction horizon instead of only Mode 1. $\mathbf{x}_{(0,\infty)} = \{x_{(0,0)}, x_{(0,1)}, \dots\}$ and $\mathbf{u}_{(0,\infty)} = \{u_{(0,0)}, u_{(0,1)}, \dots\}$ are the 0^{th} partial state and control sequences, which account for the nominal dynamics

$$\begin{aligned} x_{(0,0)} &= x_t, \\ x_{(0,k+1)} &= Ax_{(0,k)} + Bu_{(0,k)}, \quad k = 0, 1, \dots, \end{aligned} \tag{4.5}$$

and $\mathbf{x}_{(j,\infty)} = \{x_{(j,j)}, x_{(j,j+1)}, \dots\}$, $\mathbf{u}_{(j,\infty)} = \{u_{(j,j)}, u_{(j,j+1)}, \dots\}$, for $j = 1, 2, \dots$, are the j^{th} partial state and control sequences describing the dynamical contribution of the disturbance $w_{(j-1)}$ which acts on the system at the $(j-1)^{th}$ prediction instant. These dynamics are given by

$$\begin{aligned} x_{(j,j)} &= w_{(j-1)}, \\ x_{(j,k+1)} &= Ax_{(j,k)} + Bu_{(j,k)}, \quad k = j, j+1, \dots \end{aligned} \tag{4.6}$$

Recall that for simplicity the notation of the type $x_{t+k|t}$ is not used for the k -step ahead predictions made at time t , and instead $(x_{(0,k)}, u_{(0,k)})$ and $(x_{(j,k)}, u_{(j,k)})$ are the k -step-ahead predictions made at each time t associated to the 0^{th} partial sequences and the j^{th} partial sequences, while $(x_{(k)}, u_{(k)})$ are the full k -step- ahead predictions made at each time t . Accordingly, the full predictions are given by

$$x_{(k)} = \sum_{j=0}^k x_{(j,k)}, \quad u_{(k)} = \sum_{j=0}^k u_{(j,k)}, \quad k = 0, 1, \dots, \tag{4.7}$$

with

$$\begin{aligned} x_{(j,k)} \in X_{(j,k)} &= \begin{cases} x_{(0,k)} & k = 0 \\ \text{conv}\{x_{(i,j,k)} : i = 1, \dots, q\} & k = 1, 2, \dots \end{cases} \\ u_{(j,k)} \in U_{(j,k)} &= \begin{cases} u_{(0,k)} & k = 0 \\ \text{conv}\{u_{(i,j,k)} : i = 1, \dots, q\} & k = 1, 2, \dots \end{cases} \end{aligned} \quad (4.8)$$

Eq.(4.7) defines a triangular prediction scheme, as in the Mode 1 section of Figure 3.1 but in this case throughout the entire prediction horizon, according to which

$$\begin{aligned} x_{(k)} \in X_{(k)} &= \bigoplus_{j=0}^k X_{(j,k)}, \quad k = 0, 1, \dots, \\ u_{(k)} \in U_{(k)} &= \bigoplus_{j=0}^k U_{(j,k)}, \quad k = 0, 1, \dots, \end{aligned} \quad (4.9)$$

where $X_{(k)}$ and $U_{(k)}$ are the state and control-cross sections of the state tube $\mathbf{X}_\infty = \{X_{(0)}, X_{(1)}, \dots\}$ and control tube $\mathbf{U}_\infty = \{U_{(0)}, U_{(1)}, \dots\}$, respectively. From (4.8), the j^{th} partial state and control sequences are defined by the j^{th} extreme partial state and control sequences, $\mathbf{x}_{(i,j,\infty)} = \{x_{(i,j,j)}, x_{(i,j,j+1)}, \dots, x_{(i,j,\infty)}\}$ and $\mathbf{u}_{(i,j,\infty)} = \{u_{(i,j,j)}, u_{(i,j,j+1)}, \dots, u_{(i,j,\infty)}\}$ respectively, which describe the effect of the extreme disturbances acting at the prediction instant $j - 1$. The dynamics of the j^{th} partial extreme state and control sequences for $i = 1, \dots, q$ are given by

$$\begin{aligned} x_{(i,j,j)} &= \tilde{w}_i, & j &= 1, 2, \dots, \\ x_{(i,j,k+1)} &= Ax_{(i,j,k)} + Bu_{(i,j,k)}, & j &= 1, 2, \dots, \quad k \geq j. \end{aligned} \quad (4.10)$$

In [73], $x_{(i,j,k)}$, for $j = 1, \dots, N$, $k = j, \dots, N$ and $u_{(i,j,N)}$ for $j = 1, \dots, N - 1$, $k = j, \dots, N - 1$ are chosen freely for different j , while those associated with Mode 2 are determined by the evolution dictated by the state feedback $u = Kx$, thus inducing the triangular prediction scheme as in Figure 3.1. Instead, here a striped structure is

invoked such that $u_{(i,j,k)}$ is set to satisfy

$$u_{(i,j,k)} := u_{(i,1,k-j+1)}, \quad i = 1, \dots, q, \quad j = 1, 2, \dots, \quad k \geq j, \quad (4.11)$$

which implies that

$$x_{(i,j,k)} = x_{(i,1,k-j+1)}, \quad i = 1, \dots, q, \quad j = 1, 2, \dots, \quad k \geq j, \quad (4.12)$$

and therefore

$$X_{(j,k)} = X_{(1,k-j+1)}, \quad U_{(j,k)} = U_{(1,k-j+1)}, \quad i = 1, \dots, q, \quad j = 1, 2, \dots, \quad k \geq j. \quad (4.13)$$

As a result all the j^{th} partial extreme state and control sequences are fully defined in terms of the 1^{st} partial extreme state and control sequences $\mathbf{x}_{(i,1,\infty)}$ and $\mathbf{u}_{(i,1,\infty)}$. This implies that (4.9) can be re-written as

$$\begin{aligned} x_{(k)} \in X_{(k)} &= \begin{cases} x_{(0,0)} & k = 0 \\ x_{(0,k)} \oplus \bigoplus_{j=1}^k X_{(1,j)} & k = 1, 2, \dots \end{cases} \\ u_{(k)} \in U_{(k)} &= \begin{cases} u_{(0,0)} & k = 0 \\ u_{(0,k)} \oplus \bigoplus_{j=1}^k U_{(1,j)} & k = 1, 2, \dots \end{cases} \end{aligned} \quad (4.14)$$

thus inducing a striped triangular prediction scheme as shown in Figure 4.1, where the predictions of $x_{(k)}$ and $u_{(k)}$ are contained in the cross-sections $X_{(k)}$ and $U_{(k)}$ that are given by the Minkowski sum of the elements in the k -th column. Due to this striped structure, the strategies based on this prediction scheme are referred to as Striped PTMPC or SPTMPC.

					Mode 1	Mode 2		
$x_{(0,0)}$	$x_{(0,1)}$	$x_{(0,2)}$	$\cdots$	$x_{(0,N-1)}$	$x_{(0,N)}$	$x_{(0,N+1)}$	$x_{(0,N+2)}$	$\cdots$
	$X_{(1,1)}$	$X_{(1,2)}$	$\cdots$	$X_{(1,N-1)}$	$X_{(1,N)}$	$X_{(1,N+1)}$	$X_{(1,N+2)}$	$\cdots$
		$X_{(1,1)}$	$\cdots$	$X_{(1,N-2)}$	$X_{(1,N-1)}$	$X_{(1,N)}$	$X_{(1,N+1)}$	$\cdots$
			$\ddots$	$\cdots$		$\vdots$	$\vdots$	
				$X_{(1,1)}$	$X_{(1,2)}$	$X_{(1,3)}$	$X_{(1,4)}$	$\cdots$
					$X_{(1,1)}$	$X_{(1,2)}$	$X_{(1,3)}$	$\cdots$
						$X_{(1,1)}$	$X_{(1,2)}$	$\cdots$
							$X_{(1,1)}$	$\cdots$
$X_{(0)}$	$X_{(1)}$	$X_{(2)}$	$\cdots$	$X_{(N-1)}$	$X_{(N)}$	$X_{(N+1)}$	$X_{(N+2)}$	$\cdots$

					Mode 1	Mode 2		
$u_{(0,0)}$	$u_{(0,1)}$	$u_{(0,2)}$	$\cdots$	$u_{(0,N-1)}$	$u_{(0,N)}$	$u_{(0,N+1)}$	$u_{(0,N+2)}$	$\cdots$
	$U_{(1,1)}$	$U_{(1,2)}$	$\cdots$	$U_{(1,N-1)}$	$U_{(1,N)}$	$U_{(1,N+1)}$	$U_{(1,N+2)}$	$\cdots$
		$U_{(1,1)}$	$\cdots$	$U_{(1,N-2)}$	$U_{(1,N-1)}$	$U_{(1,N+1)}$	$U_{(1,N+2)}$	$\cdots$
			$\ddots$	$\vdots$	$\vdots$	$\vdots$		
				$U_{(1,1)}$	$U_{(1,2)}$	$U_{(1,3)}$	$U_{(1,3)}$	$\cdots$
					$U_{(1,1)}$	$U_{(1,2)}$	$U_{(1,3)}$	$\cdots$
						$U_{(1,1)}$	$U_{(1,2)}$	$\cdots$
							$U_{(1,1)}$	$\cdots$
$U_{(0)}$	$U_{(1)}$	$U_{(2)}$	$\cdots$	$U_{(N-1)}$	$U_{(N)}$	$U_{(N+1)}$	$U_{(N+2)}$	$\cdots$

Figure 4.1: Striped triangular prediction scheme of SPTMPC

The striped separable prediction scheme, as reflected by (4.13), and (4.14) induces a separable non-linear predicted control policy $\pi_{N-1} = \{\mu_0(\cdot), \mu_1(\cdot), \dots\}$ given by

$$\begin{aligned}
 \mu_k(x_{(k)}) &= \sum_{j=1}^k \mu_{(j,k)}(x_{(j,k)}), \quad k = 0, 1, \dots, \quad \text{with} \\
 x_{(k)} &= x_{(0,k)} + \sum_{j=1}^k x_{(j,k)}, \quad x_{(j,k)} = \sum_{i=1}^q \lambda_{(i,j)}(w_{(j-1)}) x_{(i,1,k-j+1)}, \quad j = 1, 2, \dots, \\
 \mu_{(0,k)}(x_{(0,k)}) &= u_{(0,k)}, \quad k = 0, 1, \dots, \\
 \mu_{(j,k)}(x_{(j,k)}) &= \sum_{i=1}^q \lambda_{(i,j)}(w_{(j-1)}) u_{(i,1,k-j+1)}, \quad j = 1, 2, \dots, k = j, j+1, \dots
 \end{aligned} \tag{4.15}$$

where $\lambda_{(i,j)}(w_{(j-1)})$ are convex interpolation parameters that satisfy $\lambda_{(i,j)}(w_{(j-1)}) \geq 0$, $\sum_{i=1}^q \lambda_{(i,j)}(w_{(j-1)}) = 1$ and $w_{(j-1)} = \sum_{i=1}^q \lambda_{(i,j)}(w_{(j-1)})$. A well-defined and unique selection of such convex multipliers $\{\lambda_{(i,j)}, i = 1, \dots, q\}$ can be obtained by employ-

ing, for instance, a least squares procedure. Note how the expressions of $x_{(j,k)}$ and $\mu_{(j,k)}(x_{(j,k)})$ in (4.15) imply that $x_{(j,k)} \in X_{(1,k-j+1)}$ and $u_{(j,k)} \in U_{(1,k-j+1)}$ as previously indicated by (4.14).

It can be seen from (4.15) and Figure 4.1 that all the predicted inputs $u_{(k)} = \mu_k(x_{(k)})$, $k = 0, 1, \dots$, depend on the on the 1^{st} partial extreme sequence $\mathbf{u}_{(i,1,\infty)}$; more precisely, the predicted $u_{(k)}$ is defined by the elements $u_{(i,1,j)}$, $j = 1, \dots, k$. Of these the elements $u_{(i,1,j)}$, $j = 1, \dots, N-1$, as will be seen next are degrees of freedom of the control policy, which implies that these provide disturbance compensation over the infinite horizon. This structure is different than just reducing the number of variables by forcing a striped structure in the parameterization of [73], since in such a parameterization the control policy would have degrees of freedom only in Mode 1, i.e. for $k = 0, \dots, N-1$.

4.3 Striped PTMPC: Predictions and constraints

This section introduces the strategies for guaranteeing constraint satisfaction under the striped prediction scheme presented in Section 4.2.

Remark 4.1. *The presentation of the chapter from this point will assume that $N \geq 2$. The reason for this assumption is that if $N = 1$, then the j^{th} partial sequences are not optimization variables, so adopting a striped structure makes no difference with respect to the predictions of PTMPC. Furthermore, the case of $N = 1$ clearly is of little interest since only the first predicted input, $u_{(0,0)}$, is optimized in the online problem. However, the results can be easily extended to the case of $N = 1$ as will be indicated in Remark 4.5.*

4.3.1 Prediction strategy

It is usual practice in MPC to use a fixed stabilizing terminal control law in Mode 2. However in the strategy proposed in this chapter no fixed terminal control law

is assumed. Instead, the following structure with is used for limiting the number of decision variables: (i) the first N inputs of the 0^{th} partial control sequence are degrees of freedom and afterwards they are defined by a stabilizing linear gain ($u = Kx$); and (ii) the first $N - 1$ inputs of the 1^{st} partial extreme control sequences are degrees of freedom, and the following inputs are defined by a stabilizing linear gain ($u = Kx$). Thus, the dynamic equations for 0^{th} partial sequences and 1^{st} partial extreme sequences are given by

$$x_{(0,0)} = x_t, \tag{4.16a}$$

$$x_{(0,k+1)} = Ax_{(0,k)} + Bu_{(0,k)}, \quad k = 0, \dots, N - 1, \tag{4.16b}$$

$$x_{(0,k+1)} = \Phi x_{(0,k)}, \quad k \geq N, \tag{4.16c}$$

and

$$x_{(i,1,1)} = \tilde{w}_i, \tag{4.17a}$$

$$x_{(i,1,k+1)} = Ax_{(i,1,k)} + Bu_{(i,1,k)}, \quad k = 1, \dots, N - 1, \tag{4.17b}$$

$$x_{(i,1,k+1)} = \Phi x_{(i,1,k)}, \quad k \geq N. \tag{4.17c}$$

From (4.16) and (4.17), the striped separable prediction scheme of Figure 4.1 can be more precisely represented by Figure 4.2.

In order to deal with the mixed state and control constraints of (4.3), attention is turned to mixed tubes. The separable prediction scheme of equations (4.16),(4.17) implies that the predictions of $(x_{(k)}, u_{(k)})$ will be contained in the overall mixed state and control tube $\mathbf{Y}_\infty = \{Y_{(0)}, Y_{(1)}, \dots, Y_{(\infty)}\}$, where

$$Y_{(0)} = Y_{(0,0)}, \quad (4.18)$$

$$Y_{(k)} = Y_{(0,k)} \oplus Y_{(1,1)} \oplus \dots \oplus Y_{(1,k)}, \quad k = 1, 2, \dots,$$

where the partial mixed state and control tubes $\mathbf{Y}_{(j,\infty)} = \{Y_{(j,j)}, Y_{(j,j+1)}, \dots, Y_{(j,\infty)}\}$, $j = 0, 1$, contain the partial sequences $\mathbf{x}_{(j,\infty)}$ and $\mathbf{u}_{(j,\infty)}$, such that 0^{th} and the 1^{st} mixed state and control cross-sections are given by $Y_{(0,k)} = (x_{(0,k)}, u_{(0,k)})$ and $Y_{(1,k)} = \text{conv}\{(x_{(i,1,k)}, u_{(i,1,k)})\}$, respectively.

					Mode 1	Mode 2		
$x_{(0,0)}$	$x_{(0,1)}$	$x_{(0,2)}$	$\dots$	$x_{(0,N-1)}$	$x_{(0,N)}$	$\Phi x_{(0,N)}$	$\Phi^2 x_{(0,N)}$	$\dots$
	$X_{(1,1)}$	$X_{(1,2)}$	$\dots$	$X_{(1,N-1)}$	$X_{(1,N)}$	$\Phi X_{(1,N)}$	$\Phi^2 X_{(1,N)}$	$\dots$
		$X_{(1,1)}$	$\dots$	$X_{(1,N-2)}$	$X_{(1,N-1)}$	$X_{(1,N)}$	$\Phi X_{(1,N)}$	$\dots$
			$\ddots$	$\dots$		$\vdots$	$\vdots$	
				$X_{(1,1)}$	$X_{(1,2)}$	$X_{(1,3)}$	$X_{(1,4)}$	$\dots$
					$X_{(1,1)}$	$X_{(1,2)}$	$X_{(1,3)}$	$\dots$
						$X_{(1,1)}$	$X_{(1,2)}$	$\dots$
							$X_{(1,1)}$	$\dots$
$X_{(0)}$	$X_{(1)}$	$X_{(2)}$	$\dots$	$X_{(N-1)}$	$X_{(N)}$	$X_{(N+1)}$	$X_{(N+2)}$	$\dots$
					Mode 1	Mode 2		
$u_{(0,0)}$	$u_{(0,1)}$	$u_{(0,2)}$	$\dots$	$u_{(0,N-1)}$	$Kx_{(0,N)}$	$K\Phi x_{(0,N)}$	$K\Phi^2 x_{(0,N)}$	$\dots$
	$U_{(1,1)}$	$U_{(1,2)}$	$\dots$	$U_{(1,N-1)}$	$KX_{(1,N)}$	$K\Phi X_{(1,N)}$	$K\Phi^2 X_{(1,N)}$	$\dots$
		$U_{(1,1)}$	$\dots$	$U_{(1,N-2)}$	$U_{(1,N-1)}$	$KX_{(1,N)}$	$K\Phi X_{(1,N)}$	$\dots$
			$\ddots$	$\vdots$	$\vdots$	$\vdots$	$\vdots$	
				$U_{(1,1)}$	$U_{(1,2)}$	$U_{(1,3)}$	$U_{(1,3)}$	$\dots$
					$U_{(1,1)}$	$U_{(1,2)}$	$U_{(1,3)}$	$\dots$
						$U_{(1,1)}$	$U_{(1,2)}$	$\dots$
							$U_{(1,1)}$	$\dots$
$U_{(0)}$	$U_{(1)}$	$U_{(2)}$	$\dots$	$U_{(N-1)}$	$U_{(N)}$	$U_{(N+1)}$	$U_{(N+2)}$	$\dots$

Figure 4.2: Striped triangular prediction scheme of SPTMPC

In the online optimization the decision variables are $\mathbf{x}_{(0,N)}$, $\mathbf{x}_{(i,1,N)}$, $\mathbf{u}_{(0,N-1)}$, $\mathbf{u}_{(i,1,N-1)}$ (along with the constraint and cost tightening parameters, which will be introduced below). It would then be convenient to re-write (4.18) as a function of parameters that are decision variables of the online optimization only. Using (4.16) and (4.17), $Y_{(0,k)}$ and $Y_{(1,k)}$ can be re-expressed as

$$Y_{(0,k)} = (x_{(0,k)}, u_{(0,k)}), \quad k = 0, \dots, N-1, \quad (4.19a)$$

$$Y_{(0,k)} = (\Phi^{k-N} x_{(0,N)}, K \Phi^{k-N} x_{(0,N)}), \quad k \geq N, \quad (4.19b)$$

$$= (I, K) \Phi^{k-N} x_{(0,N)}, \quad k \geq N, \quad (4.19c)$$

and

$$Y_{(1,k)} = \text{conv}\{(x_{(i,1,k)}, u_{(i,1,k)}) : i = 1, \dots, q\}, \quad k = 1, \dots, N-1, \quad (4.20a)$$

$$Y_{(1,k)} = \text{conv}\{(\Phi^{k-N} x_{(i,1,N)}, K \Phi^{k-N} x_{(i,1,N)}) : i = 1, \dots, q\}, \quad k \geq N, \quad (4.20b)$$

$$= (I, K) \Phi^{k-N} X_{(1,N)}, \quad k \geq N, \quad (4.20c)$$

Then (4.18) can be rewritten as

$$Y_{(0)} = Y_{(0,0)}, \quad (4.21a)$$

$$Y_{(k)} = Y_{(0,k)} \oplus \bigoplus_{j=1}^k Y_{(1,j)}, \quad k = 1, \dots, N-1, \quad (4.21b)$$

$$Y_{(k)} = (I, K) \Phi^{k-N} x_{(0,N)} \oplus \bigoplus_{j=1}^{N-1} Y_{(1,j)} \oplus \bigoplus_{j=N}^k (I, K) \Phi^{j-N} X_{(1,N)}, \quad k \geq N, \quad (4.21c)$$

which only depend on $\mathbf{x}_{(0,N)}$, $\mathbf{x}_{(i,1,N)}$, $\mathbf{u}_{(0,N-1)}$, $\mathbf{u}_{(i,1,N-1)}$.

4.3.2 Constraints

The condition for constraint satisfaction is that

$$Y_{(k)} \subseteq \mathbb{Y}, \quad k = 0, 1, \dots \quad (4.22)$$

Satisfaction of (4.22) cannot be ensured explicitly for all $k = 0, 1, \dots$ because it would require an infinite number of constraints. Additionally, the usual strategy consisting on using a terminal set that is invariant and ensures feasibility of constraints for a fixed terminal control law is not possible here. On account of having disturbance compensation throughout the infinite prediction horizon, the terminal control law is not fixed. Thus, an alternative approach is considered here. A secondary horizon $N_2 \geq 0$ is used such that constraint satisfaction is enforced explicitly for $k = 0, \dots, N + N_2 - 1$, whereas satisfaction of (4.22) for $k \geq N + N_2$ is guaranteed by a set of terminal constraints: one that verifies that a set that contains the predicted $Y_{(k)}$ for $k \geq N + N_2$ is inside $\mathbb{Y}$, and two other constraints on the 0^{th} and the 1^{st} partial sequences at the $(N + N_2)^{th}$ prediction instant.

Proposition 4.2. *Let $\mathbb{X}_0$ be a positively invariant set for the nominal dynamics of (4.16c) and let $\mathbb{X}_1$ be an arbitrary polytopic set that contains the origin. Satisfaction of*

$$x_{(0, N+N_2)} \in \mathbb{X}_0, \quad (4.23a)$$

$$X_{(1, N+N_2)} \in \mathbb{X}_1, \quad (4.23b)$$

implies that for all $k \geq N + N_2$ $Y_{(k)}$ satisfies

$$Y_{(k)} \subseteq \bar{Y}_{(\infty)}, \quad (4.24)$$

where

$$\bar{Y}_{(\infty)} = (I, K)\mathbb{X}_0 \oplus \bigoplus_{j=1}^{N-1} Y_{(1,j)} \oplus \bigoplus_{j=N}^{N+N_2-1} (I, K)\Phi^{j-N} X_{(1,N)} \oplus (I, K)\Omega_{\infty}(\mathbb{X}_1), \quad (4.25)$$

and $\Omega_{\infty}(\mathbb{X}_1) = \bigoplus_{j=0}^{\infty} \Phi^j \mathbb{X}_1$ is the minimal robust invariant set for the system with dynamics $z^+ = \Phi z + w$, with $w \in \mathbb{X}_1$.

Proof. From (4.21c), the predicted $Y_{(k)}$ for $k \geq N + N_2$ can be expressed as

$$\begin{aligned}
 Y_{(k)} = & (I, K) \Phi^{k-(N+N_2)} x_{(0, N+N_2)} \oplus \bigoplus_{j=1}^{N-1} Y_{(1, j)} \oplus \bigoplus_{j=N}^{N+N_2-1} (I, K) \Phi^{j-N} X_{(1, N)} \oplus \\
 & \bigoplus_{j=N+N_2}^k (I, K) \Phi^{j-N} X_{(1, N)},
 \end{aligned} \tag{4.26}$$

Since $\mathbb{X}_0$ is invariant for the dynamics $z^+ = \Phi z$, satisfaction of (4.23a) implies that the first term of the RHS of (4.26) satisfies $\Phi^{k-(N+N_2)} x_{(0, N+N_2)} \in \mathbb{X}_0$. Additionally, satisfaction of (4.23b) implies that

$$\bigoplus_{j=N+N_2}^k \Phi^{j-N} X_{(1, N)} = \bigoplus_{j=N+N_2}^k \Phi^{j-(N+N_2)} X_{(1, N+N_2)} \subseteq \bigoplus_{j=0}^{k-(N+N_2)} \Phi^j \mathbb{X}_1 \subseteq \Omega_\infty(\mathbb{X}_1),$$

and therefore the last term of the RHS (4.26) can be bounded as

$$\bigoplus_{j=N+N_2}^k (I, K) \Phi^{j-N} X_{(1, N)} \subseteq (I, K) \Omega_\infty(\mathbb{X}_1), \quad k \geq N + N_2,$$

which yields that $Y_{(k)}$ for $k \geq N + N_2$ is bounded as

$$Y_{(k)} \subseteq (I, K) \mathbb{X}_0 \oplus \bigoplus_{j=1}^{N-1} Y_{(1, j)} \oplus \bigoplus_{j=N}^{N+N_2-1} (I, K) \Phi^{j-N} X_{(1, N)} \oplus (I, K) \Omega_\infty(\mathbb{X}_1),$$

thus proving the desired result. $\square$

The bounding above can be visualized in Figure 4.3, where the columns associated with the predictions for $k = N + N_2 - 1, \dots, N + N_2 + 2$ are displayed. The terms from $Y_{(1,1)}$ to $Y_{(1, N-1)}$, from $(I, K) X_{(1, N)}$ to $(I, K) \Phi^{N_2-1} X_{(1, N)}$ and $(I, K) \mathbb{X}_0$ appear in all the columns and thus the bound, $\bar{Y}_{(\infty)}$, contains them (as the second, third and first terms in the RHS of (4.25)). $x_{(0, N_2)}$ is required to be in $\mathbb{X}_0$, and due to its invariance the nominal state will remain there. The remaining terms, highlighted in red in Figure 4.3, due to the constraint that $X_{(1, N+N_2)} \in \mathbb{X}_1$, will be contained in $\mathbb{X}_1$,

$\mathbb{X}_1 \oplus \Phi \mathbb{X}_1$, $\mathbb{X}_1 \oplus \Phi \mathbb{X}_1 \Phi^2 \mathbb{X}_1$, $\dots$, and each term is bounded by $\Omega_\infty(\mathbb{X}_1)$.

$(I, K)\Phi^{N_2-1}x_{(0,N)}$	$(I, K)\mathbb{X}_0$	$(I, K)\mathbb{X}_0$	$(I, K)\mathbb{X}_0$
$(I, K)\Phi^{N_2-1}X_{(1,N)}$	$(I, K)\mathbb{X}_1$	$(I, K)\Phi \mathbb{X}_1$	$(I, K)\Phi^2 \mathbb{X}_1$
$\vdots$	$(I, K)\Phi^{N_2-1}X_{(1,N)}$	$(I, K)\mathbb{X}_1$	$(I, K)\Phi \mathbb{X}_1$
$(I, K)\Phi X_{(1,N)}$	$\vdots$	$(I, K)\Phi^{N_2-1}X_{(1,N)}$	$(I, K)\mathbb{X}_1$
$(I, K)X_{(1,N)}$	$(I, K)\Phi X_{(1,N)}$	$\vdots$	$(I, K)\Phi^{N_2-1}X_{(1,N)}$
$Y_{(1,N-1)}$	$(I, K)X_{(1,N)}$	$(I, K)\Phi X_{(1,N)}$	$\vdots$
$\vdots$	$Y_{(1,N-1)}$	$(I, K)X_{(1,N)}$	$(I, K)\Phi X_{(1,N)}$
$Y_{(1,1)}$	$\vdots$	$Y_{(1,N-1)}$	$(I, K)X_{(1,N)}$
	$Y_{(1,1)}$	$\vdots$	$Y_{(1,N-1)}$
		$Y_{(1,1)}$	$\vdots$
			$Y_{(1,1)}$
$Y_{(N+N_2-1)}$	$Y_{(N+N_2)}$	$Y_{(N+N_2+1)}$	$Y_{(N+N_2+2)}$

Figure 4.3: Predicted $Y_{(k)}$ after the secondary horizon

The construction of a set that contains the state and input predictions for all $k \geq N + N_2$, enables the use of a finite number of constraints to ensure constraints satisfaction:

Corollary 4.3. *Satisfaction of (4.22) is guaranteed if the 0^{th} and 1^{st} partial sequences satisfy (4.23) and*

$$Y_{(k)} \subseteq \mathbb{Y}, \quad k = 0, \dots, N + N_2 - 1, \quad (4.27)$$

$$\bar{Y}_{(\infty)} \subseteq \mathbb{Y}. \quad (4.28)$$

Condition (4.27) can be written equivalently as linear constraints on the 0^{th} partial state and control sequences by using slack variables that account for the worst case realizations associated with the 1^{st} partial extreme state and control sequences. Condition (4.27) is then equivalent to

$$F_l x_{(0,0)} + G_l u_{(0,0)} \leq 1, \quad l = 1, \dots, n_c, \quad (4.29a)$$

$$F_l x_{(0,k)} + G_l u_{(0,k)} + \sum_{j=1}^k f_{(l,1,j)} \leq 1, \quad k = 1, \dots, N-1, \quad l = 1, \dots, n_c, \quad (4.29b)$$

$$f_{(l,1,j)} \geq F_l x_{(i,1,j)} + G_l u_{(i,1,j)}, \quad j = 1, \dots, N-1, \quad i = 1, \dots, q, \quad l = 1, \dots, n_c, \quad (4.29c)$$

$$(F_l + G_l K) \Phi^{j-N} x_{(0,N)} + \sum_{j=1}^k f_{(l,1,j)} \leq 1, \quad k = N, \dots, N+N_2-1, \quad l = 1, \dots, n_c, \quad (4.29d)$$

$$f_{(l,1,j)} \geq (F_l + G_l K) \Phi^{j-N} x_{(i,1,N)}, \quad j = N, \dots, N+N_2-1, \quad i = 1, \dots, q, \quad l = 1, \dots, n_c, \quad (4.29e)$$

where F_l and G_l are the l^{th} rows of F and G , respectively.

Using the same ideas, satisfaction of (4.28) can also be imposed equivalently by

$$f_{(l,0,\infty)} + \sum_{j=1}^{N+N_2-1} f_{(l,1,j)} + f_{(l,1,\infty)} \leq 1, \quad l = 1, \dots, n_c, \quad (4.30)$$

where

$$f_{(l,0,\infty)} = \max_{x \in \mathbb{X}_0} (F_l + G_l K)x, \quad f_{(l,1,\infty)} = \max_{x \in \Omega_\infty(\mathbb{X}_1)} (F_l + G_l K)x \quad (4.31)$$

Since $\Omega_\infty(\mathbb{X}_1)$ may be defined by an infinite number of constraints, practical application indicates that it should be replaced by an outer approximation with an arbitrary precision (which can be computed by the technique of [72]) in the definition of $f_{(l,1,\infty)}$.

In this case, (4.30) is only a sufficient condition for (4.28).

Assuming that $\mathbb{X}_0 = \{x : Hx \leq \underline{1}\}$ and $\mathbb{X}_1 = \{x : Lx \leq \underline{1}\}$, where $H \in \mathbb{R}^{n_{f0} \times n}$ and $L \in \mathbb{R}^{n_{f1} \times n}$, then condition (4.23) can be expressed as

$$H_l x_{(0, N+N_2)} \leq 1, \quad l = 1, \dots, n_{f0}, \quad (4.32a)$$

$$L_l x_{(i, 1, N+N_2)} \leq 1, \quad l = 1, \dots, n_{f1}, \quad i = 1, \dots, q, \quad (4.32b)$$

where H_l and L_l are the l^{th} rows of H and L , respectively.

Assumption 4.2. $\mathbb{X}_1$ is defined by

$$\mathbb{X}_1 = \Phi^{N+N_2-1} \mathbb{W}, \quad (4.33)$$

and $\mathbb{X}_0$ is a robust positively invariant set for $z^+ = \Phi z + w$, $w \in \Phi \mathbb{X}_1$.

Remark 4.4. The assumption above is convenient in terms of Proposition 4.2 and the results related with recursive feasibility, Theorem 4.8 and Theorem 4.18, since recursive feasibility requires that $\mathbb{X}_0$ is a robust positively invariant set for $z^+ = \Phi z + w$, $w \in \Phi \mathbb{X}_1$. At time $t + 1$ the nominal predictions will consider the effect of the disturbance acting on the system at time t , w_t , since its effect will be already present at x_{t+1} . Then, the candidate solution at time $t + 1$ for the $N + N_2$ -step-ahead prediction of the nominal state $\tilde{x}_{(0, N+N_2)}$ will also have to consider the effect of the disturbance that has just acted on the system, effect which is contained in $\Phi \mathbb{X}_1$ due to (4.23b). Thus robust invariance for a system with the disturbance in $\Phi \mathbb{X}_1$ will guarantee that $\tilde{x}_{(0, N+N_2)} = \Phi x_{(0, N+N-2)} + \Phi x_{(1, N+N_2)} \in \mathbb{X}_0$ provided that at time t $x_{(0, N+N_2)} \in \mathbb{X}_0$ and $x_{(1, N+N_2)} \in \mathbb{X}_1$.

Additionally, it is worth noting that $\mathbb{X}_0$ should be chosen to be small in order to relax (4.24), but not too small since in that case a large N_2 would be required so that the terminal condition $x_{(0, N+N_2)} \in \mathbb{X}_0$ is met.

Remark 4.5. The presentation in this section assumes that $N \geq 2$ as introduced in Remark 4.1. The results shown here, and those that will be presented below, can easily

be extended to $N = 1$ by noting that in this case $\bar{Y}_{(\infty)} = (I, K)\mathbb{X}_0 \oplus \bigoplus_{j=1}^{N_2} (I, K)\Phi^{j-1}\mathbb{W} \oplus (I, K)\Omega_{\infty}(\mathbb{X}_1)$.

4.4 Striped PTMPC: Online implementation and properties - Two cases

The cost used in [73] and in Chapter 3 is not useful for the SPTMPC prediction scheme; stability cannot be proved by using that cost in the online optimal control problem. The reason that such a cost is useful for PTMPC and OPTMPC is that $u = Kx$ is used in Mode 2, thus the candidate solutions at time $t+1$ can be obtained from the optimal solution at time t , according to the *receding* procedure as explained in Section 3.3, by discarding the 0-step-ahead predictions made at time t and completing the $N-1$ -step-ahead predicted input with an extension given by $u = Kx$ (recall that the offline update of the slacks in Chapter 3 replicates this process). This allows one to obtain a decreasing cost, where the cost acts as a measure of distance to the set where the control law $u = Kx$ is feasible. In the SPTMPC setting, however, since the degrees of freedom of the control policy extend over Mode 2, replacing the N -step-ahead predicted input by $u = Kx$ is not possible unless all the degrees of freedom of the policy are modified. Then this cost will not necessarily be decreasing and stability cannot be proved.

Two variations of the SPTMPC strategy are considered next that provide different stability results. The first considers an additional terminal constraint by which the predictions beyond $N + N_2$ must be inside a robustly invariant set under the static feedback $u = Kx$, and a cost function that is based on the distance of the predictions to this terminal set. The strategy can be proved to be robustly exponentially stable with respect to the terminal set, and once inside the maximal invariant set under the static feedback $u = Kx$, the control law is chosen to be $u = Kx$ which makes the

system converge to the minimal invariant set under $u = Kx$. This extra constraint, however, can incur some conservativeness. The second alternative uses a nominal cost and does not consider this terminal constraint with the aim to enhance the freedom of the predicted control policy, but this comes at the cost of obtaining a weaker notion of stability, namely ISS. A detailed description of both alternatives is presented next.

4.4.1 Cost based on state distance to an invariant set - strong stability

Let $\mathbb{X}_f = \{x : H^f x \leq 1\}$, $H^f \in \mathbb{R}^{n_f \times n}$, be a positively robust invariant set for system (4.1) subject to (4.3) with the control law given by $u = Kx$.

Consider the additional constraint that the predicted states for $k \geq N + N_2$ are inside $\mathbb{X}_f$. This can be expressed as

$$\text{Proj}_x(\bar{Y}_{(\infty)}) = \mathbb{X}_0 \oplus \bigoplus_{j=1}^{N-1} X_{(1,j)} \oplus \bigoplus_{j=N}^{N+N_2-1} \Phi^{j-N} X_{(1,N)} \oplus \Omega_{\infty}(\mathbb{X}_1) \subseteq \mathbb{X}_f. \quad (4.34)$$

Similarly to (4.27), this condition can be written equivalently as linear constraints by using slack variables that account for the worst case realizations associated with the 1st partial extreme state sequences. Condition (4.34) is then equivalent to

$$h_{(l,0,\infty)} + \sum_{j=1}^{N+N_2-1} h_{(l,1,j)} + h_{(l,1,\infty)} \leq 1, \quad l = 1, \dots, n_f, \quad (4.35a)$$

$$h_{(l,1,j)} \geq H_l^f x_{(i,1,j)}, \quad i = 1, \dots, q, \quad j = 1, \dots, N, \quad l = 1, \dots, n_f, \quad (4.35b)$$

$$h_{(l,1,j)} \geq H_l^f \Phi^{j-N} x_{(i,1,N)}, \quad i = 1, \dots, q, \quad j = N + 1, \dots, N + N_2 - 1, \quad l = 1, \dots, n_f, \quad (4.35c)$$

where

$$h_{(l,0,\infty)} = \max_{x \in \mathbb{X}_0} H_l^f x, \quad h_{(l,1,\infty)} = \max_{x \in \Omega_{\infty}(\mathbb{X}_1)} H_l^f x. \quad (4.36)$$

Let $\mathbf{d}$ be the vector of all the decision variables in the formulation. It is composed of $\mathbf{x}_{(0,N)}$, $\mathbf{x}_{(i,1,N)}$, $\mathbf{u}_{(0,N-1)}$, $\mathbf{u}_{(i,1,N-1)}$, $\mathbf{f}_{(l,1,N+N_2-1)} = \{f_{(l,1,1)}, \dots, f_{(l,1,N+N_2-1)}\}$ and $\mathbf{h}_{(l,1,N+N_2-1)} = \{h_{(l,1,1)}, \dots, h_{(l,1,N+N_2-1)}\}$. The set of admissible decision variables $\mathbf{d}$ is given by

$$\mathcal{D}(x) = \{\mathbf{d} : (4.16a), (4.16b), (4.17a), (4.17b), (4.29), (4.30), (4.32), (4.35) \text{ hold}\}, \quad (4.37)$$

and the domain of attraction is

$$\mathcal{X} = \{x : \mathcal{D}(x) \neq \emptyset\} \quad (4.38)$$

Assumption 4.3. N , N_2 , $\mathbb{X}_0$ and $\mathbb{X}_f$ are such that $\mathcal{X} \neq \emptyset$ and $\mathcal{X}$ contains the origin in its interior.

Remark 4.6. This assumption can always be satisfied provided that Assumption 4.1 is satisfied. Set all the predicted inputs to be given by $u = Kx$, then $\bar{Y}_{(\infty)} = (I, K)\mathbb{X}_0 \oplus \bigoplus_{j=1}^{N+N_2-1} (I, K)\Phi^{j-1}\mathbb{W} \oplus (\Omega_{\infty}(\mathbb{X}_1), K\Omega_{\infty}(\mathbb{X}_1)) = (I, K)\mathbb{X}_0 \oplus \bigoplus_{j=0}^{\infty} (I, K)\Phi^j\mathbb{W}$, and $\bar{Y}_{(\infty)} = (I, K)\Omega_{\infty}(\mathbb{W})$ if $\mathbb{X}_0 = \{0\}$. But $(I, K)\Omega_{\infty}(\mathbb{W}) \in \text{interior}(\mathbb{Y})$, which implies that there exists a robustly invariant set Ω for system (4.1) subject to (4.3) with the control law given by $u = Kx$, such that $\Omega_{\infty}(\mathbb{W}) \subseteq \text{interior}(\Omega)$. Therefore there exists an $\bar{N}$ such that for all $N + N_2 \geq \bar{N}$, $\mathbb{X}_1 = \Phi^{N+N_2-1}\mathbb{W}$ is small enough so that $\mathbb{X}_0$ can be chosen to be small enough, according to Assumption 4.2, such that $(I, K)\mathbb{X}_0 \oplus \bigoplus_{j=1}^{N+N_2-1} (I, K)\Phi^{j-1}\mathbb{W} \oplus (I, K)\Omega_{\infty}(\mathbb{X}_1) \subseteq (I, K)\Omega \subseteq \mathbb{Y}$. This guarantees that for any $x \in \mathbb{X}_0$, N, N_2 such that $N + N_2 \geq \bar{N}$ and $\mathbb{X}_f = \Omega$, a solution given by $\mathbf{x}_{(0,N)} = \{x, \dots, \Phi^N x\}$, $\mathbf{u}_{(0,N-1)} = \{Kx, \dots, K\Phi^N x\}$, $\mathbf{x}_{(i,1,N)} = \{\tilde{w}_i, \dots, \Phi^{N-1}\tilde{w}_i\}$, $\mathbf{u}_{(i,1,N-1)} = \{K\tilde{w}_i, \dots, K\Phi^{N-2}\tilde{w}_i\}$ will be feasible, which guarantees that $\mathbb{X}_0 \subseteq \mathcal{X}$. $\mathbb{X}_0$ contains the origin in its interior since it is an invariant set for a contractive system, which implies that $\mathcal{X}$ contains the origin in its interior, and then Assumption 4.3 is

satisfied.

Consider a polytopic set defined by $Z = \{z : Ez \leq \underline{1}\}$, $E \in \mathbb{R}^{n_E \times n}$. The distance function of x to Z is defined by $\text{dist}_Z(x) = \max(\{E_l x - 1 : l = 1, \dots, n_E\} \cup \{0\})$, where E_l is the l -th row of E . The maximum distance function of X to Z is defined by $\text{maxdist}_Y(X) = \max_{x \in X}(\text{dist}_Z(x)) = \max(\{E_l x - 1 : x \in X, l = 1, \dots, n_E\} \cup \{0\})$.

Consider the cost function

$$V(\mathbf{d}) = \sum_{k=0}^{N+N_2-1} \text{maxdist}_{\mathbb{X}_f}(X_{(k)}). \quad (4.39)$$

Here, the stage cost $\text{maxdist}_{\mathbb{X}_f}(X_{(k)})$ measures the maximum distance to $\mathbb{X}_f$ among the elements in $X_{(k)}$. Note that $\text{maxdist}_{\mathbb{X}_f}(X_{(k)}) = 0$ for all $k \geq N + N_2$ due to constraint (4.34). Components of the type $\text{maxdist}_{K\mathbb{X}_f}(U_{(k)})$ can also be considered in order to include weighing of the control sequences. In this case constraint (4.34) should be modified so that $\bar{Y}_{(\infty)} \subseteq (I, K)\mathbb{X}_f$ is required in order to satisfy the same properties as are stated below. However, the presentation here does not include such modifications for simplicity.

The optimal control problem $\mathbb{P}(x_t)$ to be solved at each instant t is defined by

$$V^0(x_t) = \min_{\mathbf{d}} V(\mathbf{d}), \text{ s.t. } \mathbf{d} \in \mathcal{D}(x_t). \quad (4.40a)$$

$$\mathbf{d}^*(x_t) = \arg \min_{\mathbf{d}} V(\mathbf{d}), \text{ s.t. } \mathbf{d} \in \mathcal{D}(x_t). \quad (4.40b)$$

By using new variables δ_k that tightly upper bound $\text{maxdist}_{\mathbb{X}_f}(X_{(k)})$, the minimization of the cost $J(\mathbf{d})$ of (4.39) can be reformulated by minimizing

$$V(\mathbf{d}) = \sum_{k=0}^{N+N_2-1} \delta_k, \quad (4.41)$$

with the additional constraints

$$\begin{aligned} \delta_k &\geq 0, & k &= 0, \dots, N + N_2 - 1 \\ \delta_k &\geq -1 + H_l^f x_{(0,k)} + \sum_{j=1}^k h_{(l,1,j)}, & l &= 1, \dots, n_f, \quad k = 0, \dots, N + N_2 - 1. \end{aligned} \quad (4.42)$$

Clearly the optimal δ_k are equal to $\text{dist}_{\mathbb{X}_f}(X_{(k)})$ so the minimization of (4.39) and of (4.41) are equivalent. The equivalent LP form of $\mathbb{P}(x_t)$ is then given by

$$\begin{aligned} &\underset{\mathbf{d}_O}{\text{minimize}} && \sum_{k=0}^{N+N_2-1} \delta_k \\ \text{s.t.} &&& x_{(0,0)} = x_t, \quad x_{(i,1,1)} = \tilde{w}_i, \quad i = 1, \dots, q, \\ &&& x_{(0,k+1)} = Ax_{(0,k)} + Bx_{(0,k)}, \quad k = 0, \dots, N-1, \\ &&& x_{(i,1,k+1)} = Ax_{(i,1,k)} + Bu_{(i,1,k)}, \quad i = 1, \dots, q, \quad k = 0, \dots, N-1, \\ &&& F_l x_{(0,0)} + G_l u_{(0,0)} \leq 1, \quad l = 1, \dots, n_c, \\ &&& F_l x_{(0,k)} + G_l u_{(0,k)} + \sum_{j=1}^k f_{(l,1,j)} \leq 1, \quad l = 1, \dots, n_c, \quad k = 1, \dots, N-1, \\ &&& (F_l + G_l K) \Phi^{k-N} x_{(0,N)} + \sum_{j=1}^k f_{(l,1,j)} \leq 1, \quad l = 1, \dots, n_c, \\ &&& k = N, \dots, N + N_2 - 1, \\ &&& f_{(l,0,\infty)} + \sum_{j=1}^{N+N_2-1} f_{(l,1,j)} + f_{l,1,\infty} \leq 1, \quad l = 1, \dots, n_c, \\ &&& h_{(l,0,\infty)} + \sum_{j=1}^{N+N_2-1} h_{(l,1,j)} + h_{(l,1,\infty)} \leq 1, \quad l = 1, \dots, n_f, \\ &&& H_l x_{(0,N+N_2)} \leq 1, \quad l = 1, \dots, n_{f0}, \\ &&& I_l x_{(i,1,N+N_2)} \leq 1, \quad i = 1, \dots, q, \quad l = 1, \dots, n_{f1}, \\ &&& f_{(l,1,j)} \geq F_l x_{(i,1,j)} + G_l u_{(i,1,j)}, \quad l = 1, \dots, n_c, \quad i = 1, \dots, q, \quad j = 1, \dots, N-1, \\ &&& f_{(l,1,j)} \geq (F_l + G_l K) \Phi^{j-N} x_{(i,1,N)} + G_l u_{(i,1,j)}, \quad l = 1, \dots, n_c, \quad i = 1, \dots, q, \end{aligned}$$

$$\begin{aligned}
j &= N, \dots, N + N_2 - 1, \\
h_{(l,1,j)} &\geq H_l^f x_{(i,1,j)}, & l &= 1, \dots, n_f, \ i = 1, \dots, q, \\
j &= 1, \dots, N, \\
h_{(l,1,j)} &\geq H_l^f \Phi^{j-N} x_{(i,1,N)}, & l &= 1, \dots, n_f, \ i = 1, \dots, q, \\
j &= N + 1, \dots, N + N_2 - 1, \\
\delta_k &\geq 0, & k &= 0, \dots, N + N_2 - 1, \\
\delta_k &\geq -1 + H_l^f x_{(0,k)} + \sum_{j=1}^k h_{(l,1,j)}, & l &= 1, \dots, n_f, \ k = 0, \dots, N + N_2 - 1,
\end{aligned} \tag{4.43}$$

where the overall decision variable $\mathbf{d}_O$ comprises $\mathbf{d}$ and the additional variables $\boldsymbol{\delta} = \{\delta_0, \dots, \delta_{N+N_2-1}\}$. This is a linear program with $N_{var} = N(n + (1 + q) + n_u(1 + q) + n_c + n_f + 1) + N_2(n_c + n_f + 1) + n - (qn_u + n_c + n_f + 1)$ variables and $N_{cons} = N((n + n_c + n_f)(1 + q) + 1) + N_2((n_c + n_f)(1 + q) + 1) + n + n_c + n_{f0} + n_{f1}q + n_f(1 - q)$ constraints (of which $Nn(1 + q) + n$ are equality constraints and $N((n_c + n_f)(1 + q) + 1) + N_2((n_c + n_f)(1 + q) + 1) + n_c + n_{f0} + n_{f1}q + n_f(1 - q)$ are inequality constraints).

Remark 4.7. Setting $\mathbb{X}_f = \bar{\Omega}$, where $\bar{\Omega}$ is the maximal invariant set for system (4.1) under the control law of $u = Kx$ s.t. (4.3), is convenient for relaxing (4.34). On the other hand, invariant sets defined by few inequalities are computationally convenient in view of (4.34), (4.42). Thus a suitable compromise between the size of $\mathbb{X}_f$ and n_f should be chosen.

The SPTMPC control law is given by

$$\kappa^0(x_t) = u_{(0,0)}^0(x_t), \tag{4.44}$$

where $u_{(0,0)}^0(x_t)$ is part of $\mathbf{d}^0(x_t)$, which is a selection of $\mathbf{d}^*(x_t)$; since $\mathbb{P}(x_t)$ is an LP, it does not necessarily have a unique solution, and then $\mathbf{d}^*(x_t)$ might be a set instead

of a singleton. The control input at time t is then given by $u_t := \kappa^0(x_t)$.

Theorem 4.8. *For system (4.1), the closed-loop application of the control law of (4.44) guarantees feasibility of (4.3), that the optimal control problem of (4.40) is recursively feasible and that $\mathcal{X}$ is invariant.*

Proof. The candidate solution at time $t + 1$, $\tilde{\mathbf{d}}$, is built from the optimal solution at time t , $\mathbf{d}^0(x_t)$ as follows. The candidate 0^{th} partial sequences are composed by the optimal 0^{th} partial sequences of $\mathbf{d}^0(x_t)$ plus a sequence of elements in the 1^{st} partial sequences of $\mathbf{d}^0(x_t)$ according to the *receding* procedure as presented in Section 3.3. Then, the candidate 0^{th} partial sequences are defined as

$$\tilde{x}_{(0,k)} = x_{(0,k+1)}^0(x_t) + \sum_{i=1}^q \lambda_{(i,1)}^*(w_t) x_{(i,1,k+1)}^0(x_t), \quad k = 0, \dots, N-1, \quad (4.45a)$$

$$\tilde{x}_{(0,N)} = \Phi \tilde{x}_{(0,N-1)},$$

$$\tilde{u}_{(0,k)} = u_{(0,k+1)}^0(x_t) + \sum_{i=1}^q \lambda_{(i,1)}^*(w_t) u_{(i,1,k+1)}^0(x_t), \quad k = 0, \dots, N-2, \quad (4.45b)$$

$$\tilde{u}_{(0,N-1)} = K \tilde{x}_{(0,N-1)}(x_{t+1}),$$

where $\lambda_{(i,1)}^*(w_t)$ are the least squares (or any other convenient selection criterion) convex interpolation parameters that satisfy $w_t = \sum_{i=1}^q \lambda_{(i,1)}^*(w_t) \tilde{w}_i$. The candidate 1^{st} partial extreme sequences are, because of the striped structure, chosen to be the same as in the optimal solution at time t

$$\tilde{x}_{(i,1,k)} = x_{(i,1,k)}^0(x_t), \quad k = 1, \dots, N, \quad i = 1, \dots, q \quad (4.46a)$$

$$\tilde{u}_{(i,1,k)} = u_{(i,1,k)}^0(x_t), \quad k = 1, \dots, N-1, \quad i = 1, \dots, q, \quad (4.46b)$$

This construction is such that (4.16a),(4.16b),(4.17a),(4.17b) are satisfied, and that the overall cross sections associated with the candidate sequences $\tilde{Y}_{(k)} = \tilde{Y}_{(0,k)} \oplus \bigoplus_{j=1}^k \tilde{Y}_{(1,j)}$, where $\tilde{Y}_{(0,k)} = (\tilde{x}_{(0,k)}, \tilde{u}_{(0,k)})$ and $\tilde{Y}_{(1,j)} = \text{conv}\{(\tilde{x}_{(i,1,k)}, \tilde{u}_{(i,1,k)}) : i = 1, \dots, q\} = Y_{(1,j)}^0(x_t)$, satisfy $\tilde{Y}_{(k)} \subseteq Y_{(k+1)}^0(x_t)$, $k = 0, 1, \dots$. But $\mathbf{d}^0(x_t)$ is such that $Y_{(k+1)}^0(x_t) \subseteq \mathbb{Y}$, $k = 0, 1, \dots$, which implies that $\tilde{Y}_{(k)} \subseteq \mathbb{Y}$, $k = 0, 1, \dots, N + N_2 - 1$

and then $\tilde{\mathbf{d}}$ satisfies (4.29). Also, because $\tilde{Y}_{(1,j)} = Y_{(1,j)}(x_t)$ it follows that $\tilde{Y}_{(\infty)} = \bar{Y}_{(\infty)}^0(x_t)$ (where $\tilde{Y}_{(\infty)}$ and $\bar{Y}_{(\infty)}^0(x_t)$ are given by (4.25) but using the variables of $\tilde{\mathbf{d}}$ and $\mathbf{d}^0(x_t)$, respectively), and since $\mathbf{d}^0(x_t)$ is such that $\bar{Y}_{(\infty)}^0(x_t) \subseteq \mathbb{Y}$ it follows that $\tilde{Y}_{(\infty)} \subseteq \mathbb{Y}$ and then $\tilde{\mathbf{d}}$ satisfies (4.30).

Satisfaction of the terminal condition (4.23b) follows directly from the fact that $\tilde{Y}_{(1,j)} = Y_{(1,j)}^0(x_t)$, from which $\tilde{X}_{(1,N+N_2)} = X_{(1,N+N_2)}^0 \subseteq \mathbb{X}_1$, where $X_{(1,N+N_2)}^0 = \text{conv}\{\Phi^{N_2} x_{(i,1,N)}^0 : i = 1, \dots, q\}$. The satisfaction of the terminal condition of (4.23a) follows from the fact that $\tilde{x}_{(0,N+N_2)} = \Phi x_{(0,N+N_2)}^0(x_t) + \Phi x_{(1,N+N_2)}$ with $x_{(1,N+N_2)} \in X_{(1,N+N_2)}^0 \subseteq \Phi \mathbb{X}_1$ and that $\mathbb{X}_0$ is invariant for $z^+ = \Phi z + w$, where $w \in \Phi \mathbb{X}_1$. Finally, since $\tilde{Y}_{(1,j)} = Y_{(1,j)}^0(x_t)$, and $\text{Proj}_x(\bar{Y}_{(\infty)}^0(x_t)) \subseteq \mathbb{X}_f$, it follows that $\text{Proj}_x(\tilde{Y}_{(\infty)}) \subseteq \mathbb{X}_f$, and then (4.35) is satisfied. Note that constraint (4.42) is always satisfied since it is instrumental in the evaluation of the cost of (4.39).

It follows that if $\mathbb{P}(x_t)$ is feasible, it will be feasible at the next instant, and by a recursive argument it will remain feasible for all future instants. Invariance of $\mathcal{X}$ follows directly. Satisfaction of (4.3) follows from the fact that $Y_{(0)}^0(x_t) \subseteq \mathbb{Y}$. $\square$

Remark 4.9. *The treatment of recursive feasibility, as shown above, is more involved than in usual RMPC formulations, where a terminal set for x_N that is robustly invariant for the fixed terminal control law is enough. This is because allowing the disturbance compensation to extend over an infinite horizon implies that the terminal control law is not fixed. In this different setting two separate terminal sets for the nominal and the partial extreme predictions are used, $\mathbb{X}_0$ and $\mathbb{X}_1$, and as shown in the proof of Theorem 4.8 it is required that $\mathbb{X}_0$ is robustly invariant for $z^+ = \Phi z + w$, with $w \in \Phi \mathbb{X}_1$, thus justifying Assumption 4.2.*

Proposition 4.10. (i) $\mathcal{D}(x)$ is a closed polyhedral set. (ii) $\mathcal{X}$ is a polytopic set.

Proof. (i) $\mathcal{D}(x) = \{\mathbf{d} : M_{\mathbf{d}}\mathbf{d} + M_x x \leq M\}$, for some matrices $M_{\mathbf{d}}, M_x, M$ of appropriate dimensions, is an alternative way of expressing $\mathcal{D}(x)$, from which it follows that

$\mathcal{D}(x)$ is closed a polyhedron. (ii) Let $\mathcal{Z} = \{(x, \mathbf{d}) : M_{\mathbf{d}}\mathbf{d} + M_x x \leq M\}$. $\mathcal{Z}$ is clearly a closed polyhedron, and since $\mathcal{X}$ is a projection of $\mathcal{Z}$ onto a suitable subspace, $\mathcal{X}$ is also a closed polyhedron. If $x \in \mathcal{X}$ then there exists $u_{(0,0)}$ such that $(x_{(0,0)}, u_{(0,0)}) \in \mathbb{Y}$, with $x_{(0,0)} = x$. If $\mathcal{X}$ were not bounded then there would exist at least one $x \neq 0$ such that $\mu x \in \mathcal{X}$, for all $\mu \in \mathbb{R}$, but this contradicts the boundedness of $\mathbb{Y}$. Thus $\mathcal{X}$ is bounded. $\square$

Proposition 4.11. *For the optimal control problem $\mathbb{P}(x)$:*

(i) *There exists a piecewise affine and continuous function $\mathbf{d}^0(\cdot) : \mathcal{X} \rightarrow \mathbb{R}^{N_{\mathbf{d}}}$ such that for all $x \in \mathcal{X}$, $\mathbf{d}^0(x) \in \mathbf{d}^*(x)$ ($\mathbf{d}^0(x)$ being a selection of $\mathbf{d}^*(x)$).*

The value function $V^0(x)$, has the following properties:

(ii) *$V^0(\cdot)$ is a convex, continuous and piecewise affine function.*

(iii) *$V^0(x) > 0$, for all $x \in \mathcal{X} \setminus \mathbb{X}_f$, for $t = 0, 1, \dots$*

(iv) *$V^0(x) = 0$, for all $x \in \mathbb{X}_f$.*

Proof. (i) The optimal control problem $\mathbb{P}(\cdot)$ is a linear parametric program in $\mathcal{X}$, then by standard results in parametric optimization problems [5] the stated property holds.

(ii) The function $V^0(\cdot)$ is the value function of a parametric linear program in the bounded set $\mathcal{X}$. Then it is a convex, piecewise affine and continuous function in x (by standard results in parametric optimization problems [5]).

(iii) If $x \notin \mathbb{X}_f$, then $\text{dist}_{\mathbb{X}_f}(x_{(0,0)}^0(x)) > 0$, hence $V^0(x) > 0$.

(iv) If $x \in \mathbb{X}_f$, then the solution given by $\mathbf{x}_{(0,N)} = \{x, \Phi x, \dots, \Phi^{N-1}x\}$, $\mathbf{u}_{(0,N-1)} = \{Kx, K\Phi x, \dots, K\Phi^{N-1}x\}$, $\mathbf{x}_{(i,1,N)} = \{\tilde{w}_i, \Phi\tilde{w}_i, \dots, \Phi^{N-1}\tilde{w}_i\}$, and $\mathbf{u}_{(i,1,N-1)} = \{K\tilde{w}_i, K\Phi\tilde{w}_i, \dots, K\Phi^{N-1}\tilde{w}_i\}$ will be feasible, and since $\mathbb{X}_f$ is invariant for Kx , $X_{(k)}$ will be in $\mathbb{X}_f$ for all $k = 0, 1, \dots$ $\square$

Theorem 4.12. *The value function $V^0(x_t)$ satisfies*

$$V^0(x_{t+1}) - V^0(x_t) \leq -\text{dist}_{\mathbb{X}_f}(x_t), \quad t = 0, 1, \dots \quad (4.47)$$

Proof. Provided that $\mathbb{P}(x_t)$ is feasible at time t , the candidate solution $\tilde{\mathbf{d}}(x_{t+1})$ defined by (4.45) and (4.46) is feasible for $\mathbb{P}(x_{t+1})$ at time $t+1$, as shown in Theorem (4.18). This candidate solution satisfies that $\tilde{X}_{(k)}(x_{t+1}) \subseteq X_{(k+1)}^0(x_t)$, for $k = 0, 1, \dots$, which implies that $\text{dist}_{\mathbb{X}_f}(\tilde{X}_{(k)}(x_{t+1})) \leq \text{dist}_{\mathbb{X}_f}(X_{(k+1)}^0(x_t))$, and then

$$\begin{aligned} V(\tilde{\mathbf{d}}(x_{t+1})) &= \sum_{k=0}^{N+N_2-1} \text{dist}_{\mathbb{X}_f}(\tilde{X}_{(k)}(x_{t+1})) \\ &\leq \sum_{k=1}^{N+N_2} \text{dist}_{\mathbb{X}_f}(X_{(k)}^0(x_t)) = V^0(x_t) - \text{dist}_{\mathbb{X}_f}(X_{(0)}^0(x_t)). \end{aligned}$$

Finally, the result follows from the fact that optimization ensures that $\delta_0^0(x_t) = \text{dist}_{\mathbb{X}_f}(x_t)$ and $V^0(x_{t+1}) \leq V(\tilde{\mathbf{d}}(x_{t+1}))$. $\square$

Corollary 4.13. (i) The value function $V^0(x_t) \rightarrow 0$ as $t \rightarrow \infty$. Furthermore, $V^0(x_t) \leq a^t V^0(x_0)$, with $a \in (0, 1)$. (ii) The distance function $\text{dist}_{\mathbb{X}_f}(x_t) \rightarrow 0$ as $t \rightarrow \infty$. Furthermore, $\text{dist}_{\mathbb{X}_f}(x_t) \leq b^t \text{dist}_{\mathbb{X}_f}(x_0)$, with $b \in (0, 1)$.

Proof. (i) By construction it is clear that $\text{dist}_{\mathbb{X}_f}(x_{(0,0)}^0(x)) \leq V^0(x)$. In addition, $V^0(x) = 0$ and $\text{dist}_{\mathbb{X}_f}(x_{(0,0)}^0(x)) = 0$ for $x \in \mathbb{X}_f$, and $V^0(x) > 0$ and $\text{dist}_{\mathbb{X}_f}(x_{(0,0)}^0(x)) > 0$ for $x \in \mathcal{X} \setminus \mathbb{X}_f$. Then, based on the continuity of $V^0(\cdot)$ and $\text{dist}_{\mathbb{X}_f}(x_{(0,0)}^0(x))$, and that $\mathbb{X}_f$ and $\mathcal{X}$ are compact, there exists a scalar $c_1 \in [1, \infty)$ such that $V^0(x) \leq c_1 \text{dist}_{\mathbb{X}_f}(x_{(0,0)}^0(x))$ for all $x \in \mathcal{X}$. Combining this with (4.47) it follows that $V^0(x_{t+1}) \leq a_N V^0(x_t)$, with $a_N = (c_1 - 1)/c_1 < 1$, which yields the desired property.

(ii) Similarly to (i), it is possible to obtain that $\text{dist}_{\mathbb{X}_f}(x_t) = \text{dist}_{\mathbb{X}_f}(x_{(0,0)}^0(x_t)) \leq c_1((c_1 - 1)/c_1)^t \text{dist}_{\mathbb{X}_f}(x_{(0,0)}^0(x_0)) = ((c_1 - 1)/c_1)^t \text{dist}_{\mathbb{X}_f}(x_0)$. $\square$

Corollary 4.13 shows the state converges exponentially fast to $\mathbb{X}_f \subseteq \bar{\Omega}$. Once the state reaches $\bar{\Omega}$, the control law can be chosen to revert to $u = Kx$, which will remain feasible given the nature of $\bar{\Omega}$, and this will steer the state asymptotically to the minimal robustly invariant set $\Omega_\infty(\mathbb{W})$. The SPTMPC control law is then modified

to be given by

$$\kappa^0(x_t) = \begin{cases} u_{(0,0)}^0(x_t), & \text{if } x_t \notin \bar{\Omega} \\ Kx_t, & \text{if } x_t \in \bar{\Omega}, \end{cases} \quad (4.48)$$

Finally, the dependence of the domain of attraction $\mathcal{X}$ on the horizons N and N_2 , is studied; this dependence is now made explicit by denoting the domain of attraction as $\mathcal{X}_{N,N_2}$. $\mathbb{X}_1$ also depends on N and N_2 , so the dependence is also made explicit by writing $\mathbb{X}_1(N, N_2)$. The desired properties are $\mathcal{X}_{N,N_2} \subseteq \mathcal{X}_{N,N_2+1}$ and $\mathcal{X}_{N,N_2} \subseteq \mathcal{X}_{N+1,N_2}$. The following theorem presents an alternative way to construct $\mathbb{X}_0$ such that the desired nestedness properties are satisfied.

Theorem 4.14. *Let $\mathbb{X}_0$ be a robust positively invariant set for system $z^+ = \Phi z + w$, with $w \in \mathcal{W}$, where $\mathcal{W}$ is any outer approximation of $\text{conv}(\Phi^{N_3}\mathbb{W} \cup \Phi^{N_3+1}\mathbb{W} \cup \Phi^{N_3+2}\mathbb{W} \dots)$, $N_3 > 0$. Then the domain of attraction of system (4.1) under the control law of (4.44) satisfies the following properties for all N, N_2 such that $N + N_2 \geq N_3$*

- (i) $\mathcal{X}_{N,N_2} \subseteq \mathcal{X}_{N+1,N_2-1}$,
- (ii) $\mathcal{X}_{N,N_2} \subseteq \mathcal{X}_{N,N_2+1}$,
- (iii) $\mathcal{X}_{N,N_2} \subseteq \mathcal{X}_{N+1,N_2}$.

Proof. (i) The optimal control problem $\mathbb{P}(x)$ when using N and N_2 is the same one as when using $\tilde{N} = N + 1$ and $\tilde{N}_2 = N_2 - 1$, except for the fact that in the former case $u_{(0,N)}$ and $u_{(i,1,N)}$ are fixed to be $u_{(0,N)} = Kx_{(0,N)}$ and $u_{(i,1,N)} = Kx_{(i,1,N)}$, but in the latter case these are degrees of freedom. In the latter case $u_{(0,N)}$ and $u_{(i,1,N)}$ can be chosen to be given by $u_{(0,N)} = Kx_{(0,N)}$ and $u_{(i,1,N)} = Kx_{(i,1,N)}$, so that the problem for $\tilde{N}$ and $\tilde{N}_2$ is feasible when it is feasible for N and N_2 .

(ii) If $x \in \mathcal{X}_{N,N_2}$, the optimal $\mathbf{d}$ that includes $\mathbf{x}_{(0,N)}$, $\mathbf{u}_{(0,N-1)}$, $\mathbf{x}_{(i,1,N)}$, $\mathbf{u}_{(i,1,N-1)}$, $\mathbf{f}_{(1,1,N+N_2-1)}$ and $\mathbf{h}_{(1,1,N+N_2-1)}$, is such that (4.16a), (4.16b), (4.17a), (4.17b), (4.29), (4.30), (4.32), (4.35) are satisfied. Then $\tilde{\mathbf{d}}$ that is composed of the variables in $\mathbf{d}$ with the addition of $f_{(l,1,N+N_2)} = \max\{(F_l + G_l K)\Phi^{N_2}x_{(i,1,N)} : i = 1, \dots, q\}$, $l = 1, \dots, n_c$,

and $h_{(l,1,N+N_2)} = \max_{i \in \mathbb{N}_{[1,q]}} \{(H_l^f \Phi^{N_2} x_{(i,1,N)} : i = 1, \dots, q)\}$, $l = 1, \dots, n_f$, satisfies (4.16a), (4.16b), (4.17a), (4.17b) for N and $\tilde{N}_2 = N_2 + 1$ directly since none of these constraints depend on N_2 . The definition of $\mathbb{X}_1$ implies that $\mathbb{X}_1(N, N_2 + 1) = \Phi \mathbb{X}_1(N, N_2)$. Then, since $\Phi^{N_2} x_{(i,1,N)} \subseteq \mathbb{X}_1(N, N_2)$ (from the fact that (4.32a) is satisfied for N, N_2) it follows $\Phi^{N_2+1} x_{(i,1,N)} \subseteq \mathbb{X}_1(N, N_2 + 1)$ and so (4.32b) is satisfied for $N, \tilde{N}_2$. The definition of $\mathbb{X}_0$ implies that it is the same for different values of N and/or N_2 and it satisfies $\Phi \mathbb{X}_0 \subseteq \mathbb{X}_0$. Then, since (4.32b) is satisfied for N, N_2 it follows that $x_{(0,N+N_2)} \in \mathbb{X}_0$ which implies that $x_{(0,N+N_2+1)} = \Phi x_{(0,N+N_2)} \in \mathbb{X}_0$, and thus (4.32a) is satisfied for $N, \tilde{N}_2$. Satisfaction of (4.29) and (4.30) for N, N_2 implies that $Y_{(k)} \subseteq \mathbb{Y}$, $k = 0, \dots, N + N_2 - 1$, and that $\bar{Y}_{(\infty)}(N, N_2) \in \mathbb{Y}$ (where the dependence of $\bar{Y}_{(\infty)}$ on N, N_2 has been made explicit), which implies that $Y_{(k)} \subseteq \mathbb{Y}$, $k = 0, \dots, N + N_2$, thus ensuring satisfaction of (4.29) for $N, \tilde{N}_2$. Eq. (4.33) implies that $\Omega_\infty(\mathbb{X}_1(N, N_2)) = \Omega_\infty(\mathbb{X}_1(N, N_2 + 1)) \oplus \Phi^{N_2} X_{(1,N)}$, which implies that $Y_{(\infty)}(N, N_2 + 1) = (I, K) \mathbb{X}_0(N, N_2) \oplus \bigoplus_{j=1}^{N-1} Y_{(1,j)} \oplus \bigoplus_{j=N}^{N+N_2} (I, K) \Phi^{j-N} X_{(1,N)} \oplus (I, K) \Omega_\infty(\mathbb{X}_1(N, N_2 + 1)) = \bar{Y}_{(\infty)}(N, N_2) \subseteq \mathbb{Y}$, and therefore (4.30) is satisfied for $N, \tilde{N}_2$ provided it is satisfied for N, N_2 . In the same way, satisfaction of (4.35) for N, N_2 implies that $\text{Proj}_x(\bar{Y}_{(\infty)}(N, N_2)) = \text{Proj}_x(\bar{Y}_{(\infty)}(N, N_2 + 1)) \subseteq \mathbb{X}_f$, thus ensuring satisfaction of (4.35) for $N, \tilde{N}_2$. All this implies that $x \in \mathcal{X}_{N, N_2+1}$, which proves the desired result.

(iii) This follows directly from combining (i) and (ii). $\square$

Remark 4.15. In Theorem 4.14, $\mathcal{W}$ is defined as any outer approximation of (instead of being equal to) $\text{conv}(\Phi^{N_3} \mathbb{W} \cup \Phi^{N_3+1} \mathbb{W} \cup \Phi^{N_3+2} \mathbb{W} \dots)$, in order to reduce the number of vertices or facets needed to define $\mathcal{W}$, and hence $\mathbb{X}_0$.

Remark 4.16. The condition that $\mathbb{X}_0$ is invariant for $z^+ = \Phi z + w$, with $w \in \mathcal{W}$, $\mathcal{W}$ being any outer approximation of $\text{conv}(\{\Phi^{N_3} \mathbb{W} \cup \Phi^{N_3+1} \mathbb{W} \cup \Phi^{N_3+2} \mathbb{W} \dots\})$, in addition to ensuring that $\mathbb{X}_0$ is the same for different values of N, N_2 , guarantees that $\mathbb{X}_0$ satisfies Assumption 4.2 for all N, N_2 such that $N + N_2 \geq N_3$, which is key for ensuring recursive feasibility. Then recursive feasibility is guaranteed by Theorem 4.8

whenever $N + N_2 \geq N_3$. Thus, N_3 should be chosen small enough as dictated by practical values of N, N_2 , but large enough to maintain $\mathbb{X}_0$ small.

4.4.2 Nominal cost and input-to-state stability

This section presents an alternative to the strategy of Section 4.4.1. This strategy uses a nominal type of cost and does not use the terminal condition of (4.34). On account of this, this variation of the strategy has more freedom, thus allows for greater domains of attraction and has a reduced number of variables and constraints.

The vector of all the decision variables in the formulation, $\mathbf{d}$, is redefined to be composed of $\mathbf{x}_{(0,N)}$, $\mathbf{x}_{(i,1,N)}$, $\mathbf{u}_{(0,N-1)}$, $\mathbf{u}_{(i,1,N-1)}$, and $\mathbf{f}_{(l,1,N+N_2-1)}$. The new admissible set of decision variables $\mathbf{d}$ is given by

$$\mathcal{D}(x) = \{\mathbf{d} : (4.16a), (4.16b), (4.17a), (4.17b), (4.29), (4.30), (4.32) \text{ hold}\}, \quad (4.49)$$

and the domain of attraction is described by $\mathcal{X} = \{x : \mathcal{D}(x) \neq \emptyset\}$.

Assumption 4.4. $N, N_2, \mathbb{X}_0$ are such that $\mathcal{X} \neq \emptyset$ and $\mathcal{X}$ contains the origin in its interior.

Remark 4.17. By the same analysis of Remark 4.6, this assumption can always be satisfied provided that Assumption 4.1 is satisfied.

The predicted cost to be optimized online, $V(\mathbf{d})$, only considers the nominal state and input sequences:

$$V(\mathbf{d}) = \sum_{k=0}^{N-1} x_{(0,k)}^T Q x_{(0,k)} + u_{(0,k)}^T R u_{(0,k)} + x_{(0,N)}^T P x_{(0,N)} \quad (4.50)$$

where $Q, R, P \succ 0$, satisfy the Lyapunov condition $\Phi^T P \Phi - P \leq -(Q + K^T R K)$.

$$\Phi^T P \Phi - P \leq -(Q + K^T R K). \quad (4.51)$$

The optimal control problem $\mathbb{P}(x_t)$ to be solved at each instant t is given by

$$V^0(x_t) = \min_{\mathbf{d}} V(\mathbf{d}), \text{ s.t. } \mathbf{d} \in \mathcal{D}(x_t). \quad (4.52a)$$

$$\mathbf{d}^0(x_t) = \arg \min_{\mathbf{d}} V(\mathbf{d}), \text{ s.t. } \mathbf{d} \in \mathcal{D}(x_t). \quad (4.52b)$$

The optimization of (4.52) is a QP with $N_{var} = N(n(1+q) + n_u(1+q) + n_c) + N_2n_c + n - (qn_u + n_c)$ variables and $N_{cons} = N(n(1+q) + n_c(1+q)) + N_2n_c(1+q) + n + n_c + n_{f0} + n_{f1}q$ constraints (of which $Nn(1+q) + n$ are equality constraints and $(N + N_2)n_c(1+q) + n_c + n_{f0} + n_{f1}q$ are inequality constraints).

The SPTMPC control law is then given by

$$\kappa^0(x_t) = u_{(0,0)}^0(x_t), \quad (4.53)$$

where $u_{(0,0)}^0(x_t)$ is part of the optimal solution given by $\mathbf{d}^0(x_t)$.

Theorem 4.18. *For system (4.1), the closed-loop application of the control law of (4.53) guarantees feasibility of (4.3), that the optimal control problem of (4.52) is recursively feasible and that $\mathcal{X}$ is robustly invariant.*

Proof. The results here are inherited directly from Theorem 4.8. Note that the absence of (4.35) does not affect the results since its feasibility at time t is not associated with the feasibility of any other constraints at time $t + 1$. $\square$

Theorem 4.19. *Let $\mathbb{X}_0$ be a robust positively invariant set for system $z^+ = \Phi z + w$, with $w \in \mathcal{W}$, where $\mathcal{W}$ is any outer approximation of $\text{conv}(\Phi^{N_3}\mathbb{W} \cup \Phi^{N_3+1}\mathbb{W} \cup \Phi^{N_3+2}\mathbb{W} \dots)$, $N_3 \in \mathbb{N}_+$. Then the domain of attraction of system (4.1) under the control law of (4.53) satisfies the following properties for all N, N_2 such that $N + N_2 \geq N_3$*

$$(i) \mathcal{X}_{N,N_2} \subseteq \mathcal{X}_{N+1,N_2-1},$$

$$(ii) \mathcal{X}_{N,N_2} \subseteq \mathcal{X}_{N,N_2+1},$$

$$(iii) \mathcal{X}_{N,N_2} \subseteq \mathcal{X}_{N+1,N_2}.$$

Proof. The proof goes along the same lines as in Theorem 4.14. Note that the absence of (4.35) does not affect the results since its feasibility is not associated with the feasibility of any other constraints. $\square$

As pointed out at the beginning of this section, $u = Kx$ is not the terminal control law in SPTMPC, and hence the stability notions of [73] do not apply. Neither does the analysis of Section 4.4.1, since it requires use of the cost function of (4.39) and condition (4.34). Instead Input-to-state stability (ISS) [41] is considered. Recall from Section 2.2.1 that ISS implies that: (i) the system is asymptotically stable in the absence of disturbances; (ii) that state trajectories are bounded for bounded disturbances; and (iii) as $t \rightarrow \infty$ all trajectories go to the origin if $w_t \rightarrow 0$. Also recall Lemma 2.8, which provides conditions for a system to be ISS.

Lemma 4.20. *Let $\mathcal{S} \subseteq \mathbb{R}^n$ contain the origin in its interior and be a robust positively invariant set for system $x^+ = f(x, w)$. Furthermore, let there exist $\mathcal{K}_\infty$ -functions $\alpha_1(\cdot)$, $\alpha_2(\cdot)$ and $\alpha_3(\cdot)$ and a function $V : \mathcal{S} \rightarrow \mathbb{R}_+$ that is Lipschitz continuous on $\mathcal{S}$ such that for all $x \in \mathcal{S}$,*

$$\begin{aligned} \alpha_1(\|x\|) &\leq V(x) \leq \alpha_2(\|x\|), \\ V(f(x, 0)) - V(x) &\leq -\alpha_3(\|x\|). \end{aligned} \tag{2.27}$$

Then $V(\cdot)$ is an ISS-Lyapunov function and system $x^+ = f(x, w)$ is ISS with domain of attraction $\mathcal{S}$ if either: (i) $f(\cdot, \cdot)$ is Lipschitz continuous on $\mathcal{S} \times \mathbb{W}$, or (ii) $f(x, w) := g(x) + w$, where $g : \mathcal{S} \rightarrow \mathbb{R}^n$ is continuous on $\mathcal{S}$.

The following analysis parallels that of [38] in order to prove that the closed-loop application of the control law of (4.53) provides input-to-state stability.

Proposition 4.21. *(i) $\mathcal{D}(x)$ is a closed polyhedral set. (ii) $\mathcal{X}$ is a polytopic set.*

Proof. The proof is the same as that of Proposition 4.10. $\square$

Proposition 4.22. $\kappa^0(x)$ is a unique Lipschitz continuous function for all $x \in \mathcal{X}$. Furthermore, the value function $V^0(x)$ is strictly convex and Lipschitz continuous for all $x \in \mathcal{X}$.

Proof. The sequence of predicted nominal states $\mathbf{x}_{(0,N)}$ can be interpreted as a function of the predicted nominal inputs $\mathbf{u}_{(0,N-1)}$ and the current state x_t . Hence $V(\mathbf{d})$, which does not depend on the uncertain extreme sequences $\mathbf{x}_{(i,1,N)}$, $\mathbf{u}_{(i,1,N-1)}$, can be interpreted as a function of $\mathbf{u}_{(0,N-1)}$ and x_t , i.e. $J(\mathbf{d}) = V(x_t, \mathbf{u}_{(0,N)})$. The set of admissible inputs is defined as $\mathcal{U}(x) = \{(\mathbf{u}_{(0,N-1)}) : \mathbf{d} = (\mathbf{x}_{(0,N)}, \mathbf{u}_{(0,N-1)}, \dots) \in \mathcal{D}(x)\}$ and then the optimal control problem can be written equivalently as

$$V^0(x) = \min_{\mathbf{u}_{(0,N-1)}} J(x_t, \mathbf{u}_{(0,N-1)}), \text{ s.t. } \mathbf{d} \in \mathcal{D}(x) \quad (4.54a)$$

$$\begin{aligned} (\mathbf{u}_{(0,N-1)})^0(x) &= \arg \min_{x, \mathbf{u}_{(0,N-1)}} J(x, \mathbf{u}_{(0,N-1)}), \\ \text{s.t. } (x, \dots, \mathbf{u}_{(0,N-1)}, \dots) &\in \mathcal{D}(x). \end{aligned} \quad (4.54b)$$

$\mathcal{D}(x)$ is polyhedral, and since $\mathcal{U}(x)$ is a projection of $\mathcal{D}(x)$, it also is polyhedral. Then, given that $R, P \succ 0$ and that x is a constant in the optimization, $J(x, \mathbf{u}_{(0,N-1)})$ is a strictly convex function of $\mathbf{u}_{(0,N-1)}$ and (4.54) is a strictly convex QP. Thus by known results (e.g. [7]), $\mathbf{u}_{(0,N-1)}^0(x)$ and hence $\kappa^0(x)$ are continuous, piecewise affine functions on $\mathcal{X}$, and $V^0(\cdot)$ is a strictly convex, piecewise quadratic function on $\mathcal{X}$. Lipschitz continuity follows from the fact that $\kappa^0(x)$ and $V^0(x)$ have bounded derivative in $\mathcal{X}$ and that $\mathcal{X}$ is compact. $\square$

Lemma 4.23. The optimal control problem $\mathbb{P}(x)$ is such that $V^0(0) = 0$ and $\kappa^0(0) = 0$.

Proof. If $x = 0$, then $\mathbf{u}_{(0,N-1)} = \{0, \dots, 0\}$, $\mathbf{x}_{(0,N)} = \{0, \dots, 0\}$ (along with $\mathbf{x}_{(i,1,N)} = \{\tilde{w}_i, \dots, \Phi^{N-2}\tilde{w}_i\}$, $\mathbf{u}_{(i,1,N-1)} = \{K\tilde{w}_i, \dots, K\Phi^{N-1}\tilde{w}_i\}$) are feasible solutions due to

Assumption 4.3. In this case $V(\mathbf{d}) = 0$, and because $V^0(x) \geq 0$, $V^0(x) = 0$. Also, because the optimal control problem is strictly convex on $\mathbf{u}_{(0,N-1)}$, the optimal solution for the 0^{th} partial control sequence is given by $\mathbf{u}_{(0,N-1)}^0 = \{0, \dots, 0\}$ and then $\kappa^0(0) = 0$. $\square$

Theorem 4.24. *System (4.1) under the control law of (4.44) is ISS with domain of attraction $\mathcal{X}$.*

Proof. Let system (4.1) be expressed as $f(x, w) = Ax + B\kappa^0(x) + w$. From Lemma 4.23, $f(0, 0) = 0$ and from Proposition 4.22, $f(\cdot, \cdot)$ is continuous. From Assumption 4.4, $\mathcal{X}$ contains the origin in its interior and from Theorem 4.8 it is robust positively invariant for system $f(x, w) = Ax + B\kappa^0(x) + w$ under constraints (4.3). Clearly $V^0(x) \geq 0$ and from proposition 4.22, $V^0(x)$ is Lipschitz continuous on $\mathcal{X}$. By construction $V^0(x) \geq x^T Qx$ and then $V^0(x) \geq \lambda_{\min}(Q)x^T x$, so taking $\alpha_1(s) = \lambda_{\min}(Q)s^2$ gives the LHS inequality of (2.27a). From Lemma 4.23, $V^0(0) = 0$ and $V^0(x) \geq x^T Qx > 0$, for all $x \neq 0, x \in \mathcal{X}$, but $x^T Qx$ is a continuous function in $\mathcal{X}$, which is a polytopic set that contains the origin in its interior. Thus there exist a scalar c such that for all $x \in \mathcal{X}$, $V^0(x) \leq cx^T Qx$, and then $\alpha_2(s) = c\lambda_{\max}(Q)s^2$ satisfies the RHS inequality of (2.27a). From (4.51) and by taking the candidate solutions as defined in (4.45) and (4.46), it follows that $V(\tilde{\mathbf{d}}(f(x, 0))) - V^0(x) \leq -(x^T Qx + \kappa^0(x)^T R\kappa^0(x))$, which implies $V^0(f(x, 0)) - V^0(x) \leq -(x^T Qx + \kappa^0(x)^T R\kappa^0(x))$, and then (2.27b) follows directly with $\alpha_3(s) = \lambda_{\min}(Q)s^2$. Finally (i) follows directly from Lipschitz continuity of $\kappa^0(x)$ stated in Proposition 4.22. $\square$

Remark 4.25. *If $\mathbb{X}_0$ satisfies*

$$\mathbb{X}_0 \oplus \Omega_\infty(\mathbb{W}) \subseteq \bar{\Omega}, \quad (4.55)$$

then the solution given by generating the 0^{th} and 1^{st} partial state and control sequences with $u = Kx$ (i.e. $\mathbf{x}_{(0,N)} = \{x, \dots, \Phi^N x\}$, $\mathbf{u}_{(0,N-1)} = \{Kx, \dots, K\Phi^N x\}$, $\mathbf{x}_{(i,1,N)} =$

$\{\tilde{w}_i, \dots, \Phi^{N-1}\tilde{w}_i\}$, $\mathbf{u}_{(i,1,N-1)} = \{K\tilde{w}_i, \dots, K\Phi^{N-2}\tilde{w}_i\}$ is feasible for $x \in \bar{\Omega}$ since

$$(I, K)\mathbb{X}_0 \oplus \bigoplus_{j=1}^{N-1} Y_{(1,j)} \oplus \bigoplus_{j=N}^{N+N_2-1} (I, K)\Phi^{j-N}X_{(1,N)} \oplus (\Omega_\infty(\mathbb{X}_1), K\Omega_\infty(\mathbb{X}_1)) = (I, K)\mathbb{X}_0 \oplus \bigoplus_{j=1}^{N+N_2-1} (I, K)\Phi^{j-1}\mathbb{W} \oplus (\Omega_\infty(\mathbb{X}_1), K\Omega_\infty(\mathbb{X}_1)) = (I, K)\mathbb{X}_0 \oplus (I, K)\Omega_\infty(\mathbb{W}) \subset (I, K)\bar{\Omega}.$$
 Also, if K is the LQR optimal gain associated with the cost of (4.50), then it will be the optimal solution of $\mathbb{P}(x)$. Therefore, if during the closed-loop execution the state enters into $\bar{\Omega}$, the control law $\kappa^0(x)$ will revert to $u = Kx$. Based on the analysis of Remark 4.6, there will always exist N, N_2 large enough so that $\mathbb{X}_0$ is small enough such that (4.55) is met.

4.4.3 On the selection of $\mathbb{X}_0$

In Assumption 4.2 it is required that $\mathbb{X}_0$ is an invariant set for $z^+ = \Phi z + w$, $w \in \Phi\mathbb{X}_1$, in order to guarantee recursive feasibility, whereas in Theorem 4.14 and Theorem 4.19 it is required that $\mathbb{X}_0$ is an invariant set for $z^+ = \Phi z + w$, $w \in \mathcal{W}$, where $\mathcal{W}$ is any outer approximation of $\text{conv}(\Phi^{N_3}\mathbb{W} \cup \Phi^{N_3+1}\mathbb{W} \cup \Phi^{N_3+2}\mathbb{W} \dots)$ in order to guarantee nestedness of the domains of attraction.

Let $\mathbb{X}_0$ be given by

$$\mathbb{X}_0 = \alpha\bar{\Omega}_0 \oplus \tilde{\Omega}_\infty(W), \quad (4.56)$$

where $0 < \alpha \leq 1$ is a scalar, $\bar{\Omega}_0$ is the maximal invariant set for $x^+ = \Phi x$ s.t. (4.3) under $u = Kx$, and $\tilde{\Omega}_\infty(W)$ is an invariant approximation of the minimal robust invariant set [72] for $z^+ = \Phi z + w$, $w \in W$, where W is either $\Phi\mathbb{X}_1$ or $\mathcal{W}$ (so that any of the two previous conditions on $\mathbb{X}_0$ can be satisfied). $\mathbb{X}_0$ is clearly robustly invariant for $z^+ = \Phi z + w$, $w \in W$. The motivation for this construction is as follows. In the absence of disturbances $\mathbb{X}_0$ would have to be chosen to be $\bar{\Omega}_0$ so as to relax (4.23a). However, since (4.23a) is applied at $k = N + N_2$, $\mathbb{X}_0$ might be too large thereby causing (4.28) to be too tight. To avoid this the scaling factor α is introduced, whose size should be commensurate with the contraction provided by Φ^{N_2} . In the presence of disturbances however, recursive feasibility, as introduced

in Remark 4.4, requires that $\mathbb{X}_0$ accounts for these disturbances which justifies the inclusion of the second term in the RHS of (4.56). This term has to be made to be as small as possible in order to relax (4.28), hence it is taken to be the minimal robust invariant set. Since the minimal robust invariant set however may not have a finite description it has been replaced in (4.56) by a robustly invariant approximation $\tilde{\Omega}_\infty(\Phi\mathbb{X}_1)$. Similarly, to reduce computational complexity $\bar{\Omega}_0$ can be replaced by an invariant approximation.

This construction of $\mathbb{X}_0$ has provided good results in terms of the resulting domains of attraction as will be shown in Section 4.5.

4.5 Simulation results

The benefits of SPTMPC are illustrated with statistics for 50 randomly generated systems controlled by the proposed SPTMPC strategies and PTMPC, using different prediction horizons N . The SPTMPC strategy of Section 4.4.1 will be referred to as SPTMPC1, and the SPTMPC strategy of Section 4.4.2 will be referred to as SPTMPC2.

The same systems as in Chapter 3 are considered, which have been generated with $n = 2$, $n_u = 1$, A, B being random matrices where each component is selected from an independent uniform distribution between -1.25 and 1.25, but only systems where the spectral radius of A is greater than 0.95 are used. F, G are chosen to represent the constraints $-\underline{1} \leq \tilde{F}x_t \leq \underline{1}$, $\tilde{F} \in \mathbb{R}^{2 \times 2}$, and $-1 \leq u_t \leq 1$, so that $F = [\tilde{F}^T, -\tilde{F}^T, [0, 0]^T, [0, 0]^T]^T$ and $G = [0, 0, 0, 0, -1, 1]^T$, with $\tilde{F}$ being a random matrix where each component is selected from an independent uniform distribution between -0.1 and 0.1. K is chosen to be the LQ optimal gain for the cost matrices $Q = I_{2 \times 2}$ and $R = 0.1$. $\mathbb{W} = \text{conv}(\{[0.1, 0.1]^T, [0.1, -0.1]^T, [-0.1, 0.1]^T, [-0.1, -0.1]^T\})$.

The sets in the cost of PTMPC are given by $\mathcal{Q} = \{x : (1, 0)x \leq 1, (0, 1)x \leq 1, (-1, 0)x \leq 1, (0, -1)x \leq 1\}$, $\mathcal{R} = \{u : u \leq 1, -u \leq 1\}$ and $\mathcal{P} = [\mathcal{P}_1^T \quad -\mathcal{P}_1^T]$,

where $\mathcal{P}_1^T \in \mathbb{R}^{s \times n}$, $s = 2$ is computed as in [55]. $\bar{\Omega}_0$ is chosen to be the maximal invariant set for the nominal system under $u = Kx$ and subject to (4.3) and $\mathbb{X}_f$ is chosen to be the maximal invariant set for system (4.1) under $u = Kx$ and subject to (4.3). N_2 is chosen to be the maximum between 5 and the minimum value such that $(I, K)\mathbb{X}_0 \oplus \bigoplus_{j=1}^{N_2} (I, K)\Phi^{j-1}\mathbb{W} \oplus (I, K)\Omega_\infty(\mathbb{X}_1) \subseteq (I, K)\Omega \subseteq \mathbb{Y}$ is feasible. The latter condition, on account of Remark 4.6, guarantees that the domain of attraction is not-empty. However, experimentation suggest that it is convenient to take N_2 to be at least 5 so as to relax the terminal constraints for the 0^{th} predicted state. W is taken to be $\mathcal{W}$, N_3 is chosen to be equal to $N_2 + 1$ and α is chosen to be maximum of the components of $\Phi^{N_2}\underline{1}$. The simulations have been performed on Matlab 7.11.0 (R2010b) with the Gurobi solver on a machine with Windows 7 and an Intel Core 2 Duo E8400 processor at 3GHz.

The domains of attraction using SPTMP1, SPTMP2 and PTMPC are computed for each of these systems. Table 4.1 shows the mean values (over all the random systems) of areas of the domains of attraction obtained with different prediction horizons. In order to make the values comparable, these are reported relative to the area of $\mathbb{X}_f$, such that $A_N = \text{Area}(\mathcal{X}_{N, N_2}) / \text{Area}(\mathbb{X}_f)$ (since N_2 is fixed for each system in these simulations, it is omitted from the notation for A_N). The table also reports information relevant to the computational implementation, such as the mean values for the number of variables, equalities and inequalities in the online optimization and average computational times. More precisely, the reported values correspond to the average time it takes the solver to perform the online optimization for a given prediction horizon, where the average is computed over a set of initial conditions evenly distributed around the edge of the domain of attraction of OPTMPC and over all random systems.

For the same N SPTMPC1,2 yield larger domains of attraction in average than PTMPC, when $N \geq 4$. This is mainly due to the fact that disturbance compensation

Table 4.1: Domain of Attraction and computational aspects.

N		1	2	3	4	5	6	7	8	9	10
A_N	SPTMPC1	1.482	2.937	4.468	6.541	7.845	9.054	10.33	11.91	13.73	16.07
	SPTMPC2	1.482	2.937	4.59	7.440	8.723	10.42	12.52	14.81	17.47	20.63
	PTMPC	1.482	3.139	4.718	5.710	6.558	7.458	8.530	9.742	11.203	12.953
vars	SPTMPC1	70.6	96.9	123.3	149.6	175.9	202.2	228.5	254.9	281.2	307.5
	SPTMPC2	43	64	85	106	127	148	169	190	211	232
	PTMPC	24.3	71.6	145.0	244.3	369.6	520.9	698.3	901.6	1130	1386
ineqs	SPTMPC1	305.1	357.7	410.3	462.9	515.5	568.2	602.8	673.4	726.0	778.6
	SPTMPC2	268.8	316.1	363.4	410.7	458.0	505.3	552.6	599.9	647.2	694.5
	PTMPC	63.6	161.4	307.3	501.1	743.0	1033	1371	1757	2190	2672
eqs	SPTMPC1	12	22	32	42	52	62	72	82	92	102
	SPTMPC2	12	22	32	42	52	62	72	82	92	102
	PTMPC	10	28	54	88	130	180	238	304	378	460
time [ms]	SPTMPC1	2.5	5.9	26.5	30.2	40.1	45.0	49.5	50.0	53.5	54.2
	SPTMPC2	2.6	5.1	19.1	19.5	20.5	24.2	26.2	27.3	29.6	30.3
	PTMPC	1.70	2.60	3.94	6.55	12.17	20.46	36.28	56.69	100.0	147.4

is not restricted to the horizon N , but instead enters into Mode 2, thereby relaxing the terminal constraints.

PTMPC has a "full" triangular structure, with more parameters than the striped structure. Therefore, in general, it has the potential outperform SPTMPC. However, the price of the "full" triangular structure is that it implies a number of variables and constraints in the online optimization that grows quadratically with N , whereas in the online optimization of SPTMPC that number only grows only linearly with N . Thus with SPTMPC one can use longer horizons, thereby enlarging the size of the domain of attraction, at a computational cost which may still be less than that required by PTMPC. For example, while for $N \leq 4$ PTMPC leads to larger domains of attraction and faster online computations, SPTMPC1 for $N \geq 8$ and SPTMPC2 for $N \geq 7$ lead to larger domains of attraction while requiring less demanding online computations. Thus, due to the quadratic increase (in number of variables and constraints) in PTMPC versus the linear increase in SPTMPC1,2, larger values of N can be used in SPTMPC1,2, thus obtaining even larger domains of attractions and still using fewer variables and constraints (as is the case for instance for SPTMPC1 with $N = 10$ versus PTMPC with $N = 8$, or the case of SPMPTC2 with $N = 10$ versus

PTMPC with $N = 7$).

The strategy of Chapter 3, OPTMPC, is faster than SPTMPC, and in the average case larger horizons could be used for OPTMPC which could give better domains of attraction while requiring the solution of less demanding optimizations problems. However, the offline designs in OPTMPC might impose too much conservativeness for some systems. For instance, the example introduced in Section 3.4 defined by

$$A = \begin{bmatrix} 1.19 & 0.95 \\ -0.89 & 0.96 \end{bmatrix}, B = \begin{bmatrix} 0.66 \\ 0.74 \end{bmatrix}, F = \begin{bmatrix} 0.083 & 0.114 & -0.083 & -0.114 & 0 & 0 \\ -0.035 & -0.070 & 0.035 & 0.070 & 0 & 0 \end{bmatrix}^T,$$

$G = [0 \ 0 \ 0 \ 0 \ 1.44 \ -1.44]^T$ and $\mathbb{W} = \text{conv}(\{[0.1, 0.1]^T, [0.1, -0.1]^T, [-0.1, 0.1]^T, [-0.1, -0.1]^T\})$, shows this situation. The results reported in Table 3.2 are complemented by the areas of the domains of attraction obtained with SPTMPC1 and SPTMPC2.

Table 4.2: Domain of Attraction and Computational Aspects

N		1	2	3	4	5	6	7	8	9	10
A_N	OPTMPC	1.302	1.446	1.626	1.830	1.880	1.935	1.958	1.977	1.990	2.002
	SPTMPC1	1.302	1.446	1.496	2.617	2.860	2.954	3.005	3.042	3.067	3.089
	SPTMPC2	1.302	1.446	1.496	2.747	3.071	3.194	3.236	3.253	3.294	3.294
	PTMPC	1.302	1.747	2.174	2.508	2.707	2.850	2.975	3.073	3.137	3.178

In this case, increasing the value of N will not bring considerable benefits for OPTMPC, so that increasing N for OPTMPC will not allow one to obtain domains of attraction as good as one could obtain with PTMPC or SPTMPC.

Finally, note that SPTMPC2 does not constrain the predicted state to lie inside $\mathbb{X}_f$ at any instant and hence it achieves larger domains of attraction with fewer variables than SPTMPC1. However, SPTMPC1 enjoys stronger stability properties which guarantee convergence to $\mathbb{X}_f$. While, on average, the difference seems to be considerable for large values of N , it can be very small in some cases, as shown in Table 4.2.

4.6 Conclusions

The proposed RMPC strategy has a number of variables and constraints that grows only linearly with the prediction horizon. This is achieved by considering a predicted control policy with a striped structure. The predicted control policy deployed allows the disturbance compensation terms in the inputs to extend over the entire prediction horizon. This provides extra relaxation to the terminal constraints and, on account of this, the proposed SPTMPC strategy can lead to domains of attraction which are larger than those possible through the use of PTMPC with the same prediction horizon. However, given the full structure of PTMPC, the converse is also possible and PTMPC can outperform SPTMPC.

In any case, since the growth of the size of the problem in SPTMPC is only linear, compared with the quadratic growth in PTMPC, one can increase N in SPTMPC in order to obtain larger domains of attraction than PTMPC.

Unlike the case of OPTMPC, the online optimization problem in SPTMPC does consider the 1^{st} partial extreme sequences as variables in the online optimization, and as such q has an influence in the number of variables and constraints. More precisely the dimensions of the problem of SPTMPC and PTMPC scale with the same order as a function of n , since both have terms qn . However, the rate with which they do it will depend on mainly on N . As shown by simulation results, for small N SPTMPC has a larger problem size, and will grow faster than PTMPC as n increases. On the other hand, for larger N , the N^2 terms in PTMPC start to dominate, and then the problem sizes of PTMPC will grow faster than in SPTMPC as n increases.

The results show that, in general, OPTMPC can be faster than SPTMPC, so then larger prediction horizons can be used for OPTMPC which may allow one to obtain larger domains of attraction than what could be obtained with SPTMPC. However, the conservativeness imposed by the offline design in OPTMPC means that

the domains of attraction may be too small, in which case SPTMPC may be more convenient.

The presented examples show that there is no a priori way to indicate preference for OPTMPC, SPTMPC or PTMPC in terms of both domain of attraction and speed of the online optimization: none of these methods will consistently provide larger domains of attraction with faster online optimizations.

Two variations of the strategy have been proposed. SPTMPC1 is guaranteed to converge to the minimal invariant set under $u = Kx$, but does so at the expense of additional constraints and variables, thus reducing the domain of attraction and making the online optimization slower. On the other hand, SPTMPC2 does not use the terminal constraint of (4.34), has a different cost and has fewer variables. Then it allows one to obtain larger domains of attraction and faster online optimizations, but it comes at the cost of a weaker notion of stability, namely ISS.

The use of (4.34) is rather artificial, because it considers a terminal set which is invariant for a fixed control law that is not the terminal control law of the predicted policy, which is not fixed. On account of this, convergence to the control law associated to that terminal set is proven. The possibility of obtaining strong notions of stability (i.e. convergence to a set) without imposing that extra constraint, or in a way that is more directly related to the structure of the control policy, is yet to be explored and forms the subject of future research.

Finally, the fact that the terminal control law is not fixed and that it can provide extra relaxation for the constraints throughout the infinite horizon, suggests that the strategy could be applied to systems where the state feedback $u = Kx$ is not necessarily feasible. This, also, forms a subject for future research.

Chapter 5

Dual-Mode Robust MPC with optimized polytopic dynamics

Polytopic dynamics have been introduced into MPC predictions and have been optimized so as to maximise invariant ellipsoids for systems subject to multiplicative uncertainty [19]. This work was recently extended to the case of mixed additive and multiplicative uncertainty and provides a sufficient condition for the invariance of the ellipsoids [34]. Additionally, under the assumption that the uncertainty is known during a prediction horizon N , this extension was used as the terminal control law of an overall robust MPC strategy that deployed an affine-in-the-disturbances control policy in Mode 1. This assumption (of known uncertainty during the Mode 1 prediction horizon) is needed to enable the handling of constraints in the overall MPC strategy of [34]. The aim of this chapter is to remove this assumption by deploying polytopic tubes to describe the evolution of the predicted trajectories and to establish the necessity and sufficiency of the relevant invariance conditions of the polytopic dynamics.

The benefits of the strategy are illustrated by numerical simulations.

5.1 Introduction

Robustness to uncertainty, in the form of additive disturbances or multiplicative uncertainty due to imprecise knowledge of the system dynamics, is an important topic in RMPC. Early contributions considering multiplicative uncertainty include [80] (which considers gain mismatches) and [2] (which assumes a weighting sequence system description). The case of additive bounded disturbances was dealt through the use of constraint tightening [36], affine-in-the-disturbances policies [58], [38] and tube MPC [63], [75]. Recently [19] considered robustness to multiplicative polytopic uncertainty and introduced polytopic prediction dynamics. The precise evolution of the predicted dynamics was unknown but was generated by the interpolation of vertex dynamics, each vertex corresponding to a vertex of the polytopic uncertainty. The introduced dynamics were chosen so as to maximize the volume of an ellipsoid which is robustly invariant under a given control law. Such ellipsoids can be used as terminal sets and the maximization of their volume relaxes the relevant terminal constraint and hence results in larger domains of attraction for the overall MPC strategy. However, maximal invariant sets of linear systems subject to linear state and input constraints can only be approximated conservatively by ellipsoids, whereas they can be described with an arbitrary degree of accuracy by polytopes. The difficulty in ensuring constraint satisfaction is that, due to the uncertainty, the predicted states and control inputs cannot be known exactly and defining exact regions for them leads to an exponential growth in computation. To avoid this [32] proposed the use of a variation of Farkas' Lemma [31], [39] in the construction of tubes whose cross-sections were so defined as to contain the predicted states and control moves. The approach used the dual-mode prediction paradigm where the predicted control moves in Mode 1 are given by optimizing the degrees of freedom of a predicted control policy given by the closed-loop paradigm, whereas the control moves in Mode 2 are given by a

predetermined terminal control law.

The aim here is to extend [19] to the case of both polytopic multiplicative and additive uncertainty. Such an extension has already been undertaken in [34], under the assumption that the disturbance lies in a unitary infinity-norm ball, which uses the approach of [19] in conjunction with a disturbance feedback term in the control policy. As in [19], the MPC strategy can be based directly on the control law resulting from the optimized polytopic dynamics. Alternatively, [34] (Section IV) deploys this strategy as a terminal control law in Mode 2 of an MPC strategy based on the dual-mode prediction paradigm, with Mode 1 deploying an affine-in-the-disturbances control policy [38]. This setting allows for multiplicative uncertainty in Mode 2, however to handle constraints in Mode 1, [34] assumes that the multiplicative “uncertainty” is known over Mode 1. Furthermore, on account of the use of the $\mathcal{S}$ -procedure and an outer bounding of the disturbance set, the results (in terms of the definition of the optimized polytopic dynamics and the definition of an upper bound on a minimax cost) are claimed to be sufficient only.

The contributions of the strategy proposed in this chapter are that: (i) it shows that the invariance conditions of [34], for the case when the disturbance lies in a unitary infinity-norm ball, are both sufficient and necessary; (ii) it proposes necessary and sufficient conditions for invariance for the case of a general polytopic disturbance set; (iii) it develops an overall prediction dual-mode prediction based MPC strategy for which the multiplicative uncertainty is not required to be known over the prediction horizon.

The chapter is structured as follows. Section 5.2 gives the system description and provides a background on the optimized polytopic dynamics for system with multiplicative uncertainty [19] and systems with additive and multiplicative uncertainty [34]. The necessity and sufficiency of the conditions for invariance of the optimized polytopic dynamics are provided in Section 5.3. Section 5.4 presents the overall MPC

strategy, with focus on constraint satisfaction, the cost function and the online implementation. Finally, Sections 5.5 and 5.6 present simulation results and conclusions.

5.2 System description and background on optimized polytopic dynamics

Consider a linear, discrete-time system subject to multiplicative and additive uncertainty

$$x_{t+1} = A_t x_t + B_t u_t + D_t w_t, \quad (5.1)$$

where $x_t \in \mathbb{R}^n$, $w_t \subseteq \mathbb{R}^{n_w}$, and $u_t \in \mathbb{R}^{n_u}$ are the state, disturbance and control input at time t . The system matrices $A_t \in \mathbb{R}^{n \times n}$, $B_t \in \mathbb{R}^{n \times n_u}$, $D_t \in \mathbb{R}^{n \times n_w}$, and the disturbances w_t are unknown for all t , but are known to be contained in polytopic sets such that

$$(A_t, B_t, D_t) \in \Xi = \text{conv}\{(A_i, B_i, D_i) : i = 1, \dots, \nu_m\}, \quad (5.2a)$$

$$w_t \in \mathbb{W} = \text{conv}\{\tilde{w}_j : j = 1, \dots, \nu_a\}, \quad (5.2b)$$

where the ν_m vertices of the multiplicative uncertainty, (A_i, B_i, D_i) , and the ν_a vertices of the additive uncertainty vertices, $\tilde{w}_j$, are known. Throughout, the indices i and j will identify the ν_m and ν_a vertices of the multiplicative and additive uncertainty, respectively. The description of (5.2) implies that the uncertain matrices and the disturbances satisfy

$$\begin{aligned} (A_t, B_t, D_t) &= \sum_{i=1}^{\nu_m} \lambda_i (A_i, B_i, D_i), \\ w_t &= \sum_{j=1}^{\nu_a} \mu_j \tilde{w}_j, \end{aligned} \quad (5.3)$$

where the interpolation parameters λ_i, μ_j are unknown scalar interpolation parameters that satisfy $\lambda_i, \mu_j \geq 0$, $\sum_{i=1}^{\nu_m} \lambda_i = 1$, $\sum_{j=1}^{\nu_a} \mu_j = 1$.

System (5.1) is subject to mixed state and input constraints

$$(x_t, u_t) \in \mathbb{Y}, \quad (5.4)$$

where $\mathbb{Y}$ is a polytopic set that contains the origin in its interior and is given by

$$\mathbb{Y} = \{(x, u) : Fx + Gu \leq \underline{1}\}, \quad (5.5)$$

where $F \in \mathbb{R}^{n_c \times n}$ and $G \in \mathbb{R}^{n_c \times n_u}$.

Recall from Section 2.2.3.3 the case where $w = 0$, for which [19] maximizes the volume of ellipsoidal invariant sets through the use of a polytopic control policy that uses a controller state $v \in \mathbb{R}^{n_v}$, such that the predicted inputs and the dynamics of v are given by

$$u_{t+k|t} = Kx_{t+k|t} + Lv_{t+k|t}, \quad k = 0, 1, 2, \dots \quad (5.6a)$$

$$v_{t+k+1|t} = \sum_{i=1}^{\nu_m} \lambda_i M_i v_{t+k|t}, \quad k = 0, 1, 2, \dots \quad (5.6b)$$

where the λ_i are the same as in (5.3). It is noted knowledge of the λ_i is not needed. As will be seen later, only the knowledge of the evolution associated with the vertices is needed to guarantee feasibility of the predictions (via invariant sets). Also, only the first predicted input is applied to the system, and this does not depend on the interpolation parameters.

By a process of lifting in which the system and controller states form an augmented state $z \in \mathbb{R}^{n_z}$, $n_z = n + n_v$, the prediction dynamics can be written as follows

$$z_{t|t} = \begin{bmatrix} x_t \\ v_{t|t} \end{bmatrix}, \quad (5.7)$$

$$z_{t+k+1|t} = \Psi_{t+k} z_{t+k|t}, \quad k = 0, 1, 2, \dots,$$

where $\Psi_{t+k} \in \text{conv}\{\Psi_i : i = 1, \dots, \nu_m\}$, and

$$\Psi_i = \begin{bmatrix} \Phi_i & BL \\ 0 & M_i \end{bmatrix}, \quad \text{and} \quad \Phi_i = A_i + B_i K, \quad i = 1, \dots, \nu_m. \quad (5.8)$$

According to (5.7) the predicted inputs of (5.6a) are given by

$$u_{t+k|t} = [K \ L]z_{t+k|t}, \quad k = 0, 1, 2, \dots \quad (5.9)$$

K is assumed to be optimal in some sense (e.g. LQR for the nominal system), and then L and M_i are designed to optimize the polytopic dynamics of (5.7).

Robust invariance of an ellipsoid $E_z = \{z : z^T P z \leq 1\}$ where $P \in \mathbb{R}^{n_z \times n_z}$, for the dynamics of (5.7) is equivalent to

$$P - \Psi_i^T P \Psi_i \succeq 0, \quad i = 1, \dots, \nu_m, \quad (5.10)$$

while feasibility of the states inside the ellipsoid for the constraints of (5.4) is equivalent to

$$\begin{bmatrix} Z & \begin{bmatrix} F + GK & GL \end{bmatrix} \\ * & P \end{bmatrix} \succeq 0, \quad Z_{p,p} \leq 1, \quad p = 1, \dots, n_c, \quad (5.11)$$

where the symmetric matrix $Z \in \mathbb{R}^{n_c \times n_c}$ is a free variable with elements $Z_{p,q}$, $p, q = 1, \dots, n_c$, and $*$ denotes symmetric blocks.

To maximize the volume of the projection, E_x , of E_z onto the x -subspace one needs to maximize $\log \det (TP^{-1}T^T)$ over P , L , M_i subject to (5.10) and (5.11), where T is the transformation for which $x = Tz$ for all z . This optimization is non-convex because (5.10) is a BMI in P, M_i, L , however, as shown in [19] the problem can be convexified through the use of a transformation proposed in [83].

It is shown in [19] that the maximum volume E_x can be achieved by taking any

value of n_v such that $n_v \geq n$, and it has the same volume as the largest invariant ellipsoid that can be obtained under any linear state feedback.

The online optimization aims at minimizing the worst case cost (min-max problem) of $J = \sum_{i=0}^{\infty} x_{t+k|t}^T Q x_{t+k|t} + u_{t+k|t}^T R u_{t+k|t}$. Since the exact evaluation of that cost would require one to evaluate a number of trajectories that grows exponentially with N (all the trajectories associated with the vertices of Ξ), this is actually performed by minimizing over $v_{t|t}$ a quadratic upper bound of the minimax cost, given by $J(x_t, v_{t|t}) = x_t^T W_x x_t + v_{t|t}^T W_v v_{t|t}$, subject to an LMI that is equivalent to $z_{t|t} \in E_z$. Then only the first predicted input, $u_{t|t} = Kx_t + Lv_{t|t}$ is actually applied to the plant according to the receding horizon strategy. The use of $z_{t|t} \in E_z$ guarantees recursive feasibility of the optimization problem and satisfaction of the system constraints, while the cost possesses a decrease property that allows one to prove asymptotic stability of the controlled system.

The extension of this approach to the case of both additive and multiplicative uncertainty (i.e. when $w_t \neq 0$) is straightforward. Thus [34], for the special case of $\mathbb{W} = \{w : \|w\|_{\infty} \leq 1\}$, employs (5.6) endowed with disturbance feedback so that the predicted inputs and dynamics of the controller state are given by

$$u_{t+k|t} = Kx_{t+k|t} + Lv_{t+k|t}, \quad k = 0, 1, 2, \dots \quad (5.12a)$$

$$v_{t+k+1|t} = \sum_{i=1}^{\nu_m} \lambda_i (M_i v_{t+k|t} + S_i w_{t+k}), \quad k = 0, 1, 2, \dots \quad (5.12b)$$

according to which the dynamics of the augmented state z become

$$z_{t|t} = \begin{bmatrix} x_t \\ v_{t|t} \end{bmatrix} \quad (5.13)$$

$$z_{t+k+1|t} = \Psi_{t+k} z_{t+k|t} + \tilde{D}_{t+k} w_{t+k}, \quad k = 0, 1, 2, \dots$$

where $(\Psi_{t+k}, \tilde{D}_{t+k}) = \text{conv}\{(\Psi_i, \tilde{D}_i) : i = 1, \dots, \nu_m\}$ with $\tilde{D}_i = [D_i^T S_i^T]^T$, and

$w_{t+k} \in \mathbb{W}$. From (5.13) the predicted inputs of (5.12a) are also given by (5.9).

In this case L , M_i and S_i are designed to optimize the polytopic dynamics of (5.7).

Additive uncertainty precludes the use of invariance conditions (5.10). Instead, in order to find conditions for invariance and feasibility of the ellipsoid $E_z = \{z : z^T P^{-1} z \leq 1\}$, [34] defines the set $\bar{\mathbb{W}} = \{w : w^T R_i w \leq \text{tr}(R_i), i = 1, \dots, \nu_m\}$ that overbounds the admissible disturbance set of $\mathbb{W}$, where R_i is a diagonal matrix with non-negative elements, and uses the $\mathcal{S}$ -procedure to ensure that the conditions $z_{t+k|t}^T P^{-1} z_{t+k|t} \leq 1$ and $w_{t+k}^T R_i w_{t+k} \leq \text{tr}(R_i)$ imply that $z_{t+k+1|t}^T P^{-1} z_{t+k+1|t} \leq 1$. This leads to the sufficient condition

$$\begin{bmatrix} P & \Psi_i P & \tilde{D}_i \\ * & \alpha_i P & 0 \\ * & * & R_i \end{bmatrix} \succeq 0, \quad \text{tr}\{R_i\} \leq 1 - \alpha_i, \quad i = 1, \dots, \nu_m, \quad (5.14)$$

while feasibility of the states inside the ellipsoid $E_z = \{z : z^T P^{-1} z \leq 1\}$ is guaranteed by

$$\begin{bmatrix} Z & \begin{bmatrix} F + GK & GL \end{bmatrix} P \\ * & P \end{bmatrix} \succeq 0, \quad Z_{p,p} \leq 1, \quad p = 1, \dots, n_c \quad (5.15)$$

where the symmetric matrix $Z \in \mathbb{R}^{n_c \times n_c}$ is a free variable with elements $Z_{p,q}$, $p, q = 1, \dots, n_c$, and $*$ denotes symmetric blocks.

From here on, the procedure of [34] parallels that of [19] in that the same convexification technique of [83] is applied to maximize the volume of the invariant ellipsoid E_x under the dynamics of (5.12) and (5.13) by maximizing $\log \det (TPT^T)$ over P, M_i, L subject to (5.14) and (5.15). The result of the convexification procedure, however, is still bilinear due to the presence of α_i . It is then proposed to replace the α_i by a single α which, at the cost of some conservativeness, provides BMIs that are much simpler than those of (5.14) and (5.15) and yields a bilinear problem that can be solved either with available BMI solvers, or an ad-hoc iterative approach that uses

the special structure of the problem and solves a sequence of LMIs.

The online problem must consider a cost function different to that of [19] because the additive uncertainty makes it infinite. Two alternatives are then considered: an $\mathcal{H}_\infty$ type of predicted cost or a nominal cost. For relaxing the constraints and avoiding the use of LMIs in the online optimization, a polytope $\mathcal{P}_z$ that is robustly invariant for the dynamics of (5.13) and constraints (5.4) is considered instead of the ellipsoid E_z . A first alternative for the online algorithm is similar to that in [19], where the predicted cost function is minimized over $v_{t|t}$ subject to $z_{t|t} \in \mathcal{P}_z$. A second alternative, which can be used when the future uncertainty during Mode 1 is known, consists of using the dual-mode prediction paradigm such that the predicted control law of (5.12) is used as a terminal control law, while an affine-in-the-disturbances policy is used in Mode 1. In both alternatives, the type of stability obtained depends on the cost used. When using an $\mathcal{H}_\infty$ type of predicted cost, the stability result is presented in the form of an ℓ_2 -bound of the closed-loop performance of the controlled system. When using a nominal cost, on the other hand, the stability will be of the ISS type.

5.3 Necessity and sufficiency conditions for the optimized polytopic dynamics

On account of the use of the $\mathcal{S}$ -procedure and the bounding of $\mathbb{W}$ via $w^T R_i w \leq \text{tr}(R_i)$, $i = 1, \dots, \nu_m$, condition (5.14) is stated as being sufficient only. In addition, the bilinear dependence of (5.14) on P and α_i , leads to intractable computation and [34], at some extra conservatism, replaces the α_i by a single α which can be optimized by a simpler bilinear problem. This section introduces a necessary and sufficient condition for the robust invariance of an ellipsoidal set. This condition has a single α and is valid for the more general description of the disturbances of (5.2b). It is then shown that for the special case of $\mathbb{W} = \{w : \|w\|_\infty \leq 1\}$ of [34] the conditions of [34] are

both necessary and sufficient.

Theorem 5.1. *The ellipsoid $E_z = \{z : z^T P^{-1} z \leq 1\}$, with $P \succ 0$, is robustly positively invariant for the dynamics of (5.12) and (5.13) if, and only if there exists a non-negative scalar α such that:*

$$\begin{bmatrix} P & \Psi_i P & \tilde{D}_i \tilde{w}_j \\ * & \alpha P & 0 \\ * & * & 1 - \alpha \end{bmatrix} \succeq 0, \quad \begin{array}{l} i = 1, \dots, \nu_m, \\ j = 1, \dots, \nu_a. \end{array} \quad (5.16)$$

Proof. It is first proved that, according to the $\mathcal{S}$ -procedure, E_z is invariant for the dynamics of (5.12) and (5.13) if and only if there exists $\alpha > 0$ such that

$$1 - [z^T \Psi_i^T + \tilde{w}_j^T \tilde{D}_i^T] P^{-1} [\Psi_i z + \tilde{D}_i \tilde{w}_j] \geq \alpha (1 - z^T P^{-1} z), \quad i = 1, \dots, \nu_m, \quad j = 1, \dots, \nu_a. \quad (5.17)$$

This condition implies invariance since if $1 - z^T P^{-1} z \geq 0$, i.e. $z \in E_z$, then $1 - [z^T \Psi_i^T + \tilde{w}_j^T \tilde{D}_i^T] P^{-1} [\Psi_i z + \tilde{D}_i \tilde{w}_j] \geq 0$ for $i = 1, \dots, \nu_m, j = 1, \dots, \nu_a$, which implies that $1 - [z^T \Psi^T + w^T \tilde{D}^T] P^{-1} [\Psi z + \tilde{D} w] \geq 0$, i.e. $z^+ \in E_z$, for all $(\Psi, \tilde{D}) = \text{conv}\{(\Psi_i, \tilde{D}_i) : i = 1, \dots, \nu_m\}$ and $w \in \mathbb{W}$.

Invariance of E_z implies that if $z \in E_z$, then $z^+ \in E_z$ for all possible realizations of the uncertainty, and particularly for all the extreme realizations. Thus if $1 - z^T P^{-1} z \geq 0$, then $1 - [z^T \Psi_i^T + \tilde{w}_j^T \tilde{D}_i^T] P^{-1} [\Psi_i z + \tilde{D}_i \tilde{w}_j] \geq 0$, for $i = 1, \dots, \nu_m, j = 1, \dots, \nu_a$. Then, for each combination of (i, j) , $1 - z^T P^{-1} z \geq 0$ implies a single inequality, and according to the $\mathcal{S}$ -procedure applied to only two inequalities

$$1 - [z^T \Psi_i^T + \tilde{w}_j^T \tilde{D}_i^T] P^{-1} [\Psi_i z + \tilde{D}_i \tilde{w}_j] \geq \alpha_{i,j} (1 - z^T P^{-1} z) \quad (5.18)$$

is a necessary condition for each $i = 1, \dots, \nu_m, j = 1, \dots, \nu_a$. Replacing $\alpha_{i,j}$ by their minimum value, α , yields that (5.17) is a necessary condition for the invariance of

E_z .

It will be proved now that (5.17) is equivalent to (5.16), and thus (5.16) is equivalent to the invariance of E_z . Condition (5.17) can be re-written as

$$\begin{bmatrix} z \\ 1 \end{bmatrix}^T \begin{bmatrix} \alpha P^{-1} - \Psi_i^T P^{-1} \Psi_i & -\Psi_i^T P^{-1} \tilde{D}_i \tilde{w}_j \\ * & 1 - \alpha - \tilde{w}_j^T \tilde{D}_i^T P^{-1} \tilde{D}_i \tilde{w}_j \end{bmatrix} \begin{bmatrix} z \\ 1 \end{bmatrix} \geq 0, \quad \begin{matrix} i = 1, \dots, \nu_m, \\ j = 1, \dots, \nu_a, \end{matrix} \quad (5.19)$$

which implies that

$$\begin{bmatrix} \alpha P^{-1} - \Psi_i^T P^{-1} \Psi_i & -\Psi_i^T P^{-1} \tilde{D}_i \tilde{w}_j \\ * & 1 - \alpha - \tilde{w}_j^T \tilde{D}_i^T P^{-1} \tilde{D}_i \tilde{w}_j \end{bmatrix} \succeq 0, \quad \begin{matrix} i = 1, \dots, \nu_m, \\ j = 1, \dots, \nu_a. \end{matrix} \quad (5.20)$$

In turn, (5.20) can be re-written as

$$\begin{bmatrix} \alpha P^{-1} & 0 \\ 0 & 1 - \alpha \end{bmatrix} - \begin{bmatrix} \Psi_i^T \\ \tilde{w}_j^T \tilde{D}_i^T \end{bmatrix} P^{-1} \begin{bmatrix} \Psi_i & \tilde{D}_i \tilde{w}_j \end{bmatrix} \succeq 0, \quad \begin{matrix} i = 1, \dots, \nu_m, \\ j = 1, \dots, \nu_a, \end{matrix} \quad (5.21)$$

which by Schur complementation, provided that $P^{-1} \succ 0$, is equivalent to

$$\begin{bmatrix} P & \Psi_i & \tilde{D}_i \tilde{w}_j \\ * & \alpha P^{-1} & 0 \\ * & * & 1 - \alpha \end{bmatrix} \succeq 0, \quad \begin{matrix} i = 1, \dots, \nu_m, \\ j = 1, \dots, \nu_a. \end{matrix} \quad (5.22)$$

Finally, (5.16) is obtained by left and right multiplying (5.22) by a block diagonal symmetric matrix whose diagonal blocks are I , P and 1, thus proving the desired result. $\square$

Theorem 5.1 exploits the polytopic nature of w and avoids the bounding via $w^T R_i w \leq \text{tr}(R_i)$, $i = 1, \dots, \nu_m$, of [34] which introduces an extra inequality for each i , thereby leading to sufficient only conditions. Below it is proved that for

$\mathbb{W} = \{w : \|w\|_\infty \leq 1\}$ the condition of [34], (5.14), with a single α parameter, is necessary and sufficient for invariance.

Corollary 5.2. *For $\mathbb{W} = \{w : \|w\|_\infty \leq 1\}$, the ellipsoid $E_z = \{z : z^T P^{-1} z \leq 1\}$, with $P \succ 0$, is robustly positively invariant under (5.12) and (5.13) if, and only if there exist diagonal matrices R_i , $i = 1, \dots, \nu_m$ with non-negative elements and a non-negative scalar α such that*

$$\begin{bmatrix} P & \Psi_i P & \tilde{D}_i \\ * & \alpha P & 0 \\ * & * & R_i \end{bmatrix} \succeq 0, \quad \text{tr}(R_i) \leq 1 - \alpha, \quad i = 1, \dots, \nu_m \quad (5.23)$$

Proof. From the proof of Theorem 5.1 it follows that the invariance of E_z , if $\alpha P^{-1} - \Psi_i^T P^{-1} \Psi_i \succ 0$, is equivalent to

$$\begin{bmatrix} \alpha P^{-1} - \Psi_i^T P^{-1} \Psi_i & -\Psi_i^T P^{-1} \tilde{D}_i \tilde{w}_j \\ * & 1 - \alpha - \tilde{w}_j^T \tilde{D}_i^T P^{-1} \tilde{D}_i \tilde{w}_j \end{bmatrix} \succeq 0, \quad i = 1, \dots, \nu_m, \quad j = 1, \dots, \nu_a \quad (5.24)$$

which, by Schur complements, is equivalent to

$$\tilde{w}_j^T \tilde{R}_i \tilde{w}_j \leq 1 - \alpha, \quad i = 1, \dots, \nu_m, \quad j = 1, \dots, \nu_a \quad (5.25)$$

$$\tilde{R}_i = \tilde{D}^T (P^{-1} + P^{-1} \Psi_i [\alpha P^{-1} - \Psi_i^T P^{-1} \Psi_i]^{-1} \Psi_i^T P^{-1}) \tilde{D}, \quad i = 1, \dots, \nu_m \quad (5.26)$$

The treatment when $\alpha P^{-1} - \Psi_i^T P^{-1} \Psi_i \succeq 0$ is similar but deploys a Schur complementation (e.g. see [13]) that is based on the generalized Moore-Penrose inverse.

Since the elements of $\tilde{w}_j$ are ± 1 it follows that the sum of the absolute values of the elements of $\tilde{R}_i$ is less than $1 - \alpha$ and so there exist non-negative diagonal matrices

R_i such that

$$[R_i]_{p,p} \geq \sum_{q=1}^{n_w} |[\tilde{R}_i]_{p,q}|, \quad i = 1, \dots, \nu_m, \quad p = 1, \dots, n_w \quad (5.27)$$

$$\tilde{w}_j^T R_i \tilde{w}_j \leq 1 - \alpha, \quad i = 1, \dots, \nu_m, \quad j = 1, \dots, \nu_a, \quad (5.28)$$

where $[(\cdot)]_{p,q}$ denotes the element of $(\cdot)$. By (5.27) $R_i - \tilde{R}_i$ is diagonally dominant so then its eigenvalues are non-negative, thus $R_i \succeq \tilde{R}_i$. This, with $\alpha P^{-1} - \Psi_i^T P^{-1} \Psi_i \succ 0$, implies that

$$\begin{bmatrix} R_i & -\tilde{D}_i^T P^{-1} \Psi_i \\ * & \alpha P^{-1} - \Psi_i^T P^{-1} \Psi_i \end{bmatrix} \succeq 0, \quad i = 1, \dots, \nu_m, \quad j = 1, \dots, \nu_a \quad (5.29)$$

However (5.28) implies that $\text{tr}(R_i) \leq 1 - \alpha$ and (5.29) can be shown (by appropriate Schur complements) to be equivalent to the LMI in (5.23), so that the invariance of E_z implies (5.23).

The converse is also true because (5.23) implies $R_i \succeq \tilde{R}_i$ and $w \in \{w : w^T R_i w \leq 1 - \alpha, \quad i = 1, \dots, \nu_a\}$. However, because $\text{tr}(R_i) \leq 1 - \alpha$, $\{w : w^T R_i w \leq 1, \quad i = 1, \dots, \nu_m\}$ is a set that contains the polytope $\{w : \|w\|_\infty \leq 1\}$. It therefore follows that (5.25) will be satisfied, which is equivalent to (5.24) (for $\alpha P^{-1} - \Psi_i^T P^{-1} \Psi_i \succ 0$), which by Theorem 5.1 is equivalent to invariance of E_z . $\square$

P, L, M_i, S_i , are chosen to maximize $\log \det (T^T P T)$ subject to (5.15) and either of (5.16) or (5.23). When using (5.16) instead of (5.23) one does not need the degrees of freedom contained in R_i , but (5.16) implies ν_a BMIs for each i , $i = 1, \dots, \nu_m$. It is noted that (5.16) allows the consideration of the polytopic class for w of (5.2b) whereas (5.23) addresses the special case of $\mathbb{W} = \{w : \|w\|_\infty \leq 1\}$.

As in Section 5.2, this optimization is non-convex because (5.16) and (5.23) are BMIs in P, L, M_i, S_i for a fixed α . However, it was shown in [19] that the problem can

be convexified through the use of a transformation based on that proposed in [83]. The details of such transformation for (5.15) and (5.16) and the resulting optimization that is used for finding the parameters of the optimized polytopic dynamics are provided in the appendix to this chapter (Appendix A).

5.4 The online robust MPC strategy

The MPC strategy of [34] has two possible settings. The first uses the optimized dynamics of [19] extended with static disturbance feedback and optimizes online a predicted cost over $v_{t|t}$ such that $z_{t|t} \in \mathcal{P}_z$. The second uses the dual-mode prediction paradigm where the optimized dynamics are deployed in Mode 2 and an affine-in-the-disturbances control policy [38] is used in Mode 1. To handle constraints in Mode 1, [34] assumes that the multiplicative uncertainty is known as a function of time over Mode 1. The MPC scheme proposed in this chapter is also based on the dual-mode prediction paradigm but does not require knowledge of the multiplicative uncertainty over Mode 1. It uses the optimized polytopic dynamics throughout the infinite prediction horizon, and adds degrees of freedom in the form of perturbations on the control moves in Mode 1. Constraints are handled for the case when both the multiplicative and additive uncertainty are unknown through the use of a variation of Farkas' Lemma [31], [39] which enables the definition of tubes for the predictions. An $\mathcal{H}_\infty$ type of cost is also considered, and the conditions invoked for the definition of an upper bound on the predicted cost are shown to be both necessary and sufficient.

5.4.1 The control policy and constraints

The MPC strategy developed here considers a control policy with predicted inputs as in (5.12a), with $n_v = n$, but additional degrees of freedom are included in Mode 1 in the form of perturbations $c_{t+k|t}$, $k = 0, \dots, N - 1$. The predicted inputs are then

given by

$$u_{t+k|t} = Kx_{t+k|t} + Lv_{t+k|t} + c_{t+k|t}, \quad k = 0, \dots, N-1, \quad (5.30a)$$

$$u_{t+k|t} = Kx_{t+k|t} + Lv_{t+k|t}, \quad k \geq N, \quad (5.30b)$$

where the controller state $v_{t+k|t}$ evolves as in (5.12b).

Although the design described in the previous section yields invariant ellipsoids with maximum volume, in which case the additional $c_{t+k|t}$ do not bring benefits in terms of domain of attraction, here consideration is given to polyhedral constraints, in which case the additional $c_{t+k|t}$ do bring benefits.

To account for constraints without introducing conservativeness, one must determine the regions where the predicted states and inputs will lie. Finding the exact regions however would lead to an exponential growth in computation. A strategy for the satisfaction of constraints while avoiding this problem was considered in [32] for systems with multiplicative uncertainty only. An extension of the technique of [32] is deployed here for the case of multiplicative and additive uncertainty in the presence of optimized polytopic controller dynamics.

An alternative (equivalent) description of the multiplicative uncertainty is considered. Let $\Psi^{(0)}$, $\tilde{B}^{(0)}$, $\tilde{D}^{(0)}$ be the nominal values of Ψ_{t+k} , $\tilde{B}_{t+k}$, $\tilde{D}_{t+k}$, where $\tilde{B}_{t+k} = [B_{t+k}^T, 0]^T$, and let $\Delta_i^\Psi = \Psi_i - \Psi^{(0)}$, $\Delta_i^{\tilde{B}} = \tilde{B}_i - \tilde{B}^{(0)}$, $\Delta_i^{\tilde{D}} = \tilde{D}_i - \tilde{D}^{(0)}$. Then the uncertain system matrices Ψ_{t+k} , $\tilde{B}_{t+k}$, $\tilde{D}_{t+k}$ of the dynamics of (5.13) satisfy:

$$\begin{aligned} \Psi_{t+k} &= \Psi^{(0)} + \Delta_{t+k}^\Psi, \quad \tilde{B}_{t+k} = \tilde{B}^{(0)} + \Delta_{t+k}^{\tilde{B}}, \quad \tilde{D}_{t+k} = \tilde{D}^{(0)} + \Delta_{t+k}^{\tilde{D}}, \\ (\Delta_{t+k}^\Psi, \Delta_{t+k}^{\tilde{B}}, \Delta_{t+k}^{\tilde{D}}) &\in \text{conv}\{(\Delta_i^\Psi, \Delta_i^{\tilde{B}}, \Delta_i^{\tilde{D}}) : i = 1, \dots, \nu_m\}. \end{aligned} \quad (5.31)$$

The dynamics of (5.13) are suitably modified to include the perturbations $c_{t+k|t}$, and are split into “certain” and “uncertain” dynamics as:

$$\chi_{t|t} + e_{t|t} = \begin{bmatrix} x_t \\ v_{t|t} \end{bmatrix}, \quad (5.32a)$$

$$\chi_{t+k+1|t} = \Psi^{(0)} \chi_{t+k|t} + \tilde{B}^{(0)} c_{t+k|t}, \quad k = 0, \dots, N-1, \quad (5.32b)$$

$$\chi_{t+k+1|t} = \Psi^{(0)} \chi_{t+k|t}, \quad k \geq N, \quad (5.32c)$$

$$e_{t+k+1|t} = \Psi_{t+k} e_{t+k|t} + \Delta_{t+k}^\Psi \chi_{t+k|t} + \Delta_{t+k}^{\tilde{B}} c_{t+k|t} + \tilde{D}_{t+k} w_{t+k}, \quad k = 0, \dots, N-1, \quad (5.32d)$$

$$e_{t+k+1|t} = \Psi_{t+k} e_{t+k|t} + \Delta_{t+k}^\Psi \chi_{t+k|t} + \tilde{D}_{t+k} w_{t+k}, \quad k \geq N, \quad (5.32e)$$

where $\chi_{t+k|t}$ and $e_{t+k|t}$ are the certain and uncertain components of the predictions, respectively. Thus the predicted state and controller state are $\chi_{t+k|t} + e_{t+k|t} = [x_{t+k|t}^T, v_{t+k|t}^T]^T$, and the predicted inputs of (5.30) are given by

$$u_{t+k|t} = [K \ L](\chi_{t+k|t} + e_{t+k|t}) + c_{t+k|t}, \quad k = 0, \dots, N-1, \quad (5.33a)$$

$$u_{t+k|t} = [K \ L](\chi_{t+k|t} + e_{t+k|t}), \quad k \geq N. \quad (5.33b)$$

Using 1-step-ahead calculations it is possible to find a tube for the augmented state of (5.32) that makes the predictions for the “uncertain” $e_{t+k|t}$ all lie inside polytopic cross-sections of the type

$$\Pi(\beta_{t+k|t}) = \{e : Ve \leq \beta_{t+k|t}\}, \quad (5.34)$$

where $V \in \mathbb{R}^{n_V \times 2n}$, provided that the initial uncertain state satisfies $e_{t|t} \in \Pi(\beta_{t|t})$. This constraint can be expressed as

$$Ve_{t|t} \leq \beta_{t|t}. \quad (5.35)$$

Then, only the initial uncertain state $e_{t|t}$, and the sequences of $\chi_{t+k|t}$, $\beta_{t+k|t}$ and $c_{t+k|t}$ are variables in the online optimization problem.

The following variation of Farkas' Lemma [31], [39] is key for the analysis to be presented.

Lemma 5.3. *Given two polytopic sets $\mathcal{P}_1 = \{x : F^1 x \leq g^1\}$ and $\mathcal{P}_2 = \{x : F^2 x \leq g^2\}$, the inclusion $\mathcal{P}_1 \subseteq \mathcal{P}_2$ holds if, and only if there exists a nonnegative matrix H such that*

$$\begin{aligned} HF^1 &= F^2, \\ Hg^1 &\leq g^2. \end{aligned} \tag{5.36}$$

Theorem 5.4. *If $e_{t+k|t} \in \Pi(\beta_{t+k|t})$, then $e_{t+k+1|t} \in \Pi(\beta_{t+k+1|t})$, $k = 0, \dots, N-1$, if there exist non-negative matrices $H_{i,j}$, $i = 1, \dots, \nu_m$, $j = 1, \dots, \nu_a$ such that*

$$H_{i,j}V = V(\Psi^{(0)} + \Delta_i^\Psi), \quad i = 1, \dots, \nu_m, \quad j = 1, \dots, \nu_a, \tag{5.37}$$

$$\begin{aligned} H_{i,j}\beta_{t+k|t} &\leq \beta_{t+k+1|t} - V \left(\Delta_i^\Psi \chi_{t+k|t} + \Delta_i^{\tilde{B}} c_{t+k|t} \right) - V \tilde{D}_i \tilde{w}_j, \quad k = 0, \dots, N-1, \\ i &= 1, \dots, \nu_m, \quad j = 1, \dots, \nu_a. \end{aligned} \tag{5.38}$$

Proof. This follows directly from applying Lemma 5.3 to the polytopes $\Pi(\beta_{t+k|t}) = \{e : Ve \leq \beta_{t+k|t}\}$ and $\{e : V(\Psi^i e + \tilde{B}_i c_{t+k|t} + \tilde{D}_i \tilde{w}_j) \leq \beta_{t+k+1|t}\}$, for $i = 1, \dots, \nu_m$, $j = 1, \dots, \nu_a$. $\square$

Using Theorem 5.4 it is possible to constrain the initial condition, $v_{t|t}$, and the degrees of freedom $c_{t+k|t}$ so that the cross-sections $\Pi(\beta_{t+k|t})$ are such that $Fx_{t+k|t} + Gu_{t+k|t} \leq 1$ for $k = 0, \dots, N-1$.

Corollary 5.5. *If $e_{t+k|t} \in \Pi(\beta_{t+k|t})$ for $k = 0, \dots, N-1$, the predictions of (5.32) satisfy constraints (5.4) if there exists a non-negative matrix H such that*

$$HV = [F + GK \quad GL], \tag{5.39}$$

$$H\beta_{t+k|t} \leq \underline{1} - [F + GK \quad GL]\chi_{t+k|t} - Gc_{t+k|t}, \quad k = 0, \dots, N-1. \quad (5.40)$$

Proof. This corollary follows from a direct application of Lemma 5.3 for the polytopes $\Pi(\beta_{t+k|t}) = \{e : Ve \leq \beta_{t+k|t}\}$ and $\{e : F(\chi_{t+k|t} + e) + G([K \quad L]\chi_{t+k|t} + e + c_{t+k|t}) \leq \underline{1}\}$, where $F(\chi_{t+k|t} + e_{t+k|t}) + G([K \quad L]\chi_{t+k|t} + e_{t+k|t} + c_{t+k|t}) \leq \underline{1}$ is the condition that is required for the satisfaction of constraints (5.4). $\square$

As usual in MPC, consideration is given to the dual-mode prediction paradigm whereby a terminal set is used in order to guarantee satisfaction of the constraints in Mode 2. Instead of the invariant ellipsoid E_z , the polytopic set

$$\Pi_f = \{z : V_f z \leq \underline{1}\}, \quad (5.41)$$

is considered, where Π_f is the maximal robust positively invariant set [12] of the dynamics of (5.13) subject to constraint (5.4). This guarantees that the predictions satisfy constraints in Mode 2.

Using Theorem 5.4, as in Corollary 5.5, it is possible to constrain the initial condition, $v_{t|t}$, and the degrees of freedom $c_{t+k|t}$ so that the cross-section $\Pi(\beta_{t+N|t})$ is such that $\chi_{t+N|t} \oplus \Pi(\beta_{t+N|t}) \subseteq \Pi_f$.

Corollary 5.6. *For $e_{t+N|t} \in \Pi(\beta_{t+N|t})$ the predictions of (5.32) satisfy constraints $\chi_{t+k|t} + e_{t+k|t} \in \Pi_f$ if there exists a non-negative matrix H_f such that*

$$H_f V = V_f, \quad (5.42)$$

$$H_f \beta_{t+k|t} \leq \underline{1} - V_f \chi_{t+N|t}. \quad (5.43)$$

Remark 5.7. *The conditions of Theorem 5.4, Corollary 5.5 and Corollary 5.6 would be necessary as well as sufficient if $H, H_{i,j}, H_f$ were computed online and were permitted to be different for each prediction instant (since, in order to be necessary*

conditions, they would have to depend on $\chi_{t+k|t}$ and $c_{t+k|t}$). However, for reasons of practical implementation $V, H, H_f, H_{i,j}$ are computed offline. In the interest of relaxing the constraints on the degrees of freedom (i.e. (5.38), (5.40) and (5.43)) a sensible way to choose $H, H_f, H_{i,j}$ is to minimize their induced l_∞ norm subject to (5.37), (5.39) and (5.42), respectively. In this selection $H_{i,j}$ does not depend on $\tilde{w}_j$ and thus, in the sequel $H_{i,j}$ will be replaced by H_i . Then, for all practical purposes, (5.38) can be replaced by

$$H_i \beta_{t+k|t} \leq \beta_{t+k+1|t} - V \left(\Delta_i^\Psi \chi_{t+k|t} + \Delta_i^{\tilde{B}} c_{t+k|t} \right) - d_i, \quad k = 0, \dots, N-1, \quad (5.44)$$

$$i = 1, \dots, \nu_m.$$

Here $d_i = [d_{i,1}, \dots, d_{i,n_V}]^T$, where $d_{i,p} = \max_{j=1, \dots, \nu_a} V_p \tilde{D}_i \tilde{w}_j \in \mathbb{R}^{n_V}$ for $i = 1, \dots, \nu_m$ and $p = 1, \dots, n_V$.

The offline computation of H_i , H , H_f , and V matrices does not take explicit account of the initial condition x_t and therefore will result in a certain degree of conservativeness. This implies that the terminal condition (5.43) is only sufficient but not necessary for $z_{t+N|t} \in \Pi_f$. Thus, although the satisfaction of (5.43) ensures $z_{t+N|t} \in \Pi_f$ which in turn implies $z_{t+N+1|t} \in \Pi_f$, it does not necessarily ensure $H_f \beta_{t+N+1|t} \leq \underline{1} - V_f \chi_{t+N+1|t}$. As a result condition (5.43) is not necessarily invariant, and therefore an MPC strategy using only that constraint as a terminal condition would not necessarily be recursively feasible.

Some approaches have been proposed to overcome this problem [27], [33]. Here, consideration will be given to the one presented in [33]. This approach was presented for systems with multiplicative uncertainty only, and here it is extended to systems with additive uncertainty as well. Following the ideas of [35], instead of a terminal invariant set for both $\chi_{t+k|t}$ and $\beta_{t+k|t}$, the constraints are enforced explicitly in Mode 2 with (5.38) and (5.40) for $k = N, \dots, N + N_2 - 1$, where $c_{t+k|t} = 0$ and N_2 is such

that, with some additional terminal conditions, (5.4) is guaranteed to be satisfied from the prediction instant $N + N_2$ onwards. This is summarized in the next theorem.

Theorem 5.8. *The predictions of (5.32) satisfy constraints (5.4) for $k \geq N$ if*

$$\begin{aligned} H_i \beta_{t+k|t} &\leq \beta_{t+k+1|t} - V \Delta_i^\Psi (\Psi^{(0)})^{k-N} \chi_{t+N|t} - d_i, \quad i = 1, \dots, \nu_m, \\ k &= N, \dots, N + N_2 - 1, \end{aligned} \quad (5.45)$$

$$H \beta_{t+k|t} \leq \underline{1} - [F + GK, GK] (\Psi^{(0)})^{k-N} \chi_{t+N|t}, \quad k = N, \dots, N + N_2 - 1, \quad (5.46)$$

$$\chi_{t+N|t} \in \Pi_f, \quad (5.47)$$

$$\|\beta_{t+N+N_2|t} - \beta_\infty\|_\infty \leq \delta, \quad (5.48)$$

where $\beta_\infty = \lim_{k \rightarrow \infty} \beta_k$ such that β_k satisfies the dynamics

$$\beta_{k+1} = \max_{i=1, \dots, \nu_m} \{H_i \beta_k + d_i\}, \quad (5.49)$$

and $N_2 \geq 0$ is large enough such that $\delta > 0$ satisfies

$$\delta \geq \frac{a}{1-h}, \quad (5.50)$$

$$\delta \leq \frac{1 - \|b\|_\infty}{\|H\|_\infty}, \quad (5.51)$$

where $a = \|\max_{i=1, \dots, \nu_m, \chi \in (\Psi^{(0)})^{N_2} \Pi_f} \{(H_i - I) \beta_\infty + V \Delta_i^\Psi \chi + d_i\}\|_\infty$, $h = \max_i \|H_i\|_\infty < 1$ and $b = \max_{\chi \in (\Psi^{(0)})^{N_2} \Pi_f} \{H \beta_\infty + [F + GK, GL] \chi\}$ ¹.

Proof. The satisfaction of (5.4) for $k = N, \dots, N + N_2 - 1$ follows from (5.45), (5.46) and the same analysis as in Corollary 5.5.

Condition (5.47) implies directly that $\chi_{t+N+N_2|t} = (\Psi^{(0)})^{N_2} \chi_{t+N|t} \in (\Psi^{(0)})^{N_2} \Pi_f$ and due to the invariance of Π_f it implies that $\chi_{t+N+N_2+k|t} = (\Psi^{(0)})^{N_2+k} \chi_{t+N|t} \in$

¹max is applied row-wise everywhere in Theorem 5.8

$(\Psi^{(0)})^{N_2} \Pi_f$ for all $k = 0, 1, \dots$. Condition (5.48) on the other hand is not inherently invariant, but (5.50) and the invariance of $\chi_{t+N+N_2+k|t} \in (\Psi^{(0)})^{N_2} \Pi_f$ make it invariant, as will be shown next. Note that (5.38) induces the following dynamics for $\beta_{t+k|t}$ for $k \geq N$: $\beta_{t+k+1|t} = \max_i (H_i \beta_{t+k|t} + V \Delta_i^\Psi \chi_{t+k|t} + d_i)$, from which $\beta_{t+k+1|t} - \beta_\infty = \max_i ((H_i - I) \beta_\infty + H_i (\beta_{t+k|t} - \beta_\infty) + V \Delta_i^\Psi \chi_{t+k|t} + d_i)$. Considering the infinity norm and applying the triangular inequality, this yields yields for $k = 0, 1, \dots$ $\|\beta_{t+N+N_2+k+1|t} - \beta_\infty\|_\infty \leq \|\max_i ((H_i - I) \beta_\infty + V \Delta_i^\Psi \chi_{t+N+N_2+k|t} + d_i)\| + \|H_i (\beta_{t+N+N_2+k|t} - \beta_\infty)\|$, and since $\chi_{t+N+N_2+k|t} = (\Psi^{(0)})^{N_2+k} \chi_{t+N|t} \in (\Psi^{(0)})^{N_2} \Pi_f$ for $k = 0, 1, \dots$ it follows that $\|\beta_{t+N+N_2+k+1|t} - \beta_\infty\|_\infty \leq a + h\delta$, which by condition (5.50) gives $\|\beta_{t+N+N_2+k+1|t} - \beta_\infty\|_\infty \leq a + h\delta \leq \delta$. This proves that if $\|\beta_{t+N+N_2|t} - \beta_\infty\|_\infty \leq \delta$ then $\|\beta_{t+N+N_2+k|t} - \beta_\infty\|_\infty \leq \delta$ for all $k = 0, 1, \dots$

Satisfaction of (5.51) and (5.48), by using its invariance, imply that for $k = 0, 1, \dots, 1 \geq \|H\|_\infty \delta + \|b\|_\infty \geq \|H\|_\infty \|\beta_{t+N+N_2+k|t} - \beta_\infty\|_\infty + \|b\|_\infty \geq \|H \beta_{t+N+N_2+k|t} - H \beta_\infty\|_\infty + \|b\|_\infty$, which in conjunction with (5.47), and its invariance, imply $1 \geq \|H \beta_{t+N+N_2+k|t} - H \beta_\infty\|_\infty + \|H \beta_\infty + [F + GK, GL] \chi_{t+N+N_2+k|t}\|_\infty \geq \|H \beta_{t+N+N_2+k|t} + [F + GK, GL] \chi_{t+N+N_2+k|t} - \underline{1}\|_\infty$. Thus $H \beta_{t+N+N_2+k} + [F + GK, GL] \chi_{t+N+N_2+k|t} \leq \underline{1}$ for $k = 0, 1, \dots$, which as in Corollary 5.5 guarantees the satisfaction of (5.4) beyond $N + N_2$. $\square$

For computational convenience, N_2 will be chosen as the minimal value such that there exists δ so that (5.50) and (5.51) are satisfied. Then, for this value of N_2 , δ is chosen to be the maximum value such that (5.50) and (5.51) are satisfied. Note that these conditions do not depend on N , and as a result the design of N_2 does not depend on N .

5.4.1.1 On the design of H , H_i and V

As pointed out in Remark 5.7, H , H_i and V are to be designed offline (note that H_f is not included here because constraints in Mode 2 will be ensured in a different

way). In the case of H and H_i , they are designed to have a minimal induced l_∞ norm subject to (5.37) and (5.39), respectively. For a fixed V , this can be achieved by solving a sequence of LPs as shown next.

Algorithm 5.9. *Design of H .*

1. For $l = 1, \dots, n_c$, solve the LP problem

$$\begin{aligned} & \underset{h_l}{\text{minimize}} && \sum_{m=1}^{n_V} h_{l,m} \\ & \text{subject to} && h_l V = \tilde{F}_l, \quad l = 1, \dots, n_c, \\ & && h_{l,m} \geq 0, \quad l = 1, \dots, n_c, \quad m = 1, \dots, n_V, \end{aligned} \tag{5.52}$$

where $h_l = [h_{l,1}, \dots, h_{l,n_V}] \in \mathbb{R}^{1 \times n_V}$ and $\tilde{F}_l$ denotes the l -th row of $\tilde{F} = [F + GK, GL]$.

2. H is given by

$$H = \begin{bmatrix} h_1 \\ \vdots \\ h_{n_c} \end{bmatrix}. \tag{5.53}$$

Algorithm 5.10. *Design of H_i .*

1. For $l = 1, \dots, n_V$, solve the LP problem

$$\begin{aligned} & \underset{h_l}{\text{minimize}} && \sum_{m=1}^{n_V} h_{l,m} \\ & \text{subject to} && h_l V = V_l \Psi_i, \quad l = 1, \dots, n_V, \\ & && h_{l,m} \geq 0, \quad l = 1, \dots, n_V, \quad m = 1, \dots, n_V, \end{aligned} \tag{5.54}$$

where $h_l = [h_{l,1}, \dots, h_{l,n_V}] \in \mathbb{R}^{1 \times n_V}$ and V_l denotes the l -th row of V .

2. H_i is given by

$$H_i = \begin{bmatrix} h_1 \\ \vdots \\ h_{n_V} \end{bmatrix}. \quad (5.55)$$

In order to get some insight for the design of V note the following. A necessary condition for the stability of the β dynamics is that all the eigenvalues of all H_i are inside the unit circle. Furthermore, the smaller the eigenvalues, the tighter the bounding of the uncertainty e , as per (5.34), will be. Additionally, from condition (5.50) it is necessary that the induced l_∞ norm of all H_i is less than 1, and, the smaller this norm is, the more relaxed (5.50) will be. However, V and H_i are related as shown next.

Theorem 5.11. *If H_i is designed as in Algorithm 5.10, then it satisfies*

$$\|H_i\|_\infty = \lambda, \quad (5.56a)$$

$$\rho(H_i) \leq \lambda, \quad (5.56b)$$

where $\rho(\cdot)$ denotes the spectral radius and λ is the minimum value such that $\mathcal{V} = \{z : Vz \leq 1\}$ is a λ -contractive set for the system $z^+ = \Psi z$.

Proof. From [12] (Proposition 7.23) $\mathcal{V}$ is a λ -contractive set for system $z^+ = \Psi z$ if and only if there exists a matrix $M_i \geq 0$ such that $M_i V = V \Psi_i$ and $M_i \underline{1} \leq \lambda \underline{1}$. The last condition is equivalent to $\|M_i\|_\infty \leq \lambda$. But the design of H_i in Algorithm 5.10 gives a minimum infinity norm H_i such that the $H_i \geq 0$ and $H_i V = V \Psi_i$. Then clearly $\|H_i\|_\infty = \lambda$ where λ is the minimum value such that $\mathcal{V}$ is λ -contractive.

(5.56b) follows from the fact that $\rho(X) \leq \|X\|_\infty$ for any square and real matrix X . □

Remark 5.12. Based on Theorem 5.11 and on testing for different systems, it is convenient to choose V as the matrix such that $\mathcal{V}$ is the maximal invariant set for system $z^+ = (1/\bar{\lambda})\Psi z$, $\sigma(\Psi) < \bar{\lambda} < 1$ where σ denotes the joint spectral radius, under the box constraints $-\underline{1} \leq z \leq \underline{1}$. This guarantees that the set $\{z : Vz \leq \underline{1}\}$ is λ -contractive for system $z^+ = \Psi z$, with $\lambda \leq \bar{\lambda}$ and then $\|H_i\|_\infty = \lambda$ and $\rho(H_i) \leq \lambda$. Smaller values of λ make (5.34) tighter and (5.50) more relaxed, but the number of facets in V (and therefore the number of variables in $\beta_{k+l|k}$) will increase. $\bar{\lambda}$ is then chosen as a sensible compromise between the tightness of (5.34) and relaxation of (5.50), and limiting the number of rows in V .

5.4.2 Cost function

Consideration is given to the worst-case $\mathcal{H}_\infty$ type of cost

$$J(\xi_{t|t}) = \max_{(A,B,D,w)_{t+k}} \sum_{k=0}^{\infty} (x_{t+k|t}^T Q x_{t+k|t} + u_{t+k|t}^T R u_{t+k|t} - \gamma^2 w_{t+k}^T w_{t+k}), \quad (5.57)$$

where $Q, R \succ 0$ and $\gamma > 0$.

For this strategy, the prediction dynamics are generated more compactly as

$$\xi_{t|t} = \begin{bmatrix} x_t \\ v_{t|t} \\ \mathbf{c}_t \end{bmatrix}, \quad (5.58)$$

$$\xi_{t+k+1|t} = \Theta_{t+k} \xi_{t+k|t} + \hat{D}_{t+k} w_{t+k}, \quad k = 0, 1, 2, \dots,$$

where

$$\mathbf{c}_t = \begin{bmatrix} c_{t|t} \\ \vdots \\ c_{t+N-1|t} \end{bmatrix}, \quad \begin{array}{l} \Theta_{t+k} \in \text{conv}\{\Theta_i\}, \\ \hat{D}_{t+k} \in \text{conv}\{\hat{D}_i\}, \end{array} \quad \Theta_i = \begin{bmatrix} \Phi_i & B_i L & B_i E \\ 0 & M_i & 0 \\ 0 & 0 & M \end{bmatrix}, \quad \hat{D}_i = \begin{bmatrix} \tilde{D}_i \\ 0 \end{bmatrix},$$

and $E \in \mathbb{R}^{n_u \times Nn_u}$, $M \in \mathbb{R}^{Nn_u \times Nn_u}$ are given by

$$E = \begin{bmatrix} I & 0 & \dots & 0 \end{bmatrix}, \quad M = \begin{bmatrix} 0 & I & 0 & \dots \\ 0 & 0 & \ddots & 0 \\ \vdots & & \ddots & I \\ 0 & \dots & 0 & 0 \end{bmatrix}$$

where $I \in \mathbb{R}^{n_u \times n_u}$.

As in [19] and [34], since the exact evaluation of the cost (5.57) requires computation that grows exponentially with N , the online optimization is actually performed by minimizing a quadratic upper bound.

Theorem 5.13. *For a given $\gamma > 0$, the quadratic cost*

$$V(\xi_{t|t}) = \xi_{t|t}^T W \xi_{t|t} \quad (5.59)$$

is an upper bound of the predicted cost of (5.57) such that

$$J(\xi_{t|t}) \leq \xi_{t|t}^T W \xi_{t|t}, \quad (5.60)$$

where $W = W^T \succ 0$, if, for all admissible $(A, B, D)_{t+k}$ and w_{t+k} of (5.2), the quadratic bound satisfies for $k = 0, 1, 2, \dots$

$$\xi_{t+k|t}^T W \xi_{t+k|t} - \xi_{t+k+1|t}^T W \xi_{t+k+1|t} \geq x_{t+k|t}^T Q x_{t+k|t} + u_{t+k|t}^T R u_{t+k|t} - \gamma^2 w_{t+k}^T w_{t+k}, \quad (5.61)$$

Furthermore, under the assumption that $0 \in \text{conv}\{\tilde{w}_j, j = 1, \dots, \nu_a\}$, (5.61) holds if, and only if W satisfies

$$\begin{bmatrix} W & 0 \\ 0 & 0 \end{bmatrix} - \begin{bmatrix} \Theta_i^T & 0 \\ \hat{D}_i^T & I \end{bmatrix} \begin{bmatrix} W & 0 \\ 0 & 0 \end{bmatrix} \begin{bmatrix} \Theta_i & \hat{D}_i \\ 0 & I \end{bmatrix} \succeq \begin{bmatrix} \tilde{Q} & 0 \\ 0 & -\gamma^2 I \end{bmatrix}, \quad i = 1, \dots, \nu_m \quad (5.62)$$

where I is the identity matrix in $\mathbb{R}^{n_w \times n_w}$ and

$$\tilde{Q} = \begin{bmatrix} Q + K^T R K & K^T R L & K^T R E \\ * & L^T R L & L^T R E \\ * & * & E^T R E \end{bmatrix}$$

Proof. The cost bound of (5.60) follows from summing (5.61) for each prediction time from current time to infinity.

Condition (5.61) can be rewritten as

$$\begin{aligned} & \xi_{t+k|t}^T W \xi_{t+k|t} - (\Psi_{t+k} \xi_{t+k|t} + \hat{D}_{t+k} w_{t+k})^T W (\Psi_{t+k} \xi_{t+k|t} + \hat{D}_{t+k} w_{t+k}) \geq \\ & \xi_{t+k|t}^T \tilde{Q} \xi_{t+k|t} - \gamma^2 w_{t+k}^T w_{t+k}, \end{aligned}$$

which is equivalent to

$$\begin{bmatrix} \xi_{t+k|t}^T & w_{t+k}^T \end{bmatrix} \begin{bmatrix} W - \Theta_{t+k}^T W \Theta_{t+k} - \tilde{Q} & -\Theta_{t+k}^T W \hat{D}_{t+k} \\ * & \gamma^2 I - \hat{D}_{t+k}^T W \hat{D}_{t+k} \end{bmatrix} \begin{bmatrix} \xi_{t+k|t} \\ w_{t+k} \end{bmatrix} \geq 0 \quad (5.63)$$

for all admissible $(A_{t+k}, B_{t+k}, D_{t+k})$ and w_{t+k} . Since $0 \in \text{conv}\{\tilde{w}_j : j = 1, \dots, \nu_a\}$ this implies that the matrix in (5.63) is positive semi-definite, and by re-arranging terms this is equivalent to

$$\begin{bmatrix} W & 0 \\ 0 & 0 \end{bmatrix} - \begin{bmatrix} \Theta_{t+k}^T & 0 \\ \hat{D}_{t+k}^T & I \end{bmatrix} \begin{bmatrix} W & 0 \\ 0 & 0 \end{bmatrix} \begin{bmatrix} \Theta_{t+k} & \hat{D}_{t+k} \\ 0 & I \end{bmatrix} \preceq \begin{bmatrix} \tilde{Q} & 0 \\ 0 & -\gamma^2 I \end{bmatrix},$$

which, by convexity, is equivalent to (5.62) which invokes the above inequality at the vertices of the admissible (A, B, D) . $\square$

Remark 5.14. Condition (5.62) in Theorem 5.13 is necessary and sufficient for (5.61), unlike the LMI of [34] which is reported to be only sufficient for (5.61). It should be noted however that condition (5.62) is equivalent to the LMI in [34] (which

can be proved by a procedure based on Schur complementation), and therefore the condition in [34] is in fact both necessary and sufficient.

Remark 5.15. γ has the interpretation of a bound on the induced l_2 norm from the disturbance to an output given by $y = [x^T Q^{\frac{1}{2}} \quad u^T R^{\frac{1}{2}}]^T$ when (5.61) is satisfied, and should be made as small as possible, e.g. by minimizing γ subject to (5.62). On the other hand, in the interest of reducing the upper bound on the minimax cost in (5.60), a convenient procedure is to minimize $\text{tr}(W)$ subject to (5.62). One could also attempt to reduce this bound on the cost subject to γ being smaller than some upper bound. A compromise then is needed here between the disturbance rejection parameter γ and the tightness of the upper bound on the cost in (5.60). Note also that the minimum attainable value of γ depends on the matrices of the optimized polytopic dynamics M_i, S_i, L and on the state feedback K . Therefore, in the case that the minimum value of γ is not satisfactory, one could use the method of [34] which determines M_i, S_i, L along with γ , while constraining γ to be smaller than a given satisfactory upper bound.

Remark 5.16. W is positive definite because the top left block of (5.62) yields $W - \Theta_i^T W \Theta_i \succeq \tilde{Q}$, where $\tilde{Q}$ is positive definite and Θ (and Θ_i) is stable. Since the cost of (5.57) is always equal or greater than zero and the upper bound is zero when initial augmented state is $\xi = 0$, it is clear that the bound is tight at least for this case.

5.4.3 Overall MPC algorithm and properties

The offline designs in this RMPC strategy as described in Section 5.4.1 are summarized in the next Algorithm.

Algorithm 5.17 (Offline). (i) Design the state feedback gain K to be optimal in some sense (e.g. LQR for the nominal system description). (ii) Design $L, M_i, S_i, i = 1, \dots, \nu_m$ by maximization of $\log \det(T^T P T)$ s.t. (5.15) and (5.16) via the reformulated problem of (A5) as presented in Appendix A. (iii) Choose $\bar{\lambda}$ and then set V to be the maximal invariant set for system $z^+ = (1/\bar{\lambda})\Psi z$, $\sigma(\Psi) < \bar{\lambda} < 1$, under the

box constraints $-1 \leq z \leq 1$. (iv) Compute H and H_i , $i = 1, \dots, \nu_m$, according to Algorithms 5.9 and 5.10, respectively. (v) Select the minimum N_2 for which there exists δ such that (5.50) and (5.51) are satisfied. Then, for that value of N_2 , select the maximum value of δ such that (5.50) and (5.51) are satisfied. (vi) Compute W as per Remark 5.15. (vii) Choose N .

The online stage of the RMPC strategy is analysed next.

The vector of decision variables in the online optimization, $\mathbf{d}$, consists of $e_{t|t}$, the sequences (in vectorized form) $\chi_t = \{\chi_{t|t}, \dots, \chi_{t+N|t}\}$ and $\beta_t = \{\beta_{t|t}, \dots, \beta_{t+N+N_2|t}\}$, and the vector $\mathbf{c}_t$. The set of admissible variables $\mathcal{D}(x)$ is given by

$$\mathcal{D}(x) = \{\mathbf{d} : (5.32b), (5.35), (5.40), (5.44), (5.45), (5.46), (5.47), (5.48) \text{ hold}\}, \quad (5.64)$$

and the domain of attraction is given by

$$\mathcal{X} = \{x : \mathcal{D}(x) \neq \emptyset\}. \quad (5.65)$$

The online procedure consists of solving the optimal control problem $\mathbb{P}(x_t)$ that is defined by

$$V^0(x_t) = \min_{\mathbf{d}} V([\chi_{t|t}^T, \mathbf{c}_t^T]^T), \text{ s.t. } \mathbf{d} \in \mathcal{D}(x_t), \quad (5.66a)$$

$$\mathbf{d}^0(x_t) = \arg \min_{\mathbf{d}} V([\chi_{t|t}^T, \mathbf{c}_t^T]^T), \text{ s.t. } \mathbf{d} \in \mathcal{D}(x_t). \quad (5.66b)$$

Then, the RMPC control law is given by

$$\kappa^0(x_t) = Kx_t + Lv_{t|t}^0(x_t) + c_{t|t}^0(x_t), \quad (5.67)$$

where $v_{t|t}^0(x_t) = [0 \quad I](e_{t|t}^0(x_t) + \chi_{t|t}^0(x_t))$ according to (5.32a), and $e_{t|t}^0(x_t)$, $\chi_{t|t}^0(x_t)$ and $c_{t|t}^0(x_t)$ are part of the optimal $\mathbf{d}^0(x_t)$. The control input at time t is then given

by $u_t := \kappa^0(x_t)$.

The optimal control problem $\mathbb{P}(x_t)$ can be formulated as the QP:

$$\begin{aligned}
 & \underset{\mathbf{d}}{\text{minimize}} && V([x_t^T, v_{t|t}^T, \mathbf{c}_t^T]^T) \\
 & \text{subject to} && \chi_{t+k+1|t} = \Psi^{(0)}\chi_{t+k|t} + \tilde{B}^{(0)}c_{t+k|t}, \quad k = 0, \dots, N-1, \\
 & && V e_{t|t} \leq \beta_{t|t}, \\
 & && \xi_{t|t} + e_{t|t} = [x_t^T, v_{t|t}^T]^T, \\
 & && H_i \beta_{t+k|t} \leq \beta_{t+k+1|t} - V \left(\Delta_i^\Psi \chi_{t+k|t} + \Delta_i^{\tilde{B}} c_{t+k|t} \right) - V \tilde{D}_i w_j, \\
 & && \quad k = 0, \dots, N-1, \quad i = 1, \dots, \nu_m, \\
 & && H_i \beta_{t+k|t} \leq \beta_{t+k+1|t} - V \Delta_i^\Psi (\Psi^{(0)})^{k-N} \chi_{t+N|t} - d_i, \\
 & && \quad k = N, \dots, N+N_2-1, \quad i = 1, \dots, \nu_m, \\
 & && H \beta_{t+k|t} \leq \underline{1} - [F + GK, GL] \chi_{t+k|t} - G c_{t+k|t}, \quad k = 0, \dots, N-1, \\
 & && H \beta_{t+k|t} \leq \underline{1} - [F + GK, GL] (\Psi^{(0)})^{k-N} \chi_{t+N|t}, \quad k = N, \dots, N+N_2-1, \\
 & && V_f \chi_{t+N|t} \leq \underline{1}, \\
 & && \beta_{t+N+N_2|t} - \beta_\infty \leq \delta \underline{1}, \\
 & && \beta_\infty - \beta_{t+N+N_2|t} \leq \delta \underline{1},
 \end{aligned} \tag{5.68}$$

This is a QP with $n_{var} = N(4n + n_V + n_u) + N_2 n_V + 2n + n_V$ variables subject to $n_{ineq} = N(2n + \nu_m n_v + n_c) + N_2(\nu_m n_V + n_c) + n_f + 4n_V$ inequalities, and $n_{eq} = 2nN$ equalities with n_V and n_f being the number of rows in V and V_f , respectively. Clearly n_{var} and $n_{cons} = n_{ineq} + n_{eq}$ grow linearly with N .

Theorem 5.18. *For system (5.1), the closed-loop application of the control law of (5.67) guarantees feasibility of (5.4), that the optimal control problem of (5.66) is recursively feasible and that $\mathcal{X}$ is invariant.*

Proof. The candidate solution at time $t+1$, $\tilde{\mathbf{d}}$, is built from the optimal solution at

time t , $\mathbf{d}^0(x_t)$ according to the *receding horizon* strategy such that

$$\tilde{e}_{t+1|t+1} = \Psi_t e_{t|t}(x_t) + \Delta_t^\Psi \chi_{t|t}^0(x_t) + \tilde{D}_t w_t, \quad (5.69a)$$

$$\tilde{\chi}_{t+1+k|t+1} = \chi_{t+k+1|t}^0(x_t), \quad k = 0, \dots, N-1, \quad (5.69b)$$

$$\tilde{\chi}_{t+1+N|t+1} = \Psi^{(0)} \chi_{t+N|t}(x_t), \quad (5.69c)$$

$$\tilde{c}_{t+1+k|t+1} = c_{t+k+1|t}^0(x_t), \quad k = 0, \dots, N-2, \quad (5.69d)$$

$$\tilde{c}_{t+1+N|t+1} = 0, \quad (5.69e)$$

$$\tilde{\beta}_{t+1+k|t+1} = \beta_{t+k+1|t}^0(x_t), \quad k = 0, \dots, N+N_2-1, \quad (5.69f)$$

$$\tilde{\beta}_{t+1+N+N_2|t+1} = \max_{i=1, \dots, \nu_m} H_i \beta_{t+N+N_2|t}^0(x_t) + V \Delta_i^\Psi (\Psi^{(0)})^{N_2-1} \chi_{k+N|t} + d_i. \quad (5.69g)$$

This construction is such that $\tilde{\mathbf{d}}$ satisfies (5.32b), (5.44) and (5.45). Additionally from Theorem 5.4, the fact that $\mathbf{d}^0(x_t)$ satisfies (5.44) for $k=0$ implies that $e_{t+1|t} \in \Pi(\beta_{t+1|t}^0(x_t))$, and by construction $\tilde{e}_{t+1|t+1} \in \Pi(\beta_{t+1|t}^0(x_t))$ so that $\tilde{\mathbf{d}}$ satisfies (5.35).

As indicated in Theorem 5.8, conditions (5.47) and (5.48) are invariant provided that (5.50) holds, and then if $\mathbf{d}^0(x_t)$ satisfies them it follows that $\tilde{\chi}_{t+1+N|t+1} = \Psi^{(0)} \chi_{t+N|t}^0(x_t) \in \Pi_f$ and $\|\beta_{t+N+N_2+1|t} - \beta_\infty\| \leq \delta$ where $\beta_{t+N+N_2+1|t} = \tilde{\beta}_{t+1+N+N_2|t+1}$. Thus $\tilde{\mathbf{d}}$ satisfies (5.47) and (5.48).

By inspection $\tilde{\mathbf{d}}$ satisfies (5.40) and (5.46) for $k = N+1, \dots, N+N_2-2$. However, since $\mathbf{d}^0(x_t)$ satisfies (5.47) and (5.48) it follows that $\mathbf{d}^0(x_t)$ satisfies (5.46) for $k = N+N_2$, and then $\tilde{\mathbf{d}}$ satisfies (5.46) for $k = N+N_2-1$.

Finally, $\tilde{\mathbf{d}}$ is a feasible vector of decision variables for the optimal problem $\mathbb{P}(x_{t+1})$ provided that $\mathbb{P}(x_t)$ was feasible. Then, by a recursive argument, it can be established that the problem of 5.66 will remain feasible for all future instants. Invariance of $\mathcal{X}$ follows directly. Satisfaction of (4.3) follows from the satisfaction of (5.40) at each instant and Corollary 5.5. $\square$

Theorem 5.19. *For system (5.1), the closed-loop application of the control law of (5.67) guarantees the upper bound for the closed-loop behaviour of system (5.1) given*

by

$$\sum_{t=0}^{\infty} (x_t^T Q x_t + u_t^T R u_t - \gamma^2 w_t^T w_t) \leq \xi_{0|0}^0(x_0)^T W \xi_{0|0}^0(x_0). \quad (5.70)$$

Proof. The candidate solution at time $t+1$, $\tilde{\mathbf{d}}$, that is built from the optimal solution at time t , $\mathbf{d}^0(x_t)$ according to (5.69) is guaranteed to be feasible as shown in Theorem 5.18. Additionally, the predicted cost satisfies by construction

$$\begin{aligned} \xi_{t|t}^0(x_t)^T W \xi_{t|t}^0(x_t) - \xi_{t+1|t+1}^0(x_t)^T W \xi_{t+1|t+1}^0(x_t) \geq \\ (x_t^T Q x_t + \kappa^0(x_t)^T R \kappa^0(x_t) - \gamma^2 w_t^T w_t) \end{aligned}$$

which, when summed from 0 to ∞ , yields the desired result of (5.70). $\square$

The dependence of the domain of attraction $\mathcal{X}$ on the horizon N is studied next; as in Chapter 4 this dependence is now made explicit by denoting the domain of attraction as $\mathcal{X}_N$. Note that as stated in Section 5.4.1 Algorithm 5.17 the choice of N_2 depends directly on the system and the selections of V, H, H_i . Because of this, the dependence on N_2 is not studied.

Theorem 5.20. *For fixed δ and N_2 such that (5.50) and (5.51) are satisfied, the domain of attraction of system (5.1) under the control law of (5.67) satisfies*

$$\mathcal{X}_N \subseteq \mathcal{X}_{N+1}, \quad N = 0, 1, \dots \quad (5.71)$$

Proof. If $x_t \in \mathcal{X}_N$ then there exist $v_{t|t}$, $e_{t|t}$, χ_t , β_t and $\mathbf{c}_t$ such that (5.32b), (5.35), (5.40), (5.44), (5.45), (5.46), (5.47) and (5.48) are satisfied. The satisfaction of these constraints is now analysed for the case $\tilde{N} = N+1$, where the sequences are extended with $\chi_{k+\tilde{N}|t} = \Psi^{(0)} \chi_{t+N|t}$, $\beta_{t+\tilde{N}+N_2|t} = \max_{i=1, \dots, \nu_m} H_i \beta_{t+N+N_2|t} + V \Delta_i^\Psi (\Psi^{(0)})^{N_2} \chi_{t+N|t} + d_i$ and $c_{t+\tilde{N}-1|t} = 0$.

Conditions (5.32b) and (5.35) are trivially satisfied. Satisfaction of (5.44) and (5.45) when using N guarantees the satisfaction of (5.44) and of (5.45) for $k =$

$\tilde{N}, \dots, \tilde{N} + N_2 - 1$. However, as indicated in Theorem 5.8, conditions (5.47) and (5.48) are invariant provided that (5.50) holds, so then they are satisfied for $\tilde{N}$ with $\chi_{k+\tilde{N}|t} = \Psi^{(0)} \chi_{k+N|t}$ and $\beta_{t+\tilde{N}+N_2|t} = \max_{i=1, \dots, \nu_m} H_i \beta_{t+\tilde{N}_2-1|t} + V \Delta_i^\Psi (\Psi^{(0)})^{N_2-1} \chi_{k+\tilde{N}|t} + d_i$, which is such that (5.45) is satisfied for $k = N + N_2 = \tilde{N} + N_2 - 1$.

Likewise, satisfaction of (5.40) and (5.46) when using N directly guarantees the satisfaction, when using $\tilde{N}$, of (5.40) and of (5.46) for $k = \tilde{N}, \dots, \tilde{N} + N_2 - 2$. However, satisfaction of (5.47) and (5.48) when using N , provided that (5.51) holds, implies that (5.46) is satisfied for $k = N + N_2 = \tilde{N} + N_2 - 1$. As a result, the optimal control problem $\mathbb{P}(x)$ is feasible when using $\tilde{N} = N + 1$, which yields $\mathcal{X}_N \subseteq \mathcal{X}_{N+1}$. $\square$

5.5 Simulation results

Consider the system defined by

$$A_1 = \begin{bmatrix} -1.751 & -0.624 \\ -1.432 & 0.474 \end{bmatrix}, \quad A_2 = \begin{bmatrix} -1.549 & -0.624 \\ -1.432 & 0.474 \end{bmatrix}, \quad B_1 = \begin{bmatrix} 1.558 \\ 1.317 \end{bmatrix}, \quad B_2 = \begin{bmatrix} 1.790 \\ 1.593 \end{bmatrix},$$

with $D_1 = D_2 = 0.1I_{2 \times 2}$, and subject to the input constraint $|u_t| \leq 1$.

The following selections are made for performing the offline design of Algorithm 5.17. The state feedback gain K is chosen to be the LQ optimal for the nominal system, given by the average of (A_1, B_1) and (A_2, B_2) and $w_t = 0$ for all t , where $Q = I_{2 \times 2}$ and $R = 0.01$ are used for this LQ design and the cost of (5.57). According to Remark 5.15, γ is fixed at its minimum value for which (5.62) has a solution for W , and for this value of γ , then W is chosen to have minimum trace, subject to (5.62). $\bar{\lambda} = 0.89$, and as a result V is such that $V \in \mathbb{R}^{26 \times 4}$.

The domain of attraction of the proposed strategy is computed for different values of N , and their areas $A(\mathcal{X}_N)$ are displayed in Table 5.1.

It can be seen from Table 5.1 that larger domains of attraction can be obtained as N increases and this comes at the cost of a linear increase in the number of variables

N	0	1	2	3	4	5	6
Area	144.8	161.0	174.3	183.6	188.4	190.3	191.7
Vars	134	169	204	239	274	309	344
Ineqs	340	398	456	514	572	630	688
Eqs	0	4	8	12	16	20	24

Table 5.1: Domain of Attraction and computational aspects

and constraints in the online optimization problem.

It is worth noting that the projection on the x -space of Π_f is the area of the domain of attraction of the strategy where only the optimized polytopic dynamics are used, with no additional degrees of freedom, and the constraints are guaranteed by imposing that the augmented state is inside Π_f (as in Section III of [34]). This area is $A(\text{Proj}_x(\Pi_f)) = 161.5$. Clearly $A(\mathcal{X}_0) < A(\text{Proj}_x(\Pi_f))$, which is a consequence of the fact that the strategy to guarantee the feasibility of the predictions is based on Lemma 5.3 which introduces a degree of conservativeness. However, the introduction of the degrees of freedom in $\mathbf{c}_t$ improves the domain of attraction so as to make it bigger than $\text{Proj}_x(\Pi_f)$.

For $N = 4$, Figure 5.1a presents state space trajectories for 50 different realizations of the uncertainty and Figure 5.1b presents the level of satisfaction of the constraints. These show how the domain of attraction is an invariant set guaranteeing constraint satisfaction despite the presence of multiplicative and additive uncertainty.

For comparison purposes, consider a strategy where the predicted control policy is given by $u_{t+k|t} = Kx_{t+k|k} + c_{t+k|t}$, for $k = 0, \dots, N - 1$ and $u_{t+k|t} = Kx_{t+k|k}$ for $k \geq N$, and the constraints are satisfied by imposing that $z_t = [x_t^T, \mathbf{c}_t^T]^T$ is inside the maximal invariant set for the augmented dynamics of z_t . The areas of the domains of attraction obtained with this strategy, the number of variables and the number of inequalities of this strategy (which has no equality constraints) are presented in Table 5.2. It can be seen that the size of the domain of attraction for any $N \leq 12$ is less than what can be obtained with the optimized dynamics for any $N \leq 6$. Although N could be potentially increased over 12, with which one

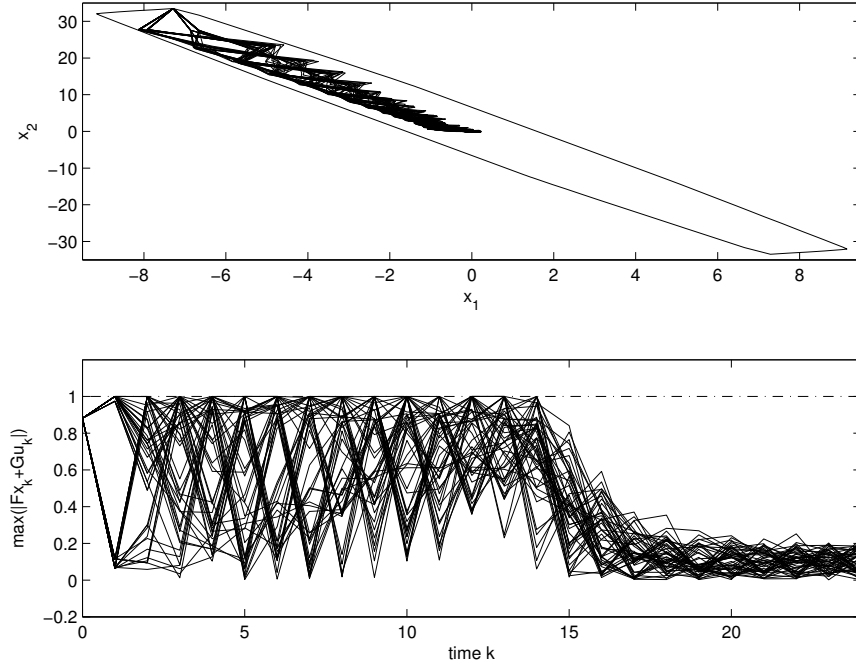


Figure 5.1: Closed-loop trajectories of (a) x_t , (b) $\max(|Fx_t + Gu_t|)$, for 50 random realizations of the uncertainty

might obtain larger domains of attraction than with the optimized dynamics, the exponential growth of the number of constraints makes this impractical.

N	0	2	4	6	8	10	12
Area	4.881	11.632	20.755	33.607	51.547	76.028	108.589
Vars	0	2	4	6	8	10	12
Ineqs	6	30	126	508	2020	7948	31260

Table 5.2: Domain of attraction and computational aspects for closed-loop paradigm

5.6 Conclusions

The approach of [19] for the optimization of prediction dynamics with the view to maximizing the volume of invariant ellipsoids has been extended to the case of mixed multiplicative and additive uncertainty. Like the approach presented in [34], this extension deploys disturbance feedback but is shown to depend on conditions which are both necessary and sufficient.

The optimized dynamics together with perturbations on the control laws over

the first N prediction steps are used to propose an effective prediction dual-mode prediction based RMPC algorithm which does not rely on the restrictive assumption that the future model uncertainty is known when predictions are made. The size of the resulting online optimization problem grows linearly with the prediction horizon, and it thus is suitable for large prediction horizons.

The optimized polytopic dynamics yield the largest ellipsoidal invariant set and the addition of new degrees of freedom could not therefore increase the size of the ellipsoidal invariant set. However, the optimized polytopic dynamics are not necessarily optimal in terms of polytopic invariant sets. Hence the use of polytopic constraints gives space for the introduction of extra degrees of freedom in the control policy which are used to improve the domain of attraction.

The strategy for constraint satisfaction uses the matrix V which, as mentioned in Remark 5.12, is selected to be the maximal invariant set for system $z^+ = (1/\bar{\lambda})\Psi z$, $\sigma(\Psi) < \bar{\lambda} < 1$, under the box constraints $-\underline{1} \leq z \leq \underline{1}$. As discussed in that remark, smaller $\bar{\lambda}$ makes the bounding of the uncertainty e , as per (5.34), tighter. Additionally, smaller $\bar{\lambda}$ make (5.50) more relaxed. The effect of smaller $\bar{\lambda}$, however, is that V will tend to have a greater number of rows and thus imply a greater number of constraints. This is because the closer the eigenvalues of $(1/\bar{\lambda})\Psi$ are to 1, the maximal invariant set above will have more inequalities. Then, for systems where $\sigma(\Psi)$ is relatively large (i.e. close to 1) the eigenvalues of $(1/\bar{\lambda})\Psi$ will be close to 1, and thus V will involve many rows and thereby rendering the RMPC strategy computationally expensive. In this sense, the complexity of the strategy does not only depend on the prediction horizons N and N_2 , but also on the eigenvalues of the optimized polytopic dynamics. This suggests two possible future research topics: that of finding alternatives for a more convenient selection of V ; or one that reduces the number of inequalities in, or finds alternative ways to impose, conditions (5.44), (5.40), (5.45) and (5.46) so that the computational demand of the implementation of the RMPC strategy is not so

sensitive on the size of V .

Appendix A

To convexify the problem of selecting the P that maximizes the volume of the invariant ellipsoid, E_x , use is made of the parameterization

$$P = \begin{bmatrix} \mathcal{W} & \mathcal{U} \\ \mathcal{U}^T & \mathcal{Q} \end{bmatrix}, \quad P^{-1} = \begin{bmatrix} \mathcal{Y} & \mathcal{V} \\ \mathcal{V}^T & \mathcal{S} \end{bmatrix}, \quad \Pi = \begin{bmatrix} I & \mathcal{Y} \\ 0 & \mathcal{V}^T \end{bmatrix}. \quad (\text{A1})$$

(5.16) is post- and pre-multiplied by a block diagonal matrix and its transpose, whose diagonal blocks are $\Pi, \Pi, 1$, and (5.15) is post- and pre-multiplied by a block diagonal matrix and its transpose, whose diagonal blocks are I, Π . Such multiplications, using that $\mathcal{W}\mathcal{Y} + \mathcal{U}\mathcal{V}^T = I$ and $\mathcal{U}^T\mathcal{Y} + \mathcal{Q}\mathcal{V}^T = 0$, which follow from $PP^{-1} = P^{-1}P = I$, yield the equivalent LMIs, for fixed α , in terms of new free variables

$$\begin{bmatrix} \mathcal{W} & I & \Phi_i \mathcal{W} + B_i \tilde{L} & \Phi_i & D_i \tilde{w}_j \\ * & \mathcal{Y} & \Gamma_i & \mathcal{Y} \Phi_i & (\mathcal{Y} D_i + \tilde{S}_i) \tilde{w}_j \\ * & * & \alpha \mathcal{W} & \alpha I & 0 \\ * & * & * & \alpha \mathcal{Y} & 0 \\ * & * & * & * & 1 - \alpha \end{bmatrix} \succeq 0, \quad (\text{A2})$$

$$\begin{bmatrix} Z & (F + GK)\mathcal{W} + G\tilde{L} & F + GK \\ * & \mathcal{W} & I \\ * & * & \mathcal{Y} \end{bmatrix} \succeq 0, \quad (\text{A3})$$

where the new free variables are defined by

$$\tilde{L} = L\mathcal{U}^T, \quad \Gamma_i = \mathcal{Y}\Phi_i\mathcal{W} + \mathcal{Y}B_iL\mathcal{U}^T + \mathcal{V}M_i\mathcal{U}^T, \quad \tilde{S}_i = \mathcal{V}S_i \quad (\text{A4})$$

from which it is clear that the parameters $\mathcal{S}$ and $\mathcal{Q}$ play no direct role.

Note that $TPT^T = \mathcal{W}$. Then, in order to find the maximum volume invariant ellipsoid, one must solve the optimization problem

$$\begin{aligned} & \underset{\mathcal{Y}, \mathcal{W}, Z, \tilde{S}, \tilde{L}, \Gamma, \alpha}{\text{minimize}} && \log \det \mathcal{W} \\ & \text{subject to} && \text{(A2) and (A3).} \end{aligned} \quad (\text{A5})$$

Since α is a variable in the design of the polytopic dynamics the resulting conditions are BMIs due to the bilinear terms $\alpha\mathcal{W}$ and $\alpha\mathcal{Y}$. One could solve the bilinear problem. Instead, a different method is proposed here. A univariate search on α is performed, where, on each iteration, an SDP is solved (which considers a fixed α), and stops when no more progress is made (with respect to a given tolerance).

Once the problem is solved, one must find M_i , S_i and L . Note that (A1) implies

$$\begin{aligned} \mathcal{W} &= (\mathcal{Y} - \mathcal{V}\mathcal{S}^{-1}\mathcal{V}^T)^{-1}, \quad \mathcal{U} = -\mathcal{Y}^{-1}\mathcal{V}(\mathcal{S} - \mathcal{V}^T\mathcal{Y}^{-1}\mathcal{V})^{-1}, \\ \mathcal{Q} &= (\mathcal{S} - \mathcal{V}^T\mathcal{Y}^{-1}\mathcal{V})^{-1} \end{aligned} \quad (\text{A6})$$

and of these the first equation implies

$$\mathcal{V}\mathcal{S}^{-1}\mathcal{V}^T = \mathcal{Y} - \mathcal{W}^{-1} = \mathcal{R}\Lambda^2\mathcal{R}^T \quad (\text{A7})$$

with $\mathcal{R}, \Lambda$ defining the eigenvector/eigenvalue matrices of $\mathcal{Y} - \mathcal{W}^{-1}$. The above implies that

$$\begin{aligned} \mathcal{V} &= \mathcal{R}\Lambda\mathcal{S}^{1/2}, \quad \mathcal{U} = -\mathcal{Y}^{-1}\mathcal{R}\Lambda(I - \Lambda\mathcal{R}^T\mathcal{Y}^{-1}\mathcal{R}\Lambda)^{-1}\mathcal{S}^{-1/2}, \\ \mathcal{Q} &= \mathcal{S}^{-1/2}(I - \Lambda\mathcal{R}^T\mathcal{Y}^{-1}\mathcal{R}\Lambda)^{-1}\mathcal{S}^{1/2} \end{aligned} \quad (\text{A8})$$

Since $\mathcal{S} \succ 0$ can be chosen freely, it is convenient to set $\mathcal{S} = rI$ with r being chosen so as to improve the conditioning of P . A sensible choice is the maximum singular value of $\mathcal{Y}$; this would avoid disparities in the volume of the projections of E_z onto x -space and v -space. Once $\mathcal{S}$ is chosen, $\mathcal{V}$ and $\mathcal{U}$ are given by (A8), and then M_i, S_i and L can be obtained from (A4).

Chapter 6

Striped affine-in-the-disturbances compensation in Stochastic MPC

Stochastic MPC (SMPC) is a family of techniques that incorporates the stochastic descriptions of the uncertainties for handling probabilistic constraints and using more realistic probabilistic measures of performance. Recent SMPC developments, however, only consider quasi closed-loop optimal control formulations, which can be conservative. A closed-loop optimal control scheme in an SMPC context is adopted in this chapter. Use is made of an affine-in-the-disturbances control policy, which like in Chapter 4, has a striped structure such that it is applied over the entire prediction horizon (extending over to infinity) thereby leading to a significant constraint relaxation. Two variations of the SMPC strategy are proposed: one where the relevant gains of the control policy are optimized offline using a sequential procedure, and another where the parameters of the control policy are optimized online. The benefits of both variations are illustrated with simulation examples.

6.1 Introduction

Robust MPC (RMPC) is a family of MPC strategies that preserve stability and performance, while guaranteeing constraint satisfaction, for all possible realizations of bounded uncertainties. However, RMPC only takes into account the bounds of the uncertainties, and not their probabilistic distribution. Thus the optimizations can be extremely conservative because the uncertainties near the bounds can be extremely unlikely. Additionally, probabilistic constraints cannot be handled effectively if no use is made of the probabilistic distribution of the uncertainty. For this reason, stochastic MPC (SMPC) has emerged as a family of techniques that incorporates the stochastic descriptions of the uncertainties for handling probabilistic constraints and using more realistic probabilistic measures of the performance. Early works in this area [57], [88], [89] considered these aspects, but the exploration of issues such as the stability and feasibility of the closed-loop implementation was conducted later, e.g. [23], [49], [21]. These works however, only consider quasi closed-loop optimal control formulations since no explicit disturbance compensation is performed. As it has been largely discussed in this thesis, this can be conservative.

This chapter considers the application of disturbance compensation in a setting of stochastic MPC (SMPC). It is shown that for a particular case the effect of the disturbances, in what concerns to constraint satisfaction, can be eliminated. This is achieved by an affine-in-the-disturbances control policy with a striped structure and an infinite number of gains. Practical implementation motivates the use of a striped affine-in-the-disturbances policy with a finite number of gains, such that the disturbance compensation extends through the infinite prediction horizon. This makes the constraints in the far horizon be more relaxed, thus improving the domain of attraction. The number of degrees of freedom of this control policy grows linearly with the prediction horizon, which is a slower rate than for other control policies [58][38][73],

thus reducing the computational burden for larger values of N .

Two strategies are proposed in this chapter. In the first the gains of the predicted control policy are optimized offline via a sequential procedure so as to improve the constraint tightening. On account of the offline design, this strategy has the same computational complexity as the strategy of [49], where the number of variables and constraints grows linearly with the prediction horizon. Additionally, the same ideas can be used to guarantee constraint satisfaction and recursive feasibility. In the second strategy all the parameters of the policy are found in the online optimization. It is shown that the striped affine-in-the-disturbances policy used here is a particular case of the control policy used in Chapter 4, and thus the ideas of Chapter 4 are used for guaranteeing constraint satisfaction. However, the implementation of these ideas is suitably modified for improved computational efficiency by exploiting the fact that the predicted policy is an affine function of the disturbances.

It is shown by means of numerical simulation that both strategies have improved control authority with respect to the strategy with a quasi closed-loop policy of [49], which is reflected in greater domains of attraction. Additionally, although one might expect the second strategy to have greater control authority on account of the online optimization of the gains, the ideas of Chapter 4 for ensuring constraint satisfaction use certain conditions that impose conservativeness. This implies that depending on the system any of the strategies could provide a better domain of attraction for a given prediction horizon.

The basic SMPC problem is defined in Section 6.2. Section 6.3 presents a particular system setting where the effect of the disturbances with respect to constraint satisfaction can be eliminated, and this is used as a motivation to introduce the striped affine-in-the-disturbances policy. The first proposed strategy is developed in Section 6.4 and the second proposed strategy is developed in Section 6.5. Section 6.7 concludes the chapter after the description of an illustrative example in Section 6.6.

6.2 System description and background

Consider a linear, time-invariant, discrete time system with additive disturbances

$$x_{t+1} = Ax_t + Bu_t + w_t, \quad (6.1)$$

where $x_t \in \mathbb{R}^n$, $w_t \in \mathbb{W} \subseteq \mathbb{R}^{n_w}$, and $u_t \in \mathbb{R}^{n_u}$ are the state, disturbance and control input at time t . The additive disturbance sequence $\{w_0, w_1, \dots\}$ is assumed to be independent and identically distributed (i.i.d.), the elements of w_t are independent and w_t has a known distribution with zero mean, variance Σ and symmetric finite support:

$$\mathbb{W} = \{w : -\bar{w} \leq w_t \leq \bar{w}\}, \quad (6.2)$$

where $\bar{w} \in \mathbb{R}^n$ and $\bar{w} \geq 0$. This symmetric description is only assumed for simplicity of exposition as the approach of this chapter can be used with any polytopic $\mathbb{W}$ (with an adequate re-definition of the operation point, if necessary, so that the disturbance has zero mean). In most practical situations finite supports are more realistic than the usual Gaussian assumption which permits the uncertainty to become arbitrarily large (albeit with small probability). System (6.1) is subject to probabilistic constraints of the type

$$\Pr\{F_j x_{t+1} + G_j u_t \leq 1\} \geq p_j, \quad j = 1, \dots, n_c, \quad t = 0, 1, \dots, \quad (6.3)$$

where F_j, G_j denote the j th rows of $F \in \mathbb{R}^{n_c \times n}$, $G \in \mathbb{R}^{n_c \times n_u}$. This formulation covers the case of state only, input only and mixed state/input constraints which can be probabilistic or deterministic (hard) since $p_j = 1$ can be chosen for some or all j .

Assumption 6.1. (i) (A, B) is stabilizable.

(ii) The state x_t is directly measurable and known at time t , but the disturbance w_t , and all its future values, are unknown at time t .

Remark 6.1. *Note that the probabilistic constraints $\Pr\{F_j x_{t+1} + G_j u_t \leq 1\} \geq p_j$ consider the input at time t and the state at time $t + 1$, i.e. u_t and x_{t+1} . Although a choice of probabilistic constraints given by $\Pr\{F_j x_t + G_j u_t \leq 1\} \geq p_j$ would have been more intuitive and standard, it cannot be handled by the current formulation and standing assumptions. The sequence of terms $F_j x_{t+k|t} + G_j u_{t+k|t}$ for $k = 1, 2, \dots$ is a stochastic process and there is no problem for handling the probabilistic constraints, however, it is not the same for $k = 0$. The term $F x_{t|t} + G u_{t|t} = F x_t + G u_t$ is not a random variable given the information that is known at time t : x_t is known at time t and $u_t = K x_t + c_{t|t}$, where $c_{t|t}$ is an optimization variable, as will be seen later. Then at time t , the constraint $F_j x_{t|t} + G_j u_{t|t}$ will be satisfied with probability $p = 0$ (it is not satisfied) or probability $p = 1$ (it is satisfied). Alternatively one might not explicitly invoke this constraint for $k = 0$, but then there would be no guarantees that the probabilistic constraints are satisfied.*

This problem has already been tackled by [49], which proposes an SMPC strategy based on the dual-mode prediction paradigm. The prediction horizon is split into the near future (Mode 1) over which the degrees of freedom of the predicted control inputs are optimized, and the far future (Mode 2) where they are given by a fixed state feedback law, $u = Kx$. Recall from Section 2.3.2.2 the strategy of [49] where the Mode 1 control policy is given by

$$u_{t+k|t} = K x_{t+k|t} + c_{t+k|t}, \quad k = 0, \dots, N - 1. \quad (6.4)$$

Then the predicted states are given by

$$x_{t|t} = x_t, \quad (6.5a)$$

$$x_{t+k+1|t} = \Phi x_{t+k|t} + B c_{t+k|t} + w_{t+k}, \quad k = 0, \dots, N - 1, \quad (6.5b)$$

$$x_{t+k+1|t} = \Phi x_{t+k|t} + w_{t+k}, \quad k \geq N, \quad (6.5c)$$

and constraints (6.3) can be expressed as

$$\begin{aligned} \Pr\{\tilde{F}_j x_{t+k|t} + \tilde{G}_j c_{t+k|t} + F_j w_{t+k} \leq 1\} &\geq p_j, \quad k = 0, \dots, N-1, j = 1, \dots, n_c, \\ \Pr\{\tilde{F}_j x_{t+k|t} + F_j w_{t+k} \leq 1\} &\geq p_j, \quad k \geq N, j = 1, \dots, n_c \end{aligned} \quad (6.6)$$

where $\Phi = A + BK$, and $\tilde{F}_j$ and $\tilde{G}_j$ denote the j th rows of $\tilde{F} = (F\Phi + GK)$ and $\tilde{G} = FB + G$, respectively.

The aim of the SMPC strategy is to minimize online, over $\mathbf{c}_t = [c_{t|t}^T, \dots, c_{t+N-1|t}^T]^T$, a predicted cost given by

$$J_t = \sum_{k=0}^{\infty} \mathbb{E}_t (x_{t+k|t}^T Q x_{t+k|t} + u_{t+k|t}^T R u_{t+k|t}), \quad (6.7)$$

where $Q \succeq 0$, $R \succ 0$. However, the limit $l_{ss} = \lim_{k \rightarrow \infty} \mathbb{E}_t (x_{t+k|t}^T Q x_{t+k|t} + u_{t+k|t}^T R u_{t+k|t})$ is non-zero due to the presence of the additive disturbance, and then J_t will be infinite. To solve this problem, the limit l_{ss} is subtracted from the stage cost so that J_t is finite:

$$J_t = \sum_{k=0}^{\infty} \mathbb{E}_t (x_{t+k|t}^T Q x_{t+k|t} + u_{t+k|t}^T R u_{t+k|t} - l_{ss}). \quad (6.8)$$

It can be proved that $l_{ss} = \text{tr}(S\Sigma)$, where S satisfies the Lyapunov equation $S - \Phi^T S \Phi = Q + K^T R K$. K is chosen to be the LQR optimal gain which, since the state is measurable, coincides with the selection that minimizes l_{ss} .

By an analysis analogous to that of [23] it can be shown that the predicted trajectories of (6.5) can be generated by the augmented system

$$z_{t|t} = \begin{bmatrix} x_t \\ \mathbf{c}_t \\ 1 \end{bmatrix}, \quad (6.9a)$$

$$z_{t+k+1|t} = \hat{\Psi}_{t+k} z_{t+k|t}, \quad k = 0, 1, \dots \quad (6.9b)$$

where

$$\hat{\Psi}_{t+k} = \begin{bmatrix} \Psi & \begin{bmatrix} w_{t+k} \\ 0 \\ 1 \end{bmatrix} \\ \begin{bmatrix} 0 & 0 \end{bmatrix} & \end{bmatrix}, \text{ in which } \Psi = \begin{bmatrix} \Phi & BE \\ 0 & M \end{bmatrix}, M = \begin{bmatrix} 0 & I & 0 & \cdots & 0 \\ 0 & 0 & I & \cdots & 0 \\ \vdots & \vdots & \ddots & \ddots & \vdots \\ 0 & \cdots & \cdots & 0 & I \\ 0 & 0 & \cdots & \cdots & 0 \end{bmatrix},$$

such that $M \in \mathbb{R}^{Nn_u \times Nn_u}$ is the matrix that shifts the elements in $\mathbf{c}_t$ according to $M\mathbf{c}_t = [c_{t+1|t}^T \cdots c_{t+N-1|t}^T 0]^T$, $E = [I \ 0 \ \cdots \ 0] \in \mathbb{R}^{n_u \times Nn_u}$ is the matrix that selects the first block component of $\mathbf{c}_t$, and I is the identity matrix in $\mathbb{R}^{n_u \times n_u}$. Then, it can be shown that the predicted cost can be expressed as the quadratic function [23]

$$J_t = \begin{bmatrix} x_t^T & \mathbf{c}_t^T & 1 \end{bmatrix} \hat{P} \begin{bmatrix} x_t^T & \mathbf{c}_t^T & 1 \end{bmatrix}^T, \quad (6.10)$$

where

$$\hat{P} - \mathbb{E}(\hat{\Psi}_{t+k}^T \hat{P} \hat{\Psi}_{t+k}) = \hat{Q} \text{ and } \hat{Q} = \begin{bmatrix} Q & 0 & 0 \\ 0 & 0 & 0 \\ 0 & 0 & -l_{ss} \end{bmatrix} + \begin{bmatrix} K^T \\ E^T \\ 0 \end{bmatrix} R \begin{bmatrix} K & E & 0 \end{bmatrix}. \quad (6.11)$$

The predicted states are reparameterized by $x_{t+k|t} = z_{t+k|t} + e_{t+k|t}$, with $e_{t|t} = 0$, where z is associated with the nominal dynamics and e is associated with the uncertain dynamics. It follows that

$$z_{t+k|t} = \Phi^k x_t + H_k \mathbf{c}_t, \text{ and } e_{t+k|t} = \Phi^{k-1} w_t + \cdots + w_{t+k-1}, \quad (6.12)$$

where $H_k = [\Phi^{k-1} B \ \cdots \ B \ 0 \ \cdots \ 0] \in \mathbb{R}^{n \times Nn_u}$ for $k = 0, \dots, N-1$ and $H_k = [\Phi^{k-1} B \ \cdots \ \Phi^{k-N} B] \in \mathbb{R}^{n \times Nn_u}$ for $k \geq N$. Thus feasibility with respect to the con-

straints (6.6) is ensured, using the arguments of [49], by

$$\tilde{F}\Phi^k x_t + (\tilde{F}H_k + \tilde{G}\hat{E}_k)\mathbf{c}_t \leq \underline{1} - \gamma_k, \quad k = 0, \dots, N-1, \quad (6.13a)$$

$$\tilde{F}\Phi^k x_t + \tilde{F}H_k \mathbf{c}_t \leq \underline{1} - \gamma_k, \quad k \geq N, \quad (6.13b)$$

where $\hat{E}_k = [0 \ 0 \ \dots 0 \ I \ 0 \ \dots 0] \in \mathbb{R}^{n_u \times n_u N}$ is a matrix of zeros except the $(k+1)$ th block component, which is given by the identity matrix in $\mathbb{R}^{n_u \times n_u}$ such that $\hat{E}_k \mathbf{c}_t = c_{t+k|t}$; and $\gamma_k \in \mathbb{R}^{n_c}$ depends on the known distribution of w , such that for $j = 1, \dots, n_c$ the j -th component, $\gamma_{k,j}$, of the vector γ_k is defined to be the minimum value such that

$$\Pr\{\tilde{F}_j \Phi^{k-1} w_t + \dots + \tilde{F}_j w_{t+k-1} + F_j w_{t+k} \leq \gamma_{k,j}\} = p_j. \quad (6.14)$$

The online use of (6.13), however, cannot guarantee recursive feasibility since it only ensures that future constraints will be satisfied with a given probability and does not account for worst cases. In particular, it is not longer the case that a candidate solution at time $t+1$ given by the *receded* $\mathbf{c}_t$ will satisfy the probabilistic constraints for all possible realizations of the disturbances. In order to have recursive feasibility, satisfaction of constraints at time t must guarantee that there will be a feasible solution at time $t+k$ for all $k \geq 0$ for any realization of the disturbances. Then, one must handle the values of $w_t, \dots, w_{t+k-1}$, by their worst case scenarios. Thus recursive feasibility and feasibility with respect to the constraints (6.6) are ensured by

$$\tilde{F}\Phi^k x_t + (\tilde{F}H_k + \tilde{G}\hat{E}_k)\mathbf{c}_t \leq \underline{1} - (\gamma_0 + \beta_k), \quad k = 0, \dots, N-1, \quad (6.15a)$$

$$\tilde{F}\Phi^k x_t + \tilde{F}H_k \mathbf{c}_t \leq \underline{1} - (\gamma_0 + \beta_k), \quad k \geq N, \quad (6.15b)$$

where β_k is given by

$$\beta_0 = 0, \quad (6.16a)$$

$$\beta_{k,j} = |\tilde{F}_j| \bar{w} + \dots + |\tilde{F}_j \Phi^{k-1}| \bar{w}, \quad k = 0, 1, \dots, \quad (6.16b)$$

such that $\beta_{k,j}$ is the j th component of β_k and $|\cdot|$ applies elementwise. Note that for all k in (6.15) the effect of w_{t+k} continues to be handled probabilistically by γ_0 . A condition for the feasibility of (6.15) is given by $\beta_\infty + \gamma \leq \underline{1}$, where $\beta_\infty = \lim_{k \rightarrow \infty} \beta_k$. If this condition is satisfied, the optimization problem will be at least feasible when $x_t = 0$ by setting $c_{t+k|t} = 0$, $k = 0, \dots, N-1$. But if it is not satisfied then there is no guarantee that the domain of attraction will not be empty.

Although (6.15) must be applied over an infinite horizon, since the sequence of β_k converges to a known limit, it is possible (see [49], [26]) to use techniques similar to those of [35] to define a terminal constraint set that ensures satisfaction of (6.15) in Mode 2.

Thus the online stage of the SMPC strategy consists of the repeated minimization at each time t , over $\mathbf{c}_t$, of the cost (6.10), subject to (6.15) over Mode 1, and to the appropriate terminal constraints at the beginning of Mode 2. The actual input to be implemented at time t is given by $u_t = Kx_t + c_t$, where c_t is the first element of the optimal $\mathbf{c}_t$. Since the cost of (6.10) is quadratic in the degrees of freedom and the constraints of (6.15) are linear, the online optimization in SMPC can be implemented as a QP.

The closed-loop application of the strategy guarantees satisfaction of constraints (6.3) and causes the mean expected value of the stage cost to be bounded from above by l_{ss} [49][26]. In addition, since K is the LQ-optimal feedback gain, the sequence $\{c_t, c_{t+1}, \dots\}$ converges to zero and the control law tends asymptotically to $u_t = Kx_t$, which steers the state to the minimal robust invariant set, $\Omega_\infty(\mathbb{W})$ [11], [72].

6.3 Closed-loop prediction strategies to compensate for the effects of disturbances

Feedback allows information available at current time to be used to improve control. Feedback can also be incorporated in the predicted control laws so as to take into account the effect of the future uncertainties on future states, that are currently unknown but whose effect will be known by the controller when the future states are measured to compute the inputs to be applied to the system. The predicted policy of (6.4), however, does not take full advantage of the feedback in the predictions. The online optimization is open-loop for the degrees of freedom in $\mathbf{c}_t$, and in spite of the presence of a feedback controller, the policy considers no special feedback treatment of future (unknown) disturbances, and can be considered to be only quasi closed-loop.

This idea was effectively exploited in [88], [58] and [38], which proposed an affine-in-the-disturbances control policy of (2.45), that is composed of a feed-forward term together with a linear mix of “known” future disturbances, KFD, where this “knowledge” will be available to the controller at any given future prediction time. Thus, whereas conventional MPC uses Nn_u degrees of freedom (in the feed-forward terms), affine-in-the-disturbances MPC deploys the extra freedom contained in a set of feedback gains $L_{k,j}$ which leads to a total number of variables and constraints in the online optimization that grows quadratically with the prediction horizon. This can cause a significant increase in computation, which may be prohibitive for fast sampling applications and systems with large dimensions or prediction horizons. In addition, no consideration was given to information on the probabilistic distribution of uncertainty. A preliminary discussion of these two issues was given in [50], and it is the aim here to present two SMPC strategies that use a type of affine-in-the-disturbances policy such that the feedback compensation for the KFD extends over Mode 2, and that the number of variables and constraints in the online optimization only grows

linearly with the prediction horizon.

[50] proposed the idea of adding an extra term to the control policy of (6.4), in order to “reduce” the effect of the KFD in constraints (6.15). Consider a control policy given by

$$u_{t+k|t} = Kx_{t+k|t} + c_{t+k|t} + v_{t+k|t}, \quad k = 0, \dots, N-1, \quad (6.17a)$$

$$u_{t+k|t} = Kx_{t+k|t} + v_{t+k|t}, \quad k \geq N, \quad (6.17b)$$

and the decomposition of the predicted states

$$e_{t|t} = 0, \quad z_{t|t} = x_t, \quad (6.18a)$$

$$z_{t+k|t} = \Phi z_{t+k|t} + Bc_{t+k|t}, \quad (6.18b)$$

$$e_{t+k|t} = \Phi e_{t+k|t} + Bv_{t+k|t} + w_{t+k} \quad (6.18c)$$

$$x_{t+k|t} = z_{t+k|t} + e_{t+k|t}, \quad (6.18d)$$

where c is used to control the nominal (disturbance-free) dynamics described by z , and v compensates for the dynamic effects of the disturbance w described by e . In particular, $v_{t+k|t}$ is to depend on the KFD sequence w_{t+j} , for $j = 0, \dots, k-1$, and is used to reduce the constraint tightening factors β_k of (6.15). To this effect consider next the predictions of variable $\psi_t = Fx_{t+1} + Gu_t$ under (6.18).

Let C_1, C_2 denote the block lower triangular Toeplitz convolution matrices formed from the elements of the impulse responses of the state space models $(\Phi, B, \tilde{F}, \tilde{G})$ and $(\Phi, I, \tilde{F}, 0)$:

$$C_1 = \begin{bmatrix} \tilde{G} & 0 & 0 & \cdots \\ \tilde{F}B & \tilde{G} & 0 & \cdots \\ \tilde{F}\Phi B & \tilde{F}B & \tilde{G} & \cdots \\ \vdots & & & \ddots \end{bmatrix}, \quad C_2 = \begin{bmatrix} 0 & 0 & 0 & \cdots \\ \tilde{F} & 0 & 0 & \cdots \\ \tilde{F}\Phi & \tilde{F} & 0 & \cdots \\ \vdots & & & \ddots \end{bmatrix}, \quad (6.19)$$

Also let $C_{1,N}$ denote the matrix containing the first N block columns of C_1 and let $D = [\tilde{G}^T \ (\tilde{G}\Phi)^T \ (\tilde{G}\Phi^2)^T \ \dots]^T$. Then by (6.3) and (6.18) it follows that $\psi_{t+k|t} = (\tilde{F}z_{t+k|t} + \tilde{G}c_{t+k|t}) + (\tilde{F}e_{t+k|t} + \tilde{G}v_{t+k|t}) + Fw_{t+k|t}$ and

$$\vec{\psi}_t = \vec{\zeta}_t + \vec{\varepsilon}_t + \vec{\eta}_t, \quad \vec{\zeta}_t = Dx_t + C_{1,N}\mathbf{c}_t, \quad \vec{\varepsilon}_t = C_1\vec{v}_t + C_2\vec{w}_t, \quad \vec{\eta}_t = \text{diag}\{F\}\vec{w}_t, \quad (6.20)$$

where $(\vec{\cdot})_t$ denotes the vector formed by the predictions of $(\cdot)_{t+k|t}$ over an infinite horizon $[(\cdot)_{t|t}^T \ (\cdot)_{t+1|t}^T \ \dots]^T$. Thus predictions have been separated into: (i) $\vec{\zeta}_t$, the nominal (disturbance-free) part of $\vec{\psi}_t$, to be controlled by c ; (ii) $\vec{\varepsilon}_t$, the stochastic part of $\vec{\psi}_k$ containing the KFD to be controlled by v ; (iii) $\vec{\eta}_t$, the stochastic part of $\vec{\psi}_t$ about which nothing can be done since it comprises the effect of future disturbances which are as yet unknown to the controller; these will be referred to as “current” future disturbances, CFD, in contrast to KFD.

First, a special case is studied, in which it is possible to choose v so as to make the tightening parameters $\beta = 0$, thus providing increased control authority to the degrees of freedom in $\mathbf{c}_t$ allowing one to obtain increased domains of attraction. This special case is then used as motivation for a control policy that will be used in two strategies to be presented in Section 6.4 and Section 6.5.

Theorem 6.2. *Let $C_1^\dagger$ denote a right inverse of C_1 and let $C_1^\perp$ denote a matrix whose columns form a basis for the kernel of C_1 . Then, for $n_c \leq n_u$, $\tilde{G}$ full-rank, and $(\Phi, B, \tilde{F}, \tilde{G})$ minimum-phase, the control policy of $u_{t+k|t} = Kx_{t+k|t} + c_{t+k|t} + v_{t+k|t}$, with*

$$\vec{v}_t = -C_1^\dagger C_2 \vec{w}_t + C_1^\perp \theta \quad (6.21)$$

(where θ is an arbitrary vector of appropriate dimensions) makes $\psi_{t+k|t}$ independent of the KFD sequence $\{w_{t|t}, \dots, w_{t+k-1|t}\}$ at each prediction time $k = 1, 2, \dots$

Proof. For $n_c \leq n_u$, the minimum-phase assumption of the theorem implies that the inverse of the dynamic map represented by $(\Phi, B, \tilde{F}, \tilde{G})$ is stable, and therefore that

C_1 is right invertible (with the elements of the right inverse, $C_1^\dagger$, tending to zero with prediction time). The expression for $\vec{v}_t$ in (6.21) achieves the cancellation of the KFD by setting $\vec{\varepsilon}_k = 0$. $\square$

The block lower triangular structure of C_1 implies that the right inverse $C_1^\dagger$ is also block lower triangular, and since C_2 has a strict block lower triangular structure, $\vec{v}_t$ depends only on the KFD and $v_{t|t} = 0$. Since $\vec{\varepsilon}_k = 0$, the probabilistic constraints take the form of (6.15) but with all β_k values set to zero:

$$Dx_t + C_{1,N}\mathbf{c}_t \leq \underline{1} - \gamma, \quad \Pr\{\tilde{G}_j w \leq \gamma_j\} = p_j, \quad j = 1, \dots, n_c, \quad (6.22)$$

where $\gamma = [\gamma_1, \dots, \gamma_{n_c}]^T := \gamma_0$.

The fact that the tightening parameters $\beta_k = 0$ means that the degrees of freedom available in $\mathbf{c}_t$ will have a greater control authority in (6.22) than in (6.15), which enables one to obtain enlarged domains of attraction.

Note that when $n_c = n_u$, (6.21) becomes

$$\vec{v}_t = -C_1^{-1}C_2\vec{w}_t = L\vec{w}_t, \quad (6.23)$$

where L is block lower triangular Toeplitz matrix, and hence $v_{t+k|t}$ depend linearly on the KFD:

$$v_{t+k|t} = L_k w_t + L_{k-1} w_{t+1} + \dots + L_1 w_{t+k-1|k}, \quad v_{t|t} = 0, \quad (6.24)$$

thus resulting in an affine-in-the-disturbances policy, but unlike the policy of (2.45) proposed in [58], [38], the matrix L has a striped lower triangular structure and the disturbance compensation extends into Mode 2. The minimum-phase assumption of Theorem 6.2 implies that the sequences formed by the elements of every column of $C_1^\dagger$ converge to zero, and therefore the values of L_k converge to zero. Hence the number

of gains L_k can be truncated so as to lead to an implementable receding horizon MPC strategy.

All this motivates the use of an affine-in-the-disturbances policy with a striped structure, and for practical reasons, use will be made only of the first $N-1$ parameters L_k . According to this, use will be made of the control policy given by

$$u_{t+k|t} = Kx_{t+k|t} + c_{t+k|t} + L_k w_t + L_{k-1} w_{t+1} + \dots + L_1 w_{t+k-1}, \quad k = 0, \dots, N-1, \quad (6.25a)$$

$$u_{t+k|t} = Kx_{t+k|t} + L_{N-1} w_{t+k-(N-1)} + \dots + L_1 w_{t+k-1}, \quad k \geq N, \quad (6.25b)$$

Note that this policy is a particular case of the policy of (4.15). For further details please see Section 6.3.1.

In the case without disturbance compensation (i.e. $L_k = 0$) the predictions of (6.20) satisfy constraints (6.6) if (6.15) holds [49], which can be re-written as

$$\hat{E}_k D x_t + \hat{E}_k C_{1,N} \mathbf{c}_t \leq \underline{1} - (\gamma + \beta_k), \quad k = 0, 1, \dots \quad (6.26)$$

where $\hat{E}_k$ is the operator that selects the $(k+1)$ th block row of D and $C_{1,N}$, so that $\hat{E}_k \vec{\zeta}_t = \zeta_{t+k|t}$. In (6.26) γ accounts for the CFD probabilistically, whereas β_k accounts for the KFD $(w_t, \dots, w_{t+k-1})$ and renders the constraint tightening robust by considering the worst-case effect of each of the KFD to $\hat{E}_k C_2 \vec{w}_t$.

This treatment of constraints clearly also applies to the case when $L_k \neq 0$. Since the tightening coefficients account for the KFD in $\vec{\varepsilon}_t$, $\tilde{\beta}_k$ must be such that

$$\vec{\tilde{\beta}} = \max_{\vec{w}_t} \vec{\varepsilon}_k = \max_{\vec{w}_t} (C_1 L + C_2) \vec{w}_k, \quad (6.27)$$

$$L = \begin{bmatrix} 0 & \cdots & & & & \\ L_1 & 0 & \cdots & & & \\ L_2 & L_1 & 0 & \cdots & & \\ \vdots & \vdots & \ddots & \ddots & & \\ L_{N-1} & L_{N-2} & \cdots & L_1 & 0 & \cdots \\ 0 & L_{N-1} & L_{N-2} & \cdots & L_1 & 0 & \cdots \\ \vdots & \vdots & \ddots & \cdots & \ddots & \ddots & \cdots \end{bmatrix}, \quad \vec{\beta} = \begin{bmatrix} 0 \\ \tilde{\beta}_1 \\ \tilde{\beta}_2 \\ \tilde{\beta}_3 \\ \vdots \end{bmatrix}, \quad (6.28)$$

where the maximization is performed element-wise, and accordingly, satisfaction of constraints (6.15) is guaranteed by

$$\hat{E}_k D x_t + \hat{E}_k C_{1,N} \mathbf{c}_t \leq \underline{1} - (\gamma + \tilde{\beta}_k). \quad k = 0, 1, \dots \quad (6.29)$$

Two strategies will be considered next. In the first, $\mathbf{L} = \{L_1, \dots, L_{N-1}\}$ and $\tilde{\beta}_k$ will be taken to be time invariant and will be designed offline so as bring about a relaxation to the online constraints of (6.29). In the second strategy, $\mathbf{L}$ and $\tilde{\beta}_k$ are variables in the online optimization.

It is useful to note that from the structures of C_1 , C_2 and L it follows that $\varepsilon_{t+k|t}$ can be rewritten as

$$\varepsilon_{t+k|t} = (\tilde{F}\Phi^{k-1} + \hat{C}_k \mathcal{L})w_t + \dots + (\tilde{F} + \hat{C}_1 \mathcal{L})w_{t+k-1}, \quad k = 0, 1, \dots \quad (6.30)$$

where $\mathcal{L} = [0 \ L_1^T \cdots L_{N-1}^T]^T$ and

$$\hat{C}_k = \begin{cases} \hat{E}_k \tilde{G} & k = 1 \\ \hat{E}_k \tilde{G} + \tilde{F}[\Phi^{k-2} B \ \cdots \ B \ 0 \cdots], & k = 2, \dots, N-1 \\ \tilde{F}[\Phi^{k-2} B \ \cdots \ B \ \Phi^{k-N} B], & k \geq N, \end{cases} \quad (6.31)$$

and then $\tilde{\beta}_k$ will be given by

$$\tilde{\beta}_k = \max_{w \in \mathbb{W}} (\tilde{F}\Phi^{k-1} + \hat{C}_k \mathcal{L})w + \dots + \max_{w \in \mathbb{W}} (\tilde{F} + \hat{C}_1 \mathcal{L})w, \quad k = 1, 2, \dots \quad (6.32a)$$

From this, the $\tilde{\beta}_k$ satisfy the recursion

$$\tilde{\beta}_0 = 0, \quad (6.33a)$$

$$\tilde{\beta}_{k+1} = \tilde{\beta}_k + \max_{w \in \mathbb{W}} (\tilde{F}\Phi^k + \hat{C}_{k+1} \mathcal{L})w, \quad k = 0, 1, \dots, \quad (6.33b)$$

where $\max_{w \in \mathbb{W}} (\tilde{F}\Phi^{k-1} + \hat{C}_k \mathcal{L})w = |\tilde{F}\Phi^{k-1} + \hat{C}_k \mathcal{L}|\bar{w}$, where $|\cdot|$ applies element-wise.

6.3.1 Connection with the control policy of Chapter 4

In the notation of Chapter 4, by setting

$$u_{(0,k)} = Kx_{(0,k)} + c_{t+k|t}, \quad k = 0, \dots, N-1, \quad (6.34a)$$

$$u_{(0,k)} = Kx_{(0,k)}, \quad k \geq N, \quad (6.34b)$$

$$u_{(i,1,k)} = Kx_{(i,1,k)} + L_k \tilde{w}_i, \quad k = 1, \dots, N-1, \quad (6.34c)$$

$$u_{(i,1,k)} = Kx_{(i,1,k)}, \quad k \geq N, \quad (6.34d)$$

where $\tilde{w}_i$, $i = 1, \dots, q$ are the vertices of $\mathbb{W}$, the control policy of (4.15) becomes

$$\mu_0(x_{(0)}) = Kx_{(0,0)} + c_{t|t}, \quad (6.35a)$$

$$\begin{aligned} \mu_k(x_{(k)}) &= Kx_{(0,k)} + c_{t+k|t} + \sum_{j=1}^k \sum_{i=1}^q \lambda_{(i,j,j)}(w_{(j-1)}) (Kx_{(i,1,k-j+1)} + L_{(i,1,k-j+1)} \tilde{w}_i), \\ &= Kx_{(0,k)} + \sum_{j=1}^k Kx_{(1,k-j+1)} + c_{t+k|t} + \sum_{j=1}^k L_{(i,1,k-j+1)} w_{(j-1)}, \\ &= Kx_{(k)} + c_{t+k|t} + L_k w_{(0)} + \dots + L_1 w_{(k-1)}, \quad k = 1, \dots, N-1, \end{aligned} \quad (6.35b)$$

and

$$\begin{aligned}
\mu_k(x_k) &= Kx_{(0,k)} + \sum_{j=1}^{k-N+1} \sum_{i=1}^q \lambda_{(i,j,j)}(w_{(j-1)}) Kx_{(i,1,k-j+1)} \\
&\quad + \sum_{j=k-N+2}^k \sum_{i=1}^q \lambda_{(i,j,j)}(w_{(j-1)}) (Kx_{(i,1,k-j+1)} + L_{(i,1,k-j+1)} \tilde{w}_i), \\
&= Kx_{(0,k)} + \sum_{j=1}^k Kx_{(1,k-j+1)} + c_{t+k|t} + \sum_{j=k-N+2}^k L_{(i,1,k-j+1)} w_{(j-1)}, \\
&= Kx_{(k)} + L_{N-1}w_{(0)} + \cdots + L_1w_{(k-1)}, \quad k \geq N,
\end{aligned} \tag{6.35c}$$

which is the same policy as that of (6.25). This proves that the control policy of (6.25) is a particular case of the policy of Chapter 4 given by (4.15).

6.4 SMPC with offline design of the disturbance compensation gains

The first SMPC strategy proposed in this chapter assumes $\mathbf{L}$ and $\tilde{\beta}_k$ to be time invariant such that their elements are designed offline, with a sequential procedure, so as to bring about a relaxation to the online constraints of (6.29).

6.4.1 Offline design of the disturbance compensation gains

It can be seen from (6.29) that the $\tilde{\beta}_k$ effect a constraint tightening on the nominal dynamics that allows one to guarantee constraint satisfaction in the presence of uncertainty. Less tight constraints imply larger domains of attraction, and thus one should seek to design $\mathbf{L}$ to relax the constraints.

The striped block lower triangular structure of L and the recursive property of (6.33) of $\tilde{\beta}_k$ imply that $\tilde{\beta}_1$ depends only on L_1 , $\tilde{\beta}_2$ depends on L_1 and L_2 , etc., which motivates a sequential design (first $\tilde{\beta}_1$, then $\tilde{\beta}_2$ etc.) analogous to that of Algorithm 3.5. This is computationally convenient, but it also makes sense in that it is consistent with the observation that constraints are more likely to be active during the early transients and less so as the predictions reach their steady state.

As a design condition, $\tilde{\beta}_k$ must satisfy

$$\tilde{\beta}_k \leq \underline{1} - \gamma, \quad k = 1, 2, \dots \quad (6.36)$$

Like (3.24) in Section 3.3, this ensures that the domain of attraction is not empty, since under this condition (6.29) is guaranteed to be feasible at least for $x_t = 0$ (and without it there is no guarantee that the domain of attraction will not be empty). Additionally, in order to ensure that the gains L_k are effectively used to relax the tightening of (6.29), it is desirable that the new tightening parameters are smaller than those obtained without compensation,

$$\tilde{\beta}_k \leq \beta_k, \quad k = 1, 2, \dots \quad (6.37)$$

Conditions (6.36) and (6.37) appear to imply an infinite number of constraints on L_k , but as is the case for the coefficients β_k in (6.16), for the disturbance compensation case considered here the sequence of $\tilde{\beta}_k$ is monotonically increasing and tends to a limit $\tilde{\beta}_\infty$ (as stated in Lemma 6.5 below). Then, a supplementary horizon N_1 is used such that (6.36) and (6.37) are imposed explicitly until $N + N_1$ in the offline design. Note that since $\mathbb{W}$ can be expressed as the set $\{w : Sw \leq h\}$, $S \in \mathbb{R}^{n_s \times n}$, as shown in [38]

$$\begin{aligned} \max_{w \in \mathbb{W}} (\hat{F}_i \Phi^{k-1} + \hat{C}_{k,i} \mathcal{L})w &= \min_{\mathbf{r}_{k,i}} h^T \mathbf{r}_{k,i} \\ \text{subject to } S^T \mathbf{r}_{k,i} &= (\hat{F}_i \Phi^{k-1} + \hat{C}_{k,i} \mathcal{L})^T, \quad \mathbf{r}_{k,i} \geq 0, \end{aligned} \quad (6.38)$$

where $\hat{C}_{k,i}$ denotes the i th row of $\hat{C}_k$, and $\mathbf{r}_{k,i} \in \mathbb{R}^{n_s}$ represents the dual variable associated with the maximization of $(\hat{F}_i \Phi^{k-1} + \hat{C}_{k,i} \mathcal{L})w$. Then conditions (6.36) and

(6.37), for $k = 1, \dots, N + N_1$, can be invoked as

$$\mathbf{R}_1^T h + \dots + \mathbf{R}_k^T h \leq \underline{1} - \gamma, \quad k = 1, \dots, N + N_1, \quad (6.39a)$$

$$\mathbf{R}_1^T h + \dots + \mathbf{R}_k^T h \leq \beta_k, \quad k = 1, \dots, N + N_1, \quad (6.39b)$$

$$(\tilde{F}\Phi^{k-1} + \hat{C}_k \mathcal{L}) = \mathbf{R}_k^T S, \quad k = 1, \dots, N + N_1, \quad (6.39c)$$

$$\mathbf{R}_k \geq 0. \quad (6.39d)$$

where $\mathbf{R}_k \in \mathbb{R}^{n_s \times n_c}$ is such that $\mathbf{R}_k^T h \geq \max_{w \in \mathbb{W}} (\tilde{F}\Phi^{k-1} + \hat{C}_k \mathcal{L})w$.

As will be seen in Lemma 6.5, condition (6.39) guarantees satisfaction of (6.37) for $k = 1, \dots, N + N_1$, and if that N_1 is sufficiently large it also guarantees satisfaction of (6.36).

Based on the considerations expressed above, the corresponding sequential optimization of $L_1, \dots, L_{N-1}$ is presented in Algorithm 6.3.

Algorithm 6.3.

1. First, solve the LP problem

$$\begin{aligned} & \underset{L_1, \dots, L_{N-1}, \mathbf{R}_1, \dots, \mathbf{R}_{N+N_1}, \mu}{\text{minimize}} && \mu \\ & \text{subject to} && \mathbf{R}_1^T h \leq \mu \underline{1}, \text{ and (6.39)}. \end{aligned} \quad (6.40)$$

Set $f_1 = \mathbf{R}_1^{*T} h$, where the super-script $*$ represents the optimal solution of the optimization.

2. For $s = 2, \dots, N - 1$, solve the LP problems

$$\begin{aligned} & \underset{L_s, \dots, L_{N-1}, \mathbf{R}_1, \dots, \mathbf{R}_{N+N_1}, \mu}{\text{minimize}} && \mu \\ & \text{subject to} && \mathbf{R}_1^T h + \dots + \mathbf{R}_k^T h \leq \mu \underline{1}, \text{ and (6.39),} \\ & && \mathbf{R}_k^T h \leq f_k, \quad k = 1, \dots, s - 1, \end{aligned} \quad (6.41)$$

Set $f_s = \mathbf{R}_s^{*T} h$.

3. Finally, the desired $\tilde{\beta}_k$ are obtained from the solution of (6.41) by $\tilde{\beta}_k = |\tilde{F} + \hat{C}_1 \mathcal{L}^* | \bar{w} + \dots + |\tilde{F} \Phi^{k-1} + \hat{C}_k \mathcal{L}^* | \bar{w}$.

Remark 6.4. The constraints $\mathbf{R}_k^T h \leq f_k$, $k = 1, \dots, s-1$ in (6.41) guarantee that once the value of $\tilde{\beta}_k$ is optimized, it is not made worse in future iterations, which is a condition needed to ensure that the sequential optimization makes sense (as discussed in 3.6). The use of this constraint is preferred to fixing the value of L_k once it has been optimized at iteration $s = k$, which would fulfil the same role. This is because the optimization in (6.41) is an LP and then several values of L_k can achieve the same minimum, so fixing the value of L_k would impose an unnecessary degree of conservativeness.

Lemma 6.5. Under the assumption that $\gamma + \beta_\infty < \underline{1}$, there exist N, N_1 such that the minimizations of Algorithm 6.3 are feasible, and the resulting sequence of constraint tightening parameters $\tilde{\beta}_k$ increases monotonically to a finite upper bound, $\tilde{\beta}_\infty$, and satisfies $\tilde{\beta}_k \leq \beta_k$ for $k = 1, \dots, N + N_1$ and $\tilde{\beta}_\infty \leq \underline{1}$.

Proof. $\gamma + \beta_\infty < \underline{1}$ ensures the existence of valid constraint tightening parameters for the case of no disturbance compensation. Therefore, $L_k = 0$ defines a feasible solution for Algorithm 6.3, which guarantees that $\tilde{\beta}_k \leq \beta_k$, $k = 1, \dots, N + N_1$.

From (6.33) it follows that $\tilde{\beta}_{k+1} = \tilde{\beta}_k + \delta_k$ where $\delta_k = |\tilde{F} \Phi^{k-1} + \hat{C}_k \mathcal{L}^* | \bar{w}$. Finally, for $\Xi = G \Phi^{N_1-N}$, $\Gamma = [\Phi^{N-2} B \quad \Phi^{N-3} B \quad \dots \quad B]$, detailed algebra shows that for $k = N_1, N_1 + 1, \dots$ the sequence of δ_k is $\{|\Xi \Gamma | \bar{w}, |\Xi \Phi \Gamma | \bar{w}, |\Xi \Phi^2 \Gamma | \bar{w} \dots\}$. Hence the infinite sum, Δ_{N+N_1} , of the δ_k can be computed using the finite l_1 norm of the impulse response of (Φ, Γ, Ξ) and $\tilde{\beta}_\infty$ can be computed as $\tilde{\beta}_\infty = \tilde{\beta}_{N+N_1} + \Delta_{N+N_1}$. Clearly Ξ and therefore Δ_{N+N_1} can be made to be as small as desired by increasing the value of N_1 , and since $\gamma + \beta_{N+N_1} < \underline{1}$ and $\tilde{\beta}_{N+N_1} \leq \beta_{N+N_1}$, it follows that there

exists a sufficiently large N_1 such that Δ_{N+N_1} is small enough for all $N + N_1 \geq \bar{N}$ so that $\beta_\infty + \delta_{N+N_1} + \gamma \leq \underline{1}$. $\square$

Remark 6.6. *Since it is assumed that $\gamma + \beta_\infty < \underline{1}$, condition (6.37) is enough to guarantee satisfaction of (6.36) provided that N_1 is sufficiently large, and then there is no need to use (6.39a) in the sequential optimization of Algorithm 6.3.*

Remark 6.7. *If the disturbances are sufficiently large such that $\gamma + \beta_\infty > \underline{1}$ then SMPC with no compensation becomes infeasible: there will not exist a valid set of β_k . However it is still possible that SMPC with disturbance compensation is feasible and the relevant $\tilde{\beta}_k$ parameters can be defined as per Algorithm 6.3 but without the constraint $\tilde{\beta}_k \leq \vec{\beta}$, imposed by (6.39b). Instead, in this case, one must use $\tilde{\beta}_k + \gamma \leq \underline{1}$, which can be imposed by (6.39a)¹.*

6.4.2 Constraints

Although (6.29) must be applied over an infinite horizon, given that the sequence of $\tilde{\beta}_k$ converges to $\tilde{\beta}_\infty$, it is possible to use techniques similar to those of [35] to define a terminal constraint set for $z_{t+N|t}$ that ensures satisfaction of (6.29) in Mode 2 [49].

Assumption 6.2. (i) $\tilde{F}$ is such that $\tilde{F}z \leq \underline{1}$ guarantees that $\tilde{F}z$ remains bounded.

(ii) The pair $(\Phi, \tilde{F})$ is observable.

(iii) $\gamma + \tilde{\beta}_\infty < \underline{1}$.

Remark 6.8. *Some conditions that guarantee that Assumption 6.2(i) holds are that $\tilde{F} = [F_1^T \quad -F_1^T]^T$ or that the set $\{z : \tilde{F}z \leq \underline{1}\}$ is bounded.*

To derive the terminal constraint, note that in Mode 2, (6.29) can be written as

$$\tilde{F}\Phi^k z_{t+N+k|t} \leq \underline{1} - (\gamma + \tilde{\beta}_{N+k}), \quad k = 0, 1, \dots, \quad (6.42)$$

¹Note however that there would not then be an a priori guarantee that SMPC with disturbance compensation will be feasible, and hence Algorithm 6.3 could become infeasible for large values of N and N_1 . Further research would be necessary to address the problem in depth.

and so the maximal admissible set $\mathbb{Z}_\infty$ is defined as $\mathbb{Z}_\infty = \{z_{t+N+k|t} : (6.42) \text{ holds}\}$.

Since $\mathbb{Z}_\infty$ is defined by an infinite number of inequalities, consider instead

$$\mathbb{Z}_k = \{z : \tilde{F}\Phi^i z \leq \underline{1} - (\gamma + \tilde{\beta}_{N+i}), \quad i = 0, \dots, k\}, \quad (6.43)$$

and let $\hat{N}$ be the minimal value that satisfies $\max_{z \in \mathbb{Z}_{\hat{N}}} \tilde{F}_j \Phi^{\hat{N}+1} z \leq 1 - \tilde{\beta}_{\hat{N}+1,j}$ for $j = 1, \dots, n_c$. The existence of a finite $\hat{N}$ that satisfies this condition is guaranteed, and $\mathbb{Z}_{\hat{N}} = \mathbb{Z}_\infty$.

Proposition 6.9. *If $\mathbb{Z}_{\hat{N}} = \mathbb{Z}_{\hat{N}+1}$, then:*

(i) *If $z \in \mathbb{Z}_{\hat{N}}$, then $\Phi z + (\Phi^N + \Phi^{N-1}BL_1 + \dots + \Phi BL_{N-1})w \in \mathbb{Z}_{\hat{N}}$ for all $w \in \mathbb{W}$, and*

(ii) $\mathbb{Z}_{\hat{N}} = \mathbb{Z}_{\hat{N}+1} = \dots = \mathbb{Z}_\infty$.

Proof. (i) If $z \in \mathbb{Z}_{\hat{N}} = \mathbb{Z}_{\hat{N}+1}$ then $\tilde{F}\Phi^k z \leq \underline{1} - (\gamma + \tilde{\beta}_{N+k})$ for $k = 0, \dots, \hat{N} + 1$. From (6.33) and (6.31), this is equivalent to $\tilde{F}\Phi^{k-1}(\Phi z + (\Phi^N + \Phi^{N-1}BL_1 + \dots + \Phi BL_{N-1})w) \leq \underline{1} - (\gamma + \tilde{\beta}_{N+k-1})$ for $k = 0, \dots, \hat{N} + 1$ and all $w \in \mathbb{W}$, which in turn is equivalent to $\tilde{F}\Phi^k(\Phi z + (\Phi^N + \Phi^{N-1}BL_1 + \dots + \Phi BL_{N-1})w) \leq \underline{1} - (\gamma + \tilde{\beta}_{N+k})$ for $k = 0, \dots, \hat{N}$ and all $w \in \mathbb{W}$. Thus $\Phi z + (\Phi^N + \Phi^{N-1}BL_1 + \dots + \Phi BL_{N-1})w \in \mathbb{Z}_{\hat{N}}$ for all $w \in \mathbb{W}$.

(ii) From (i) it follows that if $z \in \mathbb{Z}_{\hat{N}} = \mathbb{Z}_{\hat{N}+1}$, then $\tilde{F}\Phi^k(\Phi z + (\Phi^N + \Phi^{N-1}BL_1 + \dots + \Phi BL_{N-1})w) \leq \underline{1} - (\gamma + \tilde{\beta}_k)$ for $k = 0, \dots, \hat{N} + 1$. The case $k = \hat{N} + 1$, using (6.33), yields $\tilde{F}\Phi^{\hat{N}+2}z + \tilde{\beta}_{\hat{N}+1} + (\Phi^{N+\hat{N}+1} + \Phi^{N+\hat{N}}BL_1 + \dots + \Phi^{\hat{N}+2}BL_{N-1})w \leq \underline{1} - \gamma$, which implies that $\tilde{F}\Phi^{\hat{N}+2}z + \leq \underline{1} - (\gamma + \tilde{\beta}_{\hat{N}+2})$, thus $z \in \mathbb{Z}_{\hat{N}+2}$. However it is clear from (6.43) that $\mathbb{Z}_{k+1} \subseteq \mathbb{Z}_k$ for all $k \geq 0$, thus $\mathbb{Z}_{\hat{N}} = \mathbb{Z}_{\hat{N}+1} = \mathbb{Z}_{\hat{N}+2}$. The result follows by applying the same reasoning inductively. $\square$

Theorem 6.10. *There exists an $\hat{N} \geq n-1$ such that if $z \in \mathbb{Z}_{\hat{N}}$, then $\max_{z \in \mathbb{Z}_{\hat{N}}} \tilde{F}_j \Phi^{\hat{N}+1} z \leq 1 - \tilde{\beta}_{\hat{N}+1,j}$ for $j = 1, \dots, n_c$. $\mathbb{Z}_{\hat{N}} = \mathbb{Z}_{\hat{N}+1} = \dots = \mathbb{Z}_\infty$ for this $\hat{N}$.*

Proof. Since the pair $(\Phi, \tilde{F})$ is observable and $\tilde{F}\Phi^i z$, $i = 1, \dots, k$ is bounded for all $z \in \mathbb{Z}_k$, $\mathbb{Z}_k$ is bounded for all $k \geq n - 1$. This and the fact that Φ is strictly stable imply that there exists an $\hat{N} \geq n - 1$ such that $\tilde{F}\Phi^{\hat{N}+1} z \leq \underline{1} - (\gamma + \tilde{\beta}_{\hat{N}+1})$ for all $z \in \mathbb{Z}_{\hat{N}}$, thus $\max_{z \in \mathbb{Z}_{\hat{N}}} \tilde{F}_j \Phi^{\hat{N}+1} z \leq 1 - \tilde{\beta}_{\hat{N}+1,j}$ for $j = 1, \dots, n_c$. This also implies that $z \in \mathbb{Z}_{\hat{N}+1}$. However it is clear from (6.43) that $\mathbb{Z}_{k+1} \subseteq \mathbb{Z}_k$ for all $k \geq 0$, thus $\mathbb{Z}_{\hat{N}} = \mathbb{Z}_{\hat{N}+1}$. That $\mathbb{Z}_{\hat{N}} = \mathbb{Z}_{\hat{N}+1} = \dots = \mathbb{Z}_{\infty}$ follows directly from Proposition 6.9. $\square$

Corollary 6.11. *If $z_{t+N|t} \in \mathbb{Z}_{\hat{N}}$, then (6.29) is satisfied for all $k \geq N$.*

Proof. If $z_{t+N|t} \in \mathbb{Z}_{\hat{N}}$, then $z_{t+N|t} \in \mathbb{Z}_{\infty}$ which implies $\tilde{F}\Phi^k z_{t+N+k|t} \leq \underline{1} - (\gamma + \tilde{\beta}_{N+k})$ for all $k \geq 0$. $\square$

As a result, satisfaction of constraints over Mode 1 and Mode 2 can be ensured by invoking (6.29) for $k = 0, \dots, N - 1$ and the terminal constraint $z_{t+N|t} \in \mathbb{Z}_{\hat{N}}$

$$\hat{E}_k D x_t + \hat{E}_k C_{1,N} \mathbf{c}_t \leq \underline{1} - (\gamma + \tilde{\beta}_k), \quad k = 0, \dots, N - 1, \quad (6.44a)$$

$$V_f(\Phi^N x_t + H_N \mathbf{c}_t) \leq \underline{1}. \quad (6.44b)$$

where $V_f \in \mathbb{R}^{n_f \times n_c}$ is such that the terminal set can be expressed as $\mathbb{Z}_{\hat{N}} = \{z : V_f z \leq \underline{1}\}$.

6.4.3 Cost function

Since the disturbance compensation extends into Mode 2, the terminal control law is not $u = Kx$. This implies that the steady state value of the expectation of the stage cost, L_{ss} , is not necessarily equal to l_{ss} , and then the cost of (6.8) is unbounded under the control policy considered here. Thus, instead of the predicted cost of (6.8), used is made of the predicted cost

$$J_t = \sum_{k=0}^{\infty} \mathbb{E}(x_{t+k|t}^T Q x_{t+k|t} + u_{t+k|t}^T R u_{t+k|t} - L_{ss}). \quad (6.45)$$

This predicted cost can be expressed as a quadratic function of the initial state and $\mathbf{c}_t$.

Theorem 6.12. *The predicted cost of (6.45) for system (6.1) under the control policy of (6.25) is quadratic in x_t and $\mathbf{c}_t$ and has the form:*

$$J_t = x_t^T P_x x_t + \mathbf{c}_t^T P_c \mathbf{c}_t + P_1 \quad (6.46)$$

where $P = \text{diag}(P_x, P_c, 1)$ satisfies

$$P - \mathbb{E}(\tilde{\Psi}^T P \tilde{\Psi}) = \tilde{Q}, \quad (6.47)$$

in which

$$\tilde{\Psi} = \begin{bmatrix} \begin{bmatrix} \Phi & BE \\ 0 & M \\ 0 & 0 \end{bmatrix} & \begin{bmatrix} w \\ \mathcal{L}w \\ 1 \end{bmatrix} \end{bmatrix}, \quad \tilde{Q} = \begin{bmatrix} Q + K^T R K & K^T R E & 0 \\ E^T R K & E^T R E & 0 \\ 0 & 0 & -L_{ss} \end{bmatrix}. \quad (6.48)$$

Furthermore, the steady state expectation, L_{ss} , of the predicted stage cost $l_t = x_{t+k|t}^T Q x_{t+k|t} + u_{t+k|t}^T R u_{t+k|t}$, is given by

$$L_{ss} = l_{ss} + \text{tr}(\mathcal{L}^T P_c \mathcal{L} \Sigma). \quad (6.49)$$

Proof. For this strategy, the prediction dynamics can be generated more compactly by

$$\chi_{t|t} = \begin{bmatrix} x_t \\ \mathbf{c}_t \\ 1 \end{bmatrix}, \quad (6.50)$$

$$\chi_{t+k+1|t} = \tilde{\Psi}_{t+k} \chi_{t+k|t},$$

where

$$\tilde{\Psi}_{t+k} = \begin{bmatrix} \begin{bmatrix} \Phi & BE \\ 0 & M \end{bmatrix} & \begin{bmatrix} w_{t+k} \\ \mathcal{L}w_{t+k} \end{bmatrix} \\ \begin{bmatrix} 0 & 0 \end{bmatrix} & 1 \end{bmatrix} \quad (6.51)$$

For $V_{t+k|t} = \chi_{t+k|t}^T P \chi_{t+k|t}$, the above together with (6.47) imply that

$$V_{t+k|t} - \mathbb{E}_{t+k}(V_{t+k+1|t}) = l_{t+k|t} - L_{ss}. \quad (6.52)$$

Applying $\mathbb{E}_t$ to (6.52) and summing over $k = 0, 1, \dots$ gives $V_t = \lim_{k \rightarrow \infty} \mathbb{E}_t(V_{t+k+1|t}) + J_t$. However the Lyapunov condition of (6.47) does not define P_1 (as can be seen from (6.56) below), which therefore can be chosen so that $\lim_{k \rightarrow \infty} V_{t+k|t} = 0$ and therefore $V_t = J_t$. If the partition of P , conformal to the partition of the vector $[x_t^T \quad \mathbf{c}_t^T \quad 1]$, is given as

$$P = \begin{bmatrix} P_x & P_{xc} & P_{x1} \\ P_{xc}^T & P_c & P_{c1} \\ P_{x1}^T & P_{c1}^T & P_1 \end{bmatrix}, \quad (6.53)$$

then the expansion of (6.47) in terms of the partitions of P gives

$$P_x - \Phi^T P_x \Phi = Q + K^T R K \quad (6.54)$$

$$\begin{bmatrix} P_{x1} \\ P_{c1} \end{bmatrix} - \mathbb{E} \left(\begin{bmatrix} \Phi & 0 \\ E^T B^T & M^T \end{bmatrix} \left(\begin{bmatrix} P_x & P_{xc} \\ P_{xc}^T & P_c \end{bmatrix} \begin{bmatrix} w \\ \mathcal{L}w \end{bmatrix} + \begin{bmatrix} P_{x1} \\ P_{c1} \end{bmatrix} \right) \right) = 0 \quad (6.55)$$

$$P_1 - \mathbb{E} \left(\begin{bmatrix} w \\ \mathcal{L}w \end{bmatrix}^T \begin{bmatrix} P_x & P_{xc} \\ P_{xc}^T & P_c \end{bmatrix} \begin{bmatrix} w \\ \mathcal{L}w \end{bmatrix} + 2 \begin{bmatrix} P_{x1} \\ P_{c1} \end{bmatrix}^T \begin{bmatrix} w \\ \mathcal{L}w \end{bmatrix} \right) - P_1 = -L_{ss}. \quad (6.56)$$

Given that the expected value of w is zero and that the modulus of the eigenvalues of Φ and M are strictly less than 1, (6.55) implies that both P_{x1} and P_{c1} are zero. It can also be shown (see [19]) that if K is chosen to be the LQR optimal, then $P_{xc} = 0$.

Thus $J_t = V_t$ has the structure of (6.46).

Furthermore (6.56) is equivalent to

$$L_{ss} = \mathbb{E}(w^T \tilde{P}_x w) + \mathbb{E}(w^T \mathcal{L}^T \tilde{P}_c \mathcal{L} w) = \text{tr}(\tilde{P}_x \Sigma) + \text{tr}(\mathcal{L}^T \tilde{P}_c \mathcal{L} \Sigma), \quad (6.57)$$

and since $l_{ss} = \text{tr}(S\Sigma)$ where $S - \Phi^T S \Phi = Q + K^T R K$, from (6.54) it follows that $\mathbb{E}(w^T P_x w) = \text{tr}(\tilde{P}_x \Sigma) = l_{ss}$, thus L_{ss} can be written as in (6.49). $\square$

6.4.4 SMPC implementation and properties

The online optimization will only have the elements of the vector of degrees of freedom $\mathbf{c}_t$ as optimization variables. The set of admissible degrees of freedom $\mathcal{D}(x)$ is given by

$$\mathcal{D}(x) = \{\mathbf{c} : (6.44) \text{ holds}\}, \quad (6.58)$$

and the domain of attraction is given by

$$\mathcal{X} = \{x : \mathcal{D}(x) \neq \emptyset\}. \quad (6.59)$$

The online procedure consists of solving the optimal control problem $\mathbb{P}(x_t)$ that is defined by

$$V^0(x_t) = \min_{\mathbf{c}_t} J_t(x_t, \mathbf{c}_t), \text{ s.t. } \mathbf{c}_t \in \mathcal{D}(x_t), \quad (6.60a)$$

$$\mathbf{c}_t^0(x_t) = \arg \min_{\mathbf{c}_t} J_t(x_t, \mathbf{c}_t), \text{ s.t. } \mathbf{c}_t \in \mathcal{D}(x_t). \quad (6.60b)$$

Then, the SMPC control law is given by

$$\kappa^0(x_t) = Kx_t + c_{t|t}^0(x_t), \quad (6.61)$$

where $c_{t|t}^0(x_t)$ is the first block component of the optimal $\mathbf{c}_t^0(x_t)$. The control input

at time t is then given by $u_t := \kappa^0(x_t)$.

The optimal control problem $\mathbb{P}(x_t)$ can be formulated as the QP:

$$\begin{aligned} & \underset{\mathbf{c}_t}{\text{minimize}} && J_t(x_t, \mathbf{c}_t) \\ & \text{subject to} && \hat{E}_k(Dx_t + C_{1,N}\mathbf{c}_t) + \tilde{\beta}_k \leq \underline{1} - \gamma, \quad k = 0, \dots, N-1, \\ & && V_f(\Phi^N x_t + H_N \mathbf{c}_t) \leq \underline{1}. \end{aligned} \tag{6.62}$$

$\mathbb{P}(x_t)$ is a QP with $N_{var} = Nn_u$ variables subject to $N_{ineq} = Nn_c + n_f$ inequality constraints.

Theorem 6.13. *For system (6.1), the closed-loop application of the control law of (6.61) guarantees feasibility of (6.3), that the optimal control problem of (6.60) is recursively feasible and that $\mathcal{X}$ is invariant.*

Proof. The candidate solution at time $t+1$, $\tilde{\mathbf{c}}_{t+1}$, is built according to the *receding* procedure from the optimal solution at time t , $\mathbf{c}_t^0(x_t)$,

$$\tilde{c}_{t+1+k|t+1} = c_{t+k+1|t}^0(x_t) + L_{k+1}w_t, \quad k = 0, \dots, N-2, \tag{6.63a}$$

$$\tilde{c}_{t+1+N|t+1} = 0, \tag{6.63b}$$

which can alternatively be written as $\tilde{\mathbf{c}}_{t+1} = M\mathbf{c}_t + \mathcal{L}w_t$.

From (6.44) one gets

$$\hat{E}_{k+1}(Dx_k + C_{1,N}\mathbf{c}_t) \leq \underline{1} - \gamma - \tilde{\beta}_{k+1} = \underline{1} - \gamma - \tilde{\beta}_k - |\tilde{F}\Phi^k + \hat{C}_{k+1}\mathcal{L}|\bar{w}, \quad k = 0, 1, \dots \tag{6.64}$$

Then, using the facts that $\hat{E}_k D\Phi x_t = \hat{E}_{k+1} Dx_t$, $\hat{E}_k(DBc_{t|t} + C_{1,N}M\mathbf{c}_t) = C_{1,N}\mathbf{c}_t$ and $\hat{E}_k(D + C_{1,N}\mathcal{L})w_t = (\tilde{F}\Phi^k + \hat{C}_{k+1})w_t$ alongside (6.64), one gets that for $x_{t+1} =$

$\Phi x_t + Bc_t + w_t$ the receded $\tilde{\mathbf{c}}_{t+1}$ satisfies

$$\begin{aligned} \hat{E}_k D(\Phi x_t + Bc_{t|t} + w_t) \\ + \hat{E}_k C_{1,N}(M\mathbf{c}_t + \mathcal{L}w_t) &= \hat{E}_{k+1} D\Phi x_t + \hat{E}_{k+1} C_{1,N}\mathbf{c}_t + \hat{C}_{k+1}w_t \\ &\leq \underline{1} - \gamma - \tilde{\beta}_{k+1} + |\tilde{F}\Phi^k + \hat{C}_{k+1}\mathcal{L}|\bar{w} = \underline{1} - \gamma - \tilde{\beta}_k. \end{aligned} \quad (6.65)$$

and then x_{t+1} and $\tilde{\mathbf{c}}_{t+1}$ satisfy (6.44).

It follows that if $\mathbb{P}(x_t)$ is feasible, it will be feasible at the next instant, and by a recursive argument it will remain feasible for all future instants. Invariance of $\mathcal{X}$ follows directly. Satisfaction of (6.3) follows from the use of constraint (6.44) for $k = 0$. $\square$

Theorem 6.14. *For system (6.1), the closed-loop application of the control law of (6.61) guarantees the closed-loop performance bound:*

$$\lim_{\nu \rightarrow \infty} \frac{1}{\nu} \sum_{k=1}^{\nu} \mathbb{E}(x_k^T Q x_k + u_k^T R u_k) \leq L_{ss}. \quad (6.66)$$

Proof. From (6.52) for $k = 0$, considering that $V^0(x_{t+1}) \leq \chi_{t+1|t}^T V \chi_{t+1|t}$ because of the optimization involved, it follows that

$$V^0(x_t) - \mathbb{E}_t(V^0(x_{t+1})) \leq x_t^T Q x_t + u_t^T R u_t - L_{ss}. \quad (6.67)$$

Summing this for $t = 0, 1, \dots, \nu$ yields that

$$\frac{1}{\nu} \sum_{k=1}^{\nu} \mathbb{E}_0(l_k) \leq L_{ss} + \frac{1}{\nu} (\mathbb{E}_0(V_k) - \mathbb{E}_0(V_{k+\nu})), \quad (6.68)$$

and the desired result follows from taking limit of (6.68) when $\nu \rightarrow \infty$. $\square$

Remark 6.15. *In the case of no compensation, (6.66) applies with L_{ss} replaced by l_{ss} . Since the expected value of the stage cost will also always be greater than or*

equal to l_{ss} , it follows that this expected value may, in the best case, be equal to l_{ss} , something that can only happen if the control law at some point becomes $u = Kx$ thereby steering the state to the minimal robust invariant set $\Omega_\infty(\mathbb{W})$ for system (6.1) under $u = Kx$. However if the disturbance is sufficiently large so that $h < \gamma + \beta_\infty$, it becomes necessary to use disturbance compensation and therefore l_{ss} will not be attainable. When l_{ss} is attainable, the reduction in the predicted cost implied by (6.52) suggests that constraints will become less active with increasing time and that beyond a time they may become inactive. Noting that, by Theorem 6.12, $\tilde{P}$ has a block diagonal structure, it follows from (6.46) that $\mathbf{c}_t = 0$ is the unconstrained optimal of the predicted cost. This shows that if the constraints become inactive, the closed-loop control law, as is the case for no disturbance compensation, will revert to $u = Kx$ and therefore the state will be steered to $\Omega_\infty(\mathbb{W})$.

6.5 SMPC with online optimization of the disturbance compensation gains

The second proposed SMPC strategy considers $\mathbf{L}$ as variables in the online optimization.

6.5.1 Constraints and cost function

Since $\mathbf{L}$ is found online, the $\tilde{\beta}_k$ also have to be found online. This implies that the terminal set, as described in Section 6.4.2, cannot be used because the $\tilde{\beta}_k$ are unknown in the offline stage. For this reason, use will be made of the strategy of Chapter 4 for the handling of constraints. This can be done because the control policy of (6.25) is a particular case of that of (4.15) and (6.29) are a robustified version of the probabilistic constraints (6.6).

According to this, a secondary horizon N_2 is used such that the feasibility of (6.29) for $k = 0, \dots, N + N_2$ is imposed explicitly. Since the $\tilde{\beta}_k$ are unknown, (6.29) can

be invoked equivalently for $k = 0, \dots, N + N_2$ by using the dual variables $\mathbf{R}_k$ that satisfy $\mathbf{R}_1^T h + \dots + \mathbf{R}_k^T h \geq \tilde{\beta}_k$:

$$\hat{E}_0 D x_t + \hat{E}_0 C_{1,N} \mathbf{c}_t \leq \underline{1} - \gamma, \quad (6.69a)$$

$$\hat{E}_k D x_t + \hat{E}_k C_{1,N} \mathbf{c}_t + \sum_{j=1}^k \mathbf{R}_j^T h \leq \underline{1} - \gamma, \quad k = 1, \dots, N + N_2 - 1, \quad (6.69b)$$

such that the dual variables $\mathbf{R}_k$ satisfy

$$(\tilde{F}\Phi^{k-1} + \hat{C}_k \mathcal{L}) = \mathbf{R}_k^T S, \quad k = 1, \dots, N + N_2, \quad (6.70a)$$

$$\mathbf{R}_k \geq 0. \quad (6.70b)$$

The feasibility of (6.29) for $k \geq N + N_2$, on the other hand, is imposed by the conditions expressed in Proposition 4.2 and Corollary 4.3. Although these conditions are for a different control policy and for robust constraints, they can still be applied here. As shown in Section 6.3.1, the control policy of (6.25) is a particular case of that used in Chapter 4, given by (4.15), and, as will be shown next, constraints (6.29) are a robustified version of the probabilistic constraints (6.6).

In (6.29) $\tilde{\beta}_k$ account for the worst cases of $\varepsilon_{t+k|t}$, which depend on the KFD, and γ accounts for the probabilistic nature of the CFD, whose effect in (6.29) is given by Fw . Then, noting that by causality $u_{t+k|t}$ depends only on the KFD, using $z_{t+k|t} = \Phi^k x_t + H_k \mathbf{c}_t$ and $\vec{\varepsilon}_t = (C_1 L + C_2) \vec{w}_t$ and by arranging terms in (6.29), it follows that (6.29) is equivalent to

$$F A x_{t+k|t} + \tilde{G} u_{t+k|t} \leq \underline{1} - \gamma, \quad k = 0, 1, \dots, \quad (6.71)$$

which in turn is the same as $(x_{t+k|t}, u_{t+k|t}) \in \mathbb{Y}$ for $k = 0, 1, \dots$, where $\mathbb{Y} = \{F A x + \tilde{G} u \leq \underline{1} - \gamma\}$.

Let $\mathbb{X}_1 = \Phi^{N+N_2-1} \mathbb{W} = \{x : H^{(1)} x \leq \underline{1}\}$ with $H^{(1)} \in \mathbb{R}^{n_{f1} \times n}$, and let $\mathbb{X}_0 =$

$\{x : H^{(0)}x \leq \underline{1}\}$ be a positively robust invariant set for system $x^+ = \Phi x + w$ where $w \in \Phi \mathbb{X}_1$. Then, the conditions to guarantee the feasibility of (6.29) for $k \geq N + N_2$ on account of Proposition 4.2 and Corollary 4.3 are presented next.

Corollary 6.16. *If $\mathbf{c}_t$ and $\mathcal{L}$ satisfy*

$$H^{(0)}(\Phi^{N+N_2}x_t + H_{N+N_2}\mathbf{c}_t) \leq \underline{1}, \quad (6.72a)$$

$$\mathbf{P}^T h \leq \underline{1}, \quad (6.72b)$$

$$H^{(1)}\Phi^{N_2}[\Phi^{N-1} \quad \Phi^{N+N_2-2}B \quad \dots B]\mathcal{L} = \mathbf{P}^t S, \quad \mathbf{P}^T \geq 0, \quad (6.72c)$$

where $\mathbf{P} \in \mathbb{R}^{n_s \times n_{f1}}$, then, the satisfaction of (6.29) for $k \geq N + N_2$ is guaranteed by (6.70) and

$$f_{(0,\infty)} + \sum_{j=1}^k \mathbf{R}_j^T h + f_{(1,\infty)} \leq \underline{1}, \quad (6.73)$$

where $f_{(0,\infty)} = \max_{z \in \mathbb{X}_0} (FA + (FB + G)K)x$, $f_{(1,\infty)} = \max_{x \in \Omega_\infty(\mathbb{X}_1)} (FA + (FB + G)K)x$ and $\mathbf{R} \in \mathbb{R}^{n_s \times n_c}$.

Proof. From (6.34) it follows that $x_{(0,N+N_2)} = \Phi^{N+N_2}x_t + H_{N+N_2}\mathbf{c}_t$ and $X_{1,N+N_2} = \Phi^{N_2}(\Phi^{N-1} + \Phi^{N+N_2-2}BL_1 + \dots + BL_{N-1})\mathbb{W}$. Then, condition (6.72a) guarantees that $x_{(0,N+N_2)} \in \mathbb{X}_0$ and conditions (6.72b), (6.72c) guarantee that $X_{1,N+N_2} \subseteq \mathbb{X}_1$. Then by Proposition 4.2 $(x_{t+k|t}, u_{t+k|t}) \in \bar{Y}_\infty$ for all $k \geq N + N_2$. Finally (6.73) implies that $\bar{Y}_\infty \subseteq \mathbb{Y}$ which implies that $(x_{t+k|t}, u_{t+k|t}) \in \mathbb{Y}$ for all $k \geq N + N_2$, thus guaranteeing the satisfaction of (6.29) for $k \geq N + N_2$. $\square$

Remark 6.17. *Note that the satisfaction of constraints for the worst cases of w is imposed by an approach using the dual variables (such as $\mathbf{R}_k$ and $\mathbf{P}$) that use the description of the uncertainty of $\mathbb{W} = \{w : Sw \leq h\}$ as in (6.38) [38], [37], [90], [91]. This is different to Chapters 3 and 4, where the slack variables overbound the effect of the vertices $\tilde{w}_i$ of $\mathbb{W}$. This is relevant because as the dimension of the state increases, the size of S tends to grow linearly with n , whereas the number of vertices grows as*

2^n , so the approach used here is better suited to higher dimensional systems. This approach can be used here, and not in Chapters 3 and 4, due to the affine dependence on the disturbances of the control policy.

The fact that $\mathbf{L}$ is unknown in the offline stage also has implications for the cost function. In Section 6.4.3 the fixed values of $\mathbf{L}$ make the value of L_{ss} fixed, and then the cost of (6.45) can be expressed as a quadratic function of x_t and $\mathbf{c}_t$ as in (6.46). However, when $\mathbf{L}$ is a variable in the online optimization, the minimization of (6.45) makes little physical sense; the selection of $\mathbf{L}$ that maximizes L_{ss} is the same one that minimizes the cost of (6.45). Clearly this goes against the original aim of minimizing (6.8). Then, for practical reasons only, the proposed strategy will consider minimization of

$$\tilde{J}_t(x_t, \mathbf{c}_t, \mathcal{L}) = x_t^T P_x x_t + \mathbf{c}_t^T P_c \mathbf{c}_t + \alpha \operatorname{tr}(\mathcal{L}^T P_c \mathcal{L} \Sigma), \quad (6.74)$$

for some constant $\alpha \geq 0$. The cost (6.74), unlike that of (6.46), does not consider the effect of $\mathbf{L}$ in the cost of (6.45). Instead, in addition to the effect of x_t and $\mathbf{c}_t$ on (6.45), the last term of the RHS of (6.74) penalises the effect of $\mathbf{L}$ on L_{ss} (as can be seen in (6.49)), which is weighted by the factor α .

6.5.2 SMPC online implementation and properties

Let $\mathbf{d}$ be the vector of all the decision variables in the online optimization. It is composed of $\mathbf{c}_t$, $\mathbf{L}$, $\mathbf{R}_k$ and $\mathbf{P}$. The set of admissible decision variables $\mathcal{D}(x)$ is given by

$$\mathcal{D}(x) = \{\mathbf{d} : (6.69), (6.70), (6.72), (6.73) \text{ hold}\}, \quad (6.75)$$

and the domain of attraction is given by

$$\mathcal{X} = \{x : \mathcal{D}(x) \neq \emptyset\}. \quad (6.76)$$

The online procedure consists of solving the optimal control problem $\mathbb{P}(x_t)$ that is defined by

$$V^0(x_t) = \min_{\mathbf{d}} \tilde{J}_t(x_t, \mathbf{c}_t, \mathbf{L}), \text{ s.t. } \mathbf{d} \in \mathcal{D}(x_t), \quad (6.77a)$$

$$\mathbf{d}^0(x_t) = \arg \min_{\mathbf{d}} \tilde{J}_t(x_t, \mathbf{c}_t, \mathbf{L}), \text{ s.t. } \mathbf{d} \in \mathcal{D}(x_t). \quad (6.77b)$$

Then, the SMPC control law is given by

$$\kappa^0(x_t) = Kx_t + c_{t|t}^0(x_t), \quad (6.78)$$

where $c_{t|t}^0(x_t)$ is the first block component of $\mathbf{c}_t^0(x_t)$, which is part of the optimal $\mathbf{d}^0(x_t)$. The control input at time t is then given by $u_t := \kappa^0(x_t)$.

The optimal control problem $\mathbb{P}(x_t)$ can be formulated as the QP:

$$\begin{aligned} & \underset{\mathbf{d}}{\text{minimize}} && x_t^T P_x x_t + \mathbf{c}_t^T P_c \mathbf{c}_t + \alpha \text{tr}(\mathcal{L} P_c \mathcal{L} \Sigma) \\ & \text{subject to} && \hat{E}_0 D x_t + \hat{E}_0 C_{1,N} \mathbf{c}_t \leq \underline{1} - \gamma, \\ & && \hat{E}_k D x_t + \hat{E}_k C_{1,N} \mathbf{c}_t + \sum_{j=1}^k \mathbf{R}_j^T h \leq \underline{1} - \gamma, \quad k = 1, \dots, N + N_2 - 1, \\ & && H^{(0)}(\Phi^{N+N_2} x_t + H_{N+N_2} \mathbf{c}_t) \leq \underline{1}, \\ & && \mathbf{P}^T h \leq \underline{1}, \\ & && f_{(0,\infty)} + \sum_{j=1}^{N+N_2-1} \mathbf{R}_j^T h + f_{(1,\infty)} \leq \underline{1}, \\ & && (\tilde{F} \Phi^{k-1} + \hat{C}_k \mathcal{L}) = \mathbf{R}_k^T S, \quad k = 1, \dots, N + N_2 - 1, \\ & && H^{(1)} \Phi^{N_2} [\Phi^{N-1} \quad \Phi^{N+N_2-2} B \quad \dots B L_{N-1}] \mathcal{L} = \mathbf{P}^t S, \\ & && \mathbf{P}^T \geq 0, \quad \mathbf{R}_k \geq 0, \quad k = 1, \dots, N + N_2 - 1. \end{aligned} \quad (6.79)$$

$\mathbb{P}(x_t)$ is a QP with $N_{var} = N(2nn_c + n_u(1+n)) + 2nn_c N_2 + n(2n_{f1} - n_u - 2n_c)$ variables subject to $N_{ineq} = (N + N_2)n_c(1 + 2n) + n_c(1 - 2n) + n_{f0} + n_{f1}(1 + 2n) + n_c(1 - 2n)$

inequality constraints and $N_{eq} = n_c(N + N_2 - 1) + n_{f1}$ equality constraints.

Theorem 6.18. *For system (6.1), the closed-loop application of the control law of (6.78) guarantees feasibility of (6.3), that the optimal control problem of (6.77) is recursively feasible and that $\mathcal{X}$ is invariant.*

Proof. Recursive feasibility and invariance follow directly from Theorem 4.18, since constraints (6.69), (6.70), (6.72), (6.73) are equivalent to those used in the strategy of Section 4.4.2 such that $\mathbb{Y} = \{FAx + \tilde{G}u \leq \underline{1} - \gamma\}$. The feasibility of (6.3) is guaranteed by the use of (6.69) for $k = 0$. $\square$

Theorem 6.19. *For system (6.1), if $\mathcal{X}$ is bounded, the closed-loop application of the control law of (6.78) guarantees the closed-loop performance bound:*

$$\lim_{\nu \rightarrow \infty} \frac{1}{\nu} \sum_{t=1}^{\nu} \mathbb{E}(x_t^T Q x_t + u_t^T R u_t) \leq \varsigma + l_{ss}. \quad (6.80)$$

where $\varsigma = \max_{x \in \mathcal{X}} \text{tr}(\mathcal{L}^0(x)^T M^T P_c M \mathcal{L}^0(x) \Sigma)$ and $\mathcal{L}^0(x) = [0 \ L_1^0(x)^T \cdots L_{N-1}^0(x)^T]^T$, where $L_k^0(x)$, $k = 1, \dots, N - 1$ are part of the optimal $\mathbf{d}^0(x)$.

Proof. As in Theorem 6.18 (where the following construction is used implicitly) and consistently with Theorem 6.13, the candidate solution at time $t + 1$ is given by

$$\tilde{c}_{t+1+k|t+1} = c_{t+k+1|t}^0(x_t) + L_{k+1}^0(x_t), \quad k = 0, \dots, N - 2, \quad (6.81a)$$

$$\tilde{c}_{t+1+N|t+1} = 0, \quad (6.81b)$$

which can alternatively be written as $\tilde{\mathbf{c}}_{t+1} = M\mathbf{c}_t + \mathcal{L}^0(x_t)w_t$, and

$$\tilde{L}_k = L_k^0(x_t), \quad l = 1, \dots, N - 1. \quad (6.82)$$

Then, since P_x satisfies (6.54) and, from (6.47), P_x and P_c satisfy

$$P_c - (E^T B^T P_x B E + M^T P_c M) = E^T R E, \quad (6.83)$$

$$\Phi^T P_x B E = K^T R E, \quad (6.84)$$

it follows that

$$\begin{aligned} J^0(x_t) - \mathbb{E}_t(\tilde{J}_{t+1}(x_{t+1}, \tilde{\mathbf{c}}_t, \tilde{\mathcal{L}})) &\geq (x_t^T Q x_t + u_t^T R u_t) - \text{tr}(\mathcal{L}^0(x_t)^T M^T P_c M \mathcal{L}^0(x_t) \Sigma) \\ &\quad - E_t(w_t P_x w_t). \end{aligned} \quad (6.85)$$

Since $\varsigma = \max_{x \in \mathcal{X}} \text{tr}(\mathcal{L}^0(x)^T M^T P_c M \mathcal{L}^0(x) \Sigma)$, then $\varsigma \geq \text{tr}(\mathcal{L}^0(x_t)^T M^T P_c M \mathcal{L}^0(x_t) \Sigma)$ for all x_t . Using this result and $E_t(w_t P_x w_t) = \text{tr}(P_x \Sigma) = l_{ss}$, and applying $\mathbb{E}_0$ and summing (6.85) for $t = 0, 1, \dots, \nu$ gives

$$J^0(x_t) - E_0(J^0(x_\nu)) + \nu \varsigma + \nu l_{ss} \geq \sum_{t=1}^{\nu} \mathbb{E}_0(x_t^T Q x_t + u_t^T R u_t), \quad (6.86)$$

which by multiplying by $\frac{1}{\nu}$ and taking the limit of $\nu \rightarrow \infty$ yields the desired result. $\square$

Remark 6.20. *The bound of (6.80) is the same type of bound as that of (6.66), except that it accounts for the different $\mathcal{L}$ that may be used during the evolution of the system, and for this reason it is likely to be a looser bound.*

Remark 6.21. *Analogously to Remark 6.15, if $\gamma + \beta_\infty < \underline{1}$, then for x sufficiently close to the origin the solution given by $\mathbf{c}_t = 0$ and $L_k = 0$ for $k = 1, \dots, N - 1$ is feasible and optimal. This is because $\mathbf{c}_t = 0$ and $L_k = 0$ for $k = 1, \dots, N - 1$ are respectively the values that minimize the terms $\mathbf{c}_t^T P_c \mathbf{c}_t$ and $\text{tr}(\mathcal{L} P_c \mathcal{L} \Sigma)$ of (6.74). Then, if the constraints become inactive, the control law, as is the case for no disturbance compensation, will reduce to $u = Kx$ and therefore the state will be steered to $\Omega_\infty(\mathbb{W})$. On the other hand, if $\gamma + \beta_\infty > \underline{1}$, $\mathbf{c}_t = 0$ and $L_k = 0$ will never be feasible and extra compensation (i.e. $L_k \neq 0$ for at least one k) will always be needed.*

Remark 6.22. *Since the control policy of (6.25) is a particular case of the control policy of (4.15), one could also use the strategy of Section 4.4.1, which deploys a cost that measures the distance to a terminal set and an additional terminal constraint, in order to guarantee convergence to the maximal invariance set.*

6.6 Simulation results

This section presents two examples. The first illustrates the advantages of using feedback policies in SMPC over conventional approaches based on quasi closed-loop policies. This example also illustrates that in some cases the strategy with offline design of the gain can provide the same or even improved benefits, in terms of domain of attraction, as the one where the gains are optimized online. The second example, shows that in other cases the second strategy can provide larger domains of attraction.

Let SMPC0 denote the basic strategy without disturbance compensation of [49], let SMPC1 denote the strategy of Section 6.4, and let SMPC2 denote the strategy of Section 6.5. N_2 is chosen as in Section 4.5 to be the maximum between 5 and the minimum value such that $(I, K)\mathbb{X}_0 \oplus_{j=1}^{N_2} (I, K)\Phi^{j-1}\mathbb{W} \oplus (I, K)\Omega_\infty(\mathbb{X}_1) \subseteq (I, K)\Omega \subseteq \mathbb{Y}$ is feasible. The latter, on account of Remark 4.6, guarantees that the domain of attraction is not empty when $N \geq 1$. However, experimentation suggest that it is convenient to take N_2 to be at least 5 so as to relax the terminal constraints for the 0^{th} predicted state. W is taken to be $\Phi\mathbb{X}_1$, and α is chosen to be maximum of the components of $\Phi^{N_2}\underline{1}$.

In both examples described in this section the i.i.d. disturbances w_k are sampled from a truncated Gaussian distribution with bounds $\bar{w} = [0.5 \ 0.5]^T$ and variance $\Sigma = \text{diag}_{2 \times 2}\{(1/12)^2\}$ and the probability p is chosen to be $p = 0.8$. The cost is defined by $Q = I_{2 \times 2}$, where $I_{2 \times 2}$ is the identity matrix in $\mathbb{R}^{2 \times 2}$, and $R = 1$.

Example A. Consider the system defined by:

$$A = \begin{bmatrix} 0.926 & -0.0908 \\ -1.225 & 0.797 \end{bmatrix}, \quad B = \begin{bmatrix} -0.175 \\ 0.976 \end{bmatrix}, \quad F = \begin{bmatrix} 0.5 & -0.325 \\ 0.471 & -0.84 \\ -0.5 & 0.325 \\ -0.471 & 0.84 \end{bmatrix}, \quad G = \begin{bmatrix} 0.184 \\ 0.269 \\ -0.184 \\ -0.269 \end{bmatrix},$$

Firstly, Figure 6.1 shows the tightening parameters obtained with both SMPC0, without compensation, and of SMPC1, where the gains of the disturbance compensation are designed offline. It is clear that the sequential procedure of Algorithm 6.3 produces better (smaller) tightening parameters. Also, on account of the compensation that extends into Mode 2, the constraint relaxation goes into the infinite horizon. These improvements in the tightening parameters are clearly reflected in Table 6.1, which shows the enlargement of the domains of attraction.

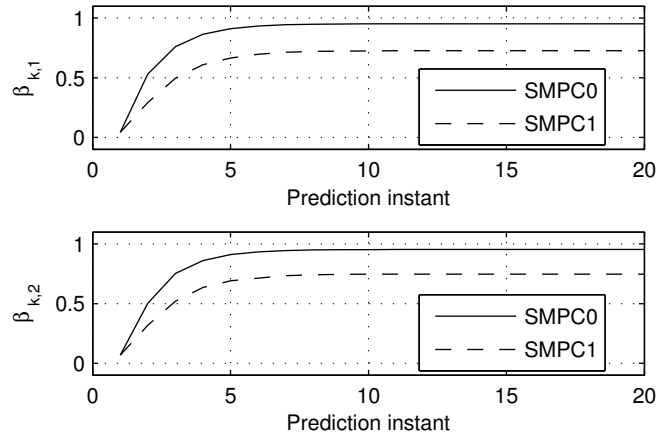


Figure 6.1: Constraint tightening with and without disturbance compensation.

Table 6.1 shows that for small values of N SMPC2 gives smaller domains than SMPC1, which seems counter-intuitive since SMPC2 computes the gains online and SMPC1 computes them offline. Furthermore, for $N = 2$, where there is one disturbance compensation gain, L_1 , SMPC2 does not provide any improvements over

SMPC0. The explanation of this is that the condition $X_{(1,N+N_2)} \subseteq \mathbb{X}_1$ is imposed in order to guarantee the satisfaction of constraints in the online optimization. While condition $x_{(0,N+N_2)} \subseteq \mathbb{X}_0$ is also imposed, it is easy to realise that the condition $X_{(1,N+N_2)} \subseteq \mathbb{X}_1$ is causing this conservativeness in this case because increasing the value of N_2 does not change the results obtained; larger values of N_2 relax $x_{(0,N+N_2)} \subseteq \mathbb{X}_0$ but not $X_{(1,N+N_2)} \subseteq \mathbb{X}_1$, so if the conservativeness was due to $x_{(0,N+N_2)} \subseteq \mathbb{X}_0$, increasing N_2 would improve the domain of attraction. However, as N increases above 2, there is more freedom in the different L_k to counteract this conservativeness, until for $N \geq 4$ SMPC2 reaches the maximum domain of attraction. In practice, amongst the examined results, for sufficiently large values of N the conservativeness $X_{(1,N+N_2)} \subseteq \mathbb{X}_1$ has always been cancelled so that the domains of attraction obtained with SMPC2 are never smaller than what can be obtained with SMPC1.

Note also that for $N = 1$ SMPC1 and SMPC2 do not have any gains to provide disturbance compensation, so there is no point in using these strategies with $N < 2$.

Table 6.1 also reports information relevant to the computational implementation, such as the number of variables, equalities and inequalities in the online optimization and average computational times. More precisely, the reported values correspond to the average time it takes the solver to solve the online optimization for a given prediction horizon, where the average is computed over a set of initial conditions evenly distributed at the edge of the domain of attraction of SMPC0. It is clear that SMPC0 and SMPC1 are much more convenient computationally on account of their reduced number of variables and constraints.

Table 6.1: Domain of Attraction and computational aspects.

N		1	2	3	4	5	6
A_N	SMPC0	147.5	197.4	208.1	210.7	210.9	210.9
	SMPC1	147.5	252.5	285.5	290.2	290.2	290.2
	SMPC2	147.5	197.4	269.8	290.2	290.2	290.2
vars	SMPC0	1	2	3	4	5	6
	SMPC1	1	2	3	4	5	6
	SMPC2	39	43	47	51	55	59
ineqs	SMPC0	14	22	30	38	46	52
	SMPC1	14	22	30	38	46	52
	SMPC2	154	174	194	214	234	254
eqs	SMPC0	0	0	0	0	0	0
	SMPC1	0	0	0	0	0	0
	SMPC2	28	32	36	40	44	48
time [ms]	SMPC0	1.3	1.5	1.7	1.8	1.5	1.6
	SMPC1	1.3	1.5	1.8	1.8	1.7	1.8
	SMPC2	2.2	11.2	12.2	15.6	17.7	20.46

Example B. Consider the system defined by:

$$A = \begin{bmatrix} -1.121 & -1.182 \\ -1.140 & -0.468 \end{bmatrix}, \quad B = \begin{bmatrix} 1.218 \\ 0.29 \end{bmatrix}, \quad F = \begin{bmatrix} 1.0 & 0.3165 \\ 0.2614 & -0.1340 \\ -1.0 & -0.3165 \\ -0.2614 & 0.1340 \end{bmatrix}, \quad G = \begin{bmatrix} 0.47 \\ 0.03 \\ -0.47 \\ -0.03 \end{bmatrix}.$$

Table 6.2 reports the areas of the domains of attraction for SMPC0, SMPC1 and SMPC2. This example shows that in some cases the conservativeness imposed by the condition $X_{(1,N+N_2)} \subseteq \mathbb{X}_1$ is not so relevant, since SMPC2 provides improvements over SMPC0 even for $N = 2$, although for this value of N no improvements were achieved in the previous example. Additionally, this example shows that sometimes Algorithm 6.3 might fail to provide improvements of the tightening parameters, since in this case $\tilde{\beta}_k = \beta_k$, for all $k = 1, 2, \dots$. While, as mentioned in Example A, in the performed simulations SMPC2 has always achieved domains of attraction as large as those of SMPC1 for a sufficiently large N , this example shows that SMPC1 does not necessarily achieve domains of attraction as large as those of SMPC2.

Table 6.2: Domain of Attraction and computational aspects.

N		2	3	4	5	6	7	8	9
A_N	SMPC0	28.84	33.30	38.80	43.59	49.30	55.16	59.70	59.70
	SMPC1	28.84	33.30	38.80	43.59	49.30	55.16	59.70	59.70
	SMPC2	29.78	43.04	55.99	69.63	69.63	69.63	69.63	69.63

6.7 Concluding remarks

Two SMPC strategies, SMPC1 and SMPC2, have been proposed which make use of an affine-in-the-disturbances policy with a striped structure. On account of the use of explicit disturbance compensation, more control authority is available in comparison with a strategy that only considers a quasi closed-loop policy [49], and this is reflected in improved domains of attraction.

SMPC1 considers an offline design of the disturbance compensation gains so as to improve the tightening parameters of the strategy of [49], whereas the gains are optimized online in SMPC2. On account of this, both strategies have a number of variables and constraints that grows linearly with the prediction horizon, but SMPC1 has fewer variables and constraints, and is as fast as the optimization in [49].

The constraint handling of SMPC1 is the same as in [49], except that the tightening parameters are smaller. Since the conditions used for ensuring satisfaction of the probabilistic constraints are equivalent to a robustification of the KFD, with an extra tightening parameter that accounts probabilistically for the CFD, the handling of constraints can be seen as a robust problem. Additionally, it has been noted that the control policy used here is a particular case of the control policy of Chapter 4. Therefore the approach of Chapter 4 is used for constraints satisfaction in SMPC2.

Although one might expect SMPC2 to have greater control authority than SMPC1 on account of the online optimization of the compensation gains, the use of the approach of Chapter 4 for constraints satisfaction imposes some conservativeness due to the extra conditions involved, so then, for small values of N , SMPC1 may have

larger domains of attraction than SMPC2. This effect, however, disappears for larger values of N as there are more L_k available, so that the maximum domains of attraction of SMPC2 are always as large or larger than those of SMPC1. These considerations, plus the fact that SMPC1 is considerably faster than SMPC2, imply that the decision between using one or another method is a compromise between control authority and computational speed.

The stability conditions for both strategies are expressed in terms of a bound on the closed-loop performance that depends on the disturbance compensation of the control policy: for the case of fixed gains (that have been designed offline) the bound is readily available and depends on those gains, while for SMPC2 the bound depends on the worst case possible value that these gains can take during the online execution. Clearly this does not yield a tight bound, thus it is desirable to pursue further research to prove a stability or convergence criterion to a terminal set or to the terminal control law $u = Kx$. Additionally, as stated in Remark 6.1, the current formulation cannot handle mixed probabilistic constraints of the type $\Pr\{F_j x_t + G_j u_t \leq 1\} \geq p_j$, so then it would be desirable to pursue further research so that this case can be covered.

Chapter 7

Concluding remarks and further research

7.1 Concluding remarks: a summary

This thesis presents a variety of strategies in the context of RMPC and SMPC with the aim of reducing online computational demand while retaining as much control authority as is conveniently possible.

Chapter 3 presents a simplification of PTMPC [73], referred to as offline PTMPC, or OPTMPC, and consists of transferring the computation of the part of the control policy associated with the uncertainty to the offline stage, so that the online optimization only deals with nominal predictions and as a result has a computational complexity comparable to that of deterministic linear MPC. Thus the control method is better suited to large prediction horizons. The offline stage consists of a sequential optimization of the parameters of the control policy that seeks to minimize the tightening parameters of the online optimization so as to relax the constraints of the nominal predictions. In order to be able to guarantee recursive feasibility and stability, the tightening parameters must be updated online based on offline calculations, so that the online optimization is time-varying. Nevertheless, after N instants the

tightening parameters will remain the same and the online optimization will become time-invariant. Although the offline computations introduce some conservativeness, this can be small enough so that often one can increase the prediction horizon used in OPTMPC so as to obtain larger domains of attraction than PTMPC (which uses a smaller N) while solving online optimizations faster and maintaining the system theoretic properties. OPTMPC thus provides a convenient alternative when the required domains of attraction are not too large: due to the offline design, in general, the maximum domain of attraction of OPTMPC will be smaller than that of PTMPC.

Chapter 4 presents a modification of PTMPC [73] such that the separable prediction scheme has a striped structure and the disturbance compensation terms extend over Mode 2 into the infinite horizon. The resulting strategy is referred to as striped PTMPC, or SPTMPC. The modification results in a reduction in the growth of the size of the optimization problem from quadratic to linear with N , so that the proposed strategy is better suited to large prediction horizons. Also, the extension of the disturbance compensation terms over Mode 2 into the infinite horizon provides extra relaxation in the terminal constraints, so that SPTMPC can potentially obtain better domains of attraction than PTMPC. Neither of the prediction schemes of PTMPC or SPTMPC is subsumed as a particular case of the other, so potentially any strategy can obtain larger domains of attraction than the other. However, the fact that the size of the online optimization only grows linearly in SPTMPC, as opposed to quadratically in PTMPC, implies that very often, by using larger prediction horizons, SPTMPC achieves better domains of attraction with faster online optimizations when compared to PTMPC. Two variations of the strategy are proposed, one that guarantees convergence of the state to the minimal invariant set under the state feedback, and another that only guarantees ISS stability. The first strategy however uses a terminal constraint that can impose some conservativeness, so the second strategy has increased control authority. The aspects previously discussed in this paragraph,

however, apply to both variations.

OPTMPC is faster than SPTMPC, and therefore larger prediction horizons could be used for OPTMPC which may allow one to obtain larger domains of attraction. However, the conservativeness imposed by the offline design in OPTMPC implies that often the domains of attraction cannot be as large as those of SPTMPC.

Chapter 5 presents a strategy that extends the approach of [19] for systems with multiplicative uncertainty, to the case of mixed multiplicative and additive uncertainty. The approach consists of the optimization of prediction dynamics with the aim of maximizing the volume of invariant ellipsoids. As in [34], this extension deploys disturbance feedback but is shown to depend on conditions which are both necessary and sufficient. The optimized dynamics together with perturbations on the control laws over the first N prediction steps is used to propose an effective prediction dual-mode prediction based RMPC algorithm which does not rely on the restrictive assumption of [34] that requires the future model uncertainty to be known over the prediction horizon. Constraint satisfaction is handled with an approach based on a variation of Farkas' Lemma, so that the resulting online optimizations have a problem size that grows linearly with the prediction horizon, and it thus suitable for large prediction horizons. The optimized polytopic dynamics yield the largest ellipsoidal invariant set in the sense that the addition of new degrees of freedom would not improve the ellipsoidal invariant set. However, the optimized polytopic dynamics are not necessarily optimal in terms of polytopic invariant sets. Therefore the use of polytopic constraints allows room for the introduction of extra degrees of freedom in the control policy which are used to improve the domain of attraction.

Unlike previous works on SMPC, Chapter 6 presents an SMPC approach that considers explicit feedback from the disturbances in the form of a striped affine-in-the-disturbances policy. Two variations of the strategy are proposed. The first variation, SMPC1, considers an offline design of the disturbance compensation gains

which, as in Chapter 3, consists of a sequential optimization of the compensation gains of the control policy that seeks to minimize the tightening parameters of the online optimization so as to relax the constraints of the nominal predictions. This results in a strategy with improved domain of attraction, while requiring online optimizations with the same computational demand as the strategy of [49], where the size of the online optimization grows linearly with the prediction horizon. The second variation, SMPC2, considers the gains as degrees of freedom of the online optimization. It is shown that the conditions for the satisfaction of the probabilistic constraints and for guaranteeing recursive feasibility can be cast as modified robust constraints, and then SMPC2 uses the approach of Chapter 4 for ensuring constraint satisfaction, but modified to account for the fact that the policy is affine in the disturbances. This allows one to impose constraints in a different way that is not so sensitive on the dimension of the additive disturbance. On account of the use of this approach for ensuring constraint satisfaction, the computational demands are greater than in SMPC1, but the size of the optimization problem still grows linearly with N . Additionally, the use of artificial terminal constraints can impose some conservativeness which may result in smaller domains of attraction than SMPC1. However, as larger prediction horizons are used this conservativeness disappears and, due to the fact that the gains are designed online, SMPC2 has greater control authority in general and yields better domains of attraction.

7.2 Further research

The first strategy proposed in Chapter 4, SPTMPC1, uses the additional constraint (4.34) to guarantee convergence to the minimal invariance set under $u = Kx$. It has been noted that use of this constraint is rather artificial because it considers a terminal set which is invariant for a fixed control law that is not the terminal control law of the predicted policy, which is not fixed. It would thus be interesting to study

the possibility of obtaining strong notions of stability (i.e. convergence to a set) without imposing that extra constraint, or in a way that is more directly related to the structure of the control policy.

In a related way, the stability conditions for both strategies of Chapter 6, SMPC1 and SMPC2, are expressed in terms of a bound on the closed-loop performance that depends on the disturbance compensation of the control policy. Thus it is desirable to pursue further research to prove a stochastic/probabilistic stability or convergence criterion to a terminal set or to the terminal control law $u = Kx$.

As it turns out, these two issues are related because both of them relate to the structure of the control policies (the control policy of Chapter 6 is a particular case of the control policy of Chapter 4), and its effect on the *receded* solutions that are typically used to prove recursive feasibility and stability. In a typical RMPC and SMPC formulation, the considered control policy and fixed terminal control law are such that the terminal control law can be implemented through a particular choice of the degrees of freedom in the control law at the $(N - 1)$ st predicted time step. This implies that the *receding* operation is valid for the entire predicted policy, thus proving recursive feasibility and that after N *receding* operations the entire control policy is given by the terminal control law. However, as explained at the beginning of Section 4.4, the striped structure of the control policy implies that the *receding* operation is not valid for the entire predicted policy (it is valid for the nominal portion, but not for the disturbance compensation terms). This was the reason behind the more involved constraint satisfaction approach in Chapter 4 and the need of the artificial condition (4.34) for proving convergence. This has to be borne in mind when pursuing the research lines proposed above.

The terminal control law of SPTMPC (and by extension also that of SMPC2) is not fixed and it can provide extra relaxation for the constraints throughout the infinite horizon, which suggests that the strategy could be applied to systems where

the state feedback $u = Kx$ is not necessarily feasible. Actually, Remark 6.15 indicates the conditions that the offline designed gains in SMPC1 need to satisfy so that the strategy is applicable, and these can be satisfied even if the state feedback $u = Kx$ is not necessarily feasible. This is a start to this interesting research direction, and further steps would involve finding conditions for the policies where its parameters are found online and performing stability analyses in an RMPC and SMPC context.

It has been discussed in Chapter 5 that the number of constraints and variables of the strategy depends on the number of rows of V , which in turn is associated with the eigenvalues of Ψ . V tends to have more rows as the eigenvalues approach to unity in absolute value, thus affecting the online computational demands. This suggests two possible research directions. One consists of finding alternatives for a more convenient selection of V . The other consists of finding alternative ways to impose the conditions (5.44), (5.40), (5.45) and (5.46), so that the online computational demands of the RMPC strategy are not so sensitive on the size of V .

Bibliography

- [1] T. Alamo, R. Tempo, and E. F. Camacho. Improved sample size bounds for probabilistic robust control design: A pack-based strategy. In *Decision and Control, 46th IEEE Conference on*, pages 6178–6183, Dec. 2007.
- [2] J. C. Allwright and G. C. Papavasiliou. On linear programming and robust model-predictive control using impulse-responses. *Systems & Control Letters*, 18(2):159–164, 1992.
- [3] Z. Artstein and S. V. Raković. Feedback and invariance under uncertainty via set-iterates. *Automatica*, 44(2):520–525, 2008.
- [4] K. J. Åström. *Introduction to Stochastic Control Theory*. Mathematics in science and engineering. Academic Press, 1970.
- [5] B. Bank, J. Guddat, D. Klatte, B. Kummer, and K. Tammer. *Non-linear Parametric Optimization*. Birkhauser, 1983.
- [6] I. Batina, A. A. Stoorvogel, and S. Weiland. Optimal control of linear, stochastic systems with state and input constraints. In *Decision and Control, Proceedings of the 41st IEEE Conference on*, volume 2, pages 1564–1569, 2002.
- [7] A. Bemporad, M. Morari, V. Dua, and E. N. Pistikopoulos. The explicit linear quadratic regulator for constrained systems. *Automatica*, 38(1):3–20, 2002.

- [8] D. Bernardini and A. Bemporad. Stabilizing model predictive control of stochastic constrained linear systems. *Automatic Control, IEEE Transactions on*, 57(6):1468–1480, June 2012.
- [9] D. Bertsekas. *Dynamic Programming and Optimal Control*. Athena Scientific, 1995.
- [10] D. Bertsimas, D. A. Iancu, and P. A. Parrilo. Optimality of affine policies in multi-stage robust optimization. In *Decision and Control, held jointly with the 28th Chinese Control Conference, Proceedings of the 48th IEEE Conference on*, pages 1131–1138, Dec. 2009.
- [11] F. Blanchini. Set invariance in control. *Automatica*, 35(11):1747–1767, 1999.
- [12] F. Blanchini and S. Miani. *Set-Theoretic Methods in Control*. Systems & control. Birkhäuser Boston, 2007.
- [13] S. Boyd, L. E. Ghaoui, E. Feron, and V. Balakrishnan. *Linear Matrix Inequalities in System and Control Theory*, volume 15 of *Studies in Applied Mathematics*. SIAM, Philadelphia, PA, June 1994.
- [14] G. C. Calafiore. On the expected probability of constraint violation in sampled convex programs. *Journal of Optimization Theory and Applications*, 143(2):405–412, 2009.
- [15] G. C. Calafiore and L. Fagiano. Robust model predictive control via scenario optimization. *Automatic Control, IEEE Transactions on*, 58(1):219–224, Jan. 2013.
- [16] G. C. Calafiore and L. Fagiano. Stochastic model predictive control of lpv systems via scenario optimization. *Automatica*, 49(6):1861–1866, 2013.

- [17] E. F. Camacho and C. Bordons. *Model Predictive Control, 2nd ed.* Advanced Textbooks in Control and Signal Processing. Springer Verlag, London, 2004.
- [18] P. J. Campo and M. Morari. Robust model predictive control. In *American Control Conference*, pages 1021–1026, June 1987.
- [19] M. Cannon and B. Kouvaritakis. Optimizing prediction dynamics for robust MPC. *Automatic Control, IEEE Transactions on*, 50(11):1892–1897, Nov. 2005.
- [20] M. Cannon, B. Kouvaritakis, and D. Ng. Probabilistic tubes in linear stochastic model predictive control. *Systems & Control Letters*, 58(1011):747–753, 2009.
- [21] M. Cannon, B. Kouvaritakis, S. V. Raković, and Q. Cheng. Stochastic tubes in model predictive control with probabilistic constraints. *Automatic Control, IEEE Transactions on*, 56(1):194–200, Jan. 2011.
- [22] M. Cannon, B. Kouvaritakis, and X. Wu. Model predictive control for systems with stochastic multiplicative uncertainty and probabilistic constraints. *Automatica*, 45(1):167–172, 2009.
- [23] M. Cannon, B. Kouvaritakis, and X. Wu. Probabilistic constrained MPC for multiplicative and additive stochastic uncertainty. *Automatic Control, IEEE Transactions on*, 54(7):1626–1632, July 2009.
- [24] M. Cannon, D. Ng, and B. Kouvaritakis. Successive linearization NMPC for a class of stochastic nonlinear systems. In L. Magni, D. M. Raimondo, and F. Allgwer, editors, *Nonlinear Model Predictive Control*, volume 384 of *Lecture Notes in Control and Information Sciences*, pages 249–262. Springer Berlin Heidelberg, 2009.
- [25] Q. Cheng. *Robust and Stochastic Model Predictive Control*. PhD thesis, Dept. of Engineering Science, University of Oxford, Oxford, United Kingdom, 2011.

- [26] Q. Cheng, M. Cannon, B. Kouvaritakis, and M. Evans. Stochastic MPC for systems with both multiplicative and additive disturbances. In *19th IFAC World Congress, Cape Town, SA*, 2014.
- [27] Q. Cheng, D. Muñoz Carpintero, M. Cannon, and B. Kouvaritakis. Efficient robust output feedback MPC. In *Control Conference (CCC), 2013 32nd Chinese*, pages 4149–4154, July 2013.
- [28] L. Chisci, J. A. Rossiter, and G. Zappa. Systems with persistent disturbances: predictive control with restricted constraints. *Automatica*, 37(7):1019–1028, 2001.
- [29] P. Couchman, B. Kouvaritakis, and M. Cannon. MPC on state space models with stochastic input map. In *Decision and Control, 45th IEEE Conference on*, pages 3216–3221, Dec. 2006.
- [30] G. de Nicolao, L. Magni, and R. Scattolini. On the robustness of receding-horizon control with terminal constraints. *Automatic Control, IEEE Transactions on*, 41(3):451–453, March 1996.
- [31] C. E. T. Dórea and J. C. Hennet. (A, B)-invariant polyhedral sets of linear discrete-time systems. *Journal of Optimization Theory and Applications*, 103(3):521–542, 1999.
- [32] M. Evans, M. Cannon, and B. Kouvaritakis. Linear stochastic MPC under finitely supported multiplicative uncertainty. In *American Control Conference*, pages 442–447, June 2012.
- [33] J. Fleming, B. Kouvaritakis, and M. Cannon. Regions of attraction and recursive feasibility in robust MPC. In *21st IEEE Mediterranean Conference on Control and Automation, Plataniass-Chania, Crete, Greece*, 2013.

- [34] A. Gautam, Y. C. Chu, and Y. C. Soh. Optimized dynamic policy for receding horizon control of linear time-varying systems with bounded disturbances. *Automatic Control, IEEE Transactions on*, 57(4):973–988, April 2012.
- [35] E. G. Gilbert and K. T. Tan. Linear systems with state and control constraints: the theory and application of maximal output admissible sets. *Automatic Control, IEEE Transactions on*, 36(9):1008–1020, Sep. 1991.
- [36] J. R. Gossner, B. Kouvaritakis, and J. A. Rossiter. Stable generalized predictive control with constraints and bounded disturbances. *Automatica*, 33(4):551–568, 1997.
- [37] P. J. Goulart and E. C. Kerrigan. Input-to-state stability of robust receding horizon control with an expected value cost. *Automatica*, 44(4):1171–1174, 2008.
- [38] P. J. Goulart, E. C. Kerrigan, and J. M. Maciejowski. Optimization over state feedback policies for robust control with constraints. *Automatica*, 42(4):523–533, 2006.
- [39] J. C. Hennet. Une extension du lemme de farkas et son application au probleme de régulation linéaire sous contraintes. *Comptes-Rendus de l’Académie des Sciendes, Série I*, 308:415–419, 1989.
- [40] L. Imsland, J. A. Rossiter, B. Pluymers, and J. Suikens. Robust triple mode MPC. *International Journal of Control*, 81(4):679–689, 2008.
- [41] Z. Jiang and Y. Wang. Input-to-state stability for discrete-time nonlinear systems. *Automatica*, 37(6):857–869, 2001.
- [42] E. C. Kerrigan and J. M. Maciejowski. Robustly stable feedback min-max model predictive control. In *American Control Conference*, volume 4, pages 3490–3495, June 2003.

- [43] B. Khan. *Efficient parameterise solutions of predictive control*. PhD thesis, Faculty of Engineering, University of Sheffield, Sheffield , United Kingdom, 2013.
- [44] I. Kolmanovsky and E. G. Gilbert. Theory and computation of disturbance invariant sets for discrete-time linear systems. *Mathematical Problems in Engineering*, 4(4):317–367, 1998.
- [45] M. V. Kothare, V. Balakrishnan, and M. Morari. Robust constrained model predictive control using linear matrix inequalities. *Automatica*, 32(10):1361–1379, 1996.
- [46] B. Kouvaritakis and M. Cannon. *Non-linear Predictive Control: Theory and Practice*. Iee Control Series, 61. The Institution of Engineering and Technology, London, UK, 2001.
- [47] B. Kouvaritakis, M. Cannon, and P. Couchman. MPC as a tool for sustainable development integrated policy assessment. *Automatic Control, IEEE Transactions on*, 51(1):145–149, Jan. 2006.
- [48] B. Kouvaritakis, M. Cannon, and D. Muñoz Carpintero. Efficient prediction strategies for disturbance compensation in stochastic MPC. *Int. Journal of Systems Science*, 44(7):1344–1353, 2013.
- [49] B. Kouvaritakis, M. Cannon, S. V. Raković, and Q. Cheng. Explicit use of probabilistic distributions in linear predictive control. *Automatica*, 46(10):1719–1724, 2010.
- [50] B. Kouvaritakis, M. Cannon, and R. Yadlin. On the centres and scalings of robust and stochastic tubes. In *UKACC Control*, 2010.
- [51] B. Kouvaritakis, J. A. Rossiter, and J Schuurmans. Efficient robust predictive control. *Automatic Control, IEEE Transactions on*, 45(8):1545–1549, 2000.

- [52] H. J. Kushner. *Introduction to stochastic control*. Holt, Rinehart and Winston, 1971.
- [53] Y. Kuwata, A Richards, and J. How. Robust receding horizon control using generalized constraint tightening. In *American Control Conference*, pages 4482–4487, July 2007.
- [54] W. Langson, I. Chrysoschoos, S. V. Raković, and D. Q. Mayne. Robust model predictive control using tubes. *Automatica*, 40(1):125–133, 2004.
- [55] M. Lazar. *Model predictive control of hybrid systems: Stability and robustness*. PhD thesis, Eindhoven University of Technology, The Netherlands, 2006.
- [56] Y. I. Lee, B. Kouvaritakis, and M. Cannon. Constrained receding horizon predictive control for nonlinear systems. *Automatica*, 38(12):2093–2102, 2002.
- [57] P. Li, M. Wendt, and G. Wozny. A probabilistically constrained model predictive controller. *Automatica*, 38(7):1171–1176, 2002.
- [58] J. Lofberg. Approximations of closed-loop minimax MPC. In *Decision and Control, Proceedings. 42nd IEEE Conference on*, volume 2, pages 1438–1442, 2003.
- [59] J. M. Maciejowski. *Predictive Control: With Constraints*. Pearson Education. Prentice Hall, 2002.
- [60] L. Magni and R. Sepulchre. Stability margins of nonlinear receding-horizon control via inverse optimality. *Systems & Control Letters*, 32(4):241–245, 1997.
- [61] D. Q. Mayne and W. Langson. Robustifying model predictive control of constrained linear systems. *Electronics Letters*, 37(23):1422–1423, 2001.

- [62] D. Q. Mayne, J. B. Rawlings, C. V. Rao, and P. O. M. Scokaert. Constrained model predictive control: Stability and optimality. *Automatica*, 36(6):789–814, 2000.
- [63] D.Q. Mayne, M.M. Seron, and S.V. Raković. Robust model predictive control of constrained linear systems with bounded disturbances. *Automatica*, 41(2):219–224, 2005.
- [64] H. Michalska and D. Q. Mayne. Robust receding horizon control of constrained nonlinear systems. *Automatic Control, IEEE Transactions on*, 38(11):1623–1633, 1993.
- [65] D. Muñoz Carpintero, M. Cannon, and B. Kouvaritakis. Recursively feasible robust MPC for linear systems with additive and multiplicative uncertainty using optimized polytopic dynamics. In *Decision and Control (CDC), 52nd IEEE Conference on, Florence, Italy*, pages 1101–1106, Dec. 2013.
- [66] D. Muñoz Carpintero, B. Kouvaritakis, and M. Cannon. On prediction strategies in stochastic MPC. In *Control and Automation (ICCA), 9th IEEE International Conference on, Santiago, Chile*, pages 249–254, Dec. 2011.
- [67] D. Muñoz Carpintero, B. Kouvaritakis, and M. Cannon. Striped Parameterized Tube Model Predictive Control. In *19th IFAC World Congress, Cape Town, SA*, Aug. 2014.
- [68] M. Ono. Joint chance-constrained model predictive control with probabilistic resolvability. In *American Control Conference*, pages 435–441, June 2012.
- [69] J. A. Primbs. Stochastic receding horizon control of constrained linear systems with state and control multiplicative noise. In *American Control Conference*, pages 4470–4475, July 2007.

- [70] J. A. Primbs and C. H. Sung. Stochastic receding horizon control of constrained linear systems with state and control multiplicative noise. *Automatic Control, IEEE Transactions on*, 54(2):221–230, Feb. 2009.
- [71] A. I. Propoi. Use of linear programming methods for synthesizing sampled-data automatic systems. *Automation and Remote Control*, 24(7):837–844, July 1963.
- [72] S. V. Raković, E. Kerrigan, K. Kouramas, and D. Mayne. Invariant approximations of the minimal robust positively invariant set. *IEEE Transactions on Automatic Control*, 50(3):406–410, 2005.
- [73] S. V. Raković, B. Kouvaritakis, M. Cannon, C. Panos, and R. Findeisen. Parameterized tube model predictive control. *Automatic Control, IEEE Transactions on*, 57(11):2746–2761, 2012.
- [74] S. V. Raković, B. Kouvaritakis, R. Findeisen, and M. Cannon. Simple homothetic tube model predictive control. In *19th Int. Symposium on Mathematical Theory of Networks and Systems, MTNS’10*, pages 1411–1418, 2010.
- [75] S. V. Raković, B. Kouvaritakis, R. Findeisen, and M. Cannon. Homothetic tube model predictive control. *Automatica*, 48(8):1631–1638, 2012.
- [76] S.V. Raković, D. Munoz-Carpintero, M. Cannon, and B. Kouvaritakis. Offline tube design for efficient implementation of parameterized tube model predictive control. In *Decision and Control (CDC), 51st IEEE Conference on, Maui, Hawaii*, pages 5176–5181, Dec. 2012.
- [77] C. V. Rao, S. J. Wright, and J. B. Rawlings. Application of interior-point methods to model predictive control. *Journal of Optimization Theory and Applications*, 99:723–757, 1998.

- [78] J. B. Rawlings and D. Q. Mayne. *Model predictive control: Theory and design*. Nob Hill Publishing, 2009.
- [79] J. Richalet, A. Rault, J. L. Testud, and J. Papon. Algorithmic control of industrial processes. In *Fourth IFAC symposium on identification and system parameter estimation*, pages 1119–1167, 1976.
- [80] J. Richalet, A. Rault, J. L. Testud, and J. Papon. Model predictive heuristic control: Applications to industrial processes. *Automatica*, 14:413–428, 1978.
- [81] A Richards and J. How. Robust stable model predictive control with constraint tightening. In *American Control Conference*, pages 6–., June 2006.
- [82] J. A. Rossiter, B. Kouvaritakis, and M. J. Rice. A numerically robust state-space approach to stable-predictive control strategies. *Automatica*, 34(1):65–73, 1998.
- [83] C. Scherer, P. Gahinet, and M. Chilali. Multiobjective output-feedback control via LMI optimization. *Automatic Control, IEEE Transactions on*, 42(7):896–911, July 1997.
- [84] J. Schuurmans and J. A. Rossiter. Robust predictive control using tight sets of predicted states. *Control Theory and Applications, IEE Proceedings -*, 147(1):13–18, 2000.
- [85] A. T. Schwarm and M. Nikolaou. Chance-constrained model predictive control. *AIChE Journal*, 45(8):1743–1752, 1999.
- [86] P. Scokaert and D. Mayne. Min-max feedback model predictive control for constrained linear systems. *Automatic Control, IEEE Transactions on*, 43(8):1136–1142, August 1998.
- [87] D. H. Van Hessem, C. W. Scherer, and O. H. Bosgra. LMI-based closed-loop economic optimization of stochastic process operation under state and input

- constraints. In *Decision and Control, Proceedings of the 40th IEEE Conference on*, volume 5, pages 4228–4233, 2001.
- [88] D.H. Van Hessem and O.H. Bosgra. Closed-loop stochastic dynamic process optimization under input and state constraints. In *American Control Conference*, volume 3, pages 2023–2028, 2002.
- [89] D.H. van Hessem and O.H. Bosgra. A conic reformulation of model predictive control including bounded and stochastic disturbances under state and input constraints. In *Decision and Control, Proceedings of the 41st IEEE Conference on*, volume 4, pages 4643–4648, 2002.
- [90] C. Wang, C. J. Ong, and M. Sim. Model predictive control using affine disturbance feedback for constrained linear system. In *Decision and Control, 2007 46th IEEE Conference on*, pages 1275–1280, Dec. 2007.
- [91] C. Wang, C. J. Ong, and M. Sim. Convergence properties of constrained linear system under MPC control law using affine disturbance feedback. *Automatica*, 45(7):1715–1720, 2009.
- [92] J. Yan and R. R. Bitmead. Incorporating state estimation into model predictive control and its application to network traffic control. *Automatica*, 41(4):595–604, 2005.
- [93] E. Zafiriou. Robust model predictive control of processes with hard constraints. *Computers & Chemical Engineering*, 14(45):359–371, 1990.
- [94] M. N. Zeilinger. *Real-time Model Predictive Control*. PhD thesis, Diss., Eidgenössische Technische Hochschule ETH Zürich, Nr. 19524, 2011.
- [95] Z. Q. Zheng and M. Morari. Robust stability of constrained model predictive control. In *American Control Conference*, pages 379–383, June 1993.