

A QUASI-INFINITE HORIZON NONLINEAR MODEL PREDICTIVE CONTROL SCHEME WITH GUARANTEED STABILITY

H. Chen

Institut für Systemdynamik und Regelungstechnik
Universität Stuttgart
Pfaffenwaldring 9, 70550 Stuttgart, Germany
chen@isr.uni-stuttgart.de

F. Allgöwer

Institut für Automatik
Eidg. Techn. Hochschule Zürich
CH-8092 Zürich, Switzerland
allgower@aut.ee.ethz.ch

Keywords: Nonlinear model predictive control, constrained system, stability.

Abstract

We present in this paper a novel nonlinear model predictive control scheme that guarantees closed-loop stability. The scheme can be applied to both stable and unstable systems with input constraints. The objective functional to be minimized consists of an integral square error (ISE) part over a finite time horizon plus a quadratic terminal cost. The terminal state penalty matrix of the terminal cost term has to be chosen as the solution of an appropriate Lyapunov equation. Furthermore, the setup includes a terminal inequality constraint that forces the states at the end of the finite horizon to lie within a prescribed terminal region. If the Jacobian linearization of the nonlinear system to be controlled is stabilizable, we prove that feasibility of the open-loop optimal control problem at time $t = 0$ implies asymptotic stability of the closed-loop system. The size of the region of attraction is only restricted by the requirement for feasibility of the optimization problem due to the input and terminal inequality constraints and is thus maximal in some sense.

1 Introduction

The history of model predictive control (MPC) began with an attempt to use the powerful computer technology to improve the control of processes that are constrained, multivariable and uncertain [5, 18]. In the last decade, many formulations have been developed for linear or nonlinear systems [6, 9], that found successful applications especially in the process industries [15, 17]. But still no completely satisfying approach is known, that gives both computational attractiveness and guaranteed closed-loop stability. This paper represents an attempt in that direction.

In general, the MPC problem is formulated as solving on-line a finite horizon open-loop optimal control prob-

lem subject to system dynamics and constraints involving states and inputs. However, as shown in [1], this general form of MPC does not guarantee closed-loop stability, because a finite horizon criterion cannot deliver asymptotic stability. Closed-loop stability can only be achieved by tuning design parameters such as prediction horizon and control horizon. Therefore, Bitmead [1] suggested an infinite horizon method (closely related to LQ control), which, however, results in an optimization problem that can generally be solved only for unconstrained linear systems.

For linear systems with constraints, the work of Rawlings and Muske [16] represents a significant leap forward in the MPC theory, where nominal closed-loop stability is guaranteed, independent of the specification of the desired control performance. For *constrained nonlinear* systems, a few restricted results can be found in the literature (e.g. [7, 11, 14, 19]), where Mayne and Michalska have contributed some very important issues. They have shown in [10, 12] that under some strong assumptions, finite horizon MPC is able to stabilize a class of nonlinear systems with constraints. The finite horizon constrained optimal control problem is posed as minimizing a standard quadratic objective functional subject to an additional terminal state equality constraint requiring the states to be zero at the end of the finite prediction horizon. However, from a computational point of view, an exact satisfaction of the terminal equality constraint requires an infinite number of iterations in the nonlinear case. An approximate satisfaction means that the achieved stability is lost inside the region of approximation. In order to avoid this, they extended their work in [13] with a terminal inequality constraint such that the states are on the boundary of a terminal region at the end of a *variable* prediction horizon. They suggested a *dual-mode* receding horizon control scheme with a local linear state feedback controller inside the terminal region and a receding horizon controller outside the terminal region. Closed-loop control with this scheme is implemented by switching between the two controllers, depending on the states being inside or outside

the terminal region.

In this paper, we introduce a quasi-infinite horizon nonlinear MPC scheme that optimizes on-line an objective functional including a finite horizon cost and a terminal cost subject to systems dynamics, input constraints and an additional terminal state inequality constraint. The feasibility of the terminal inequality constraint implies that the states at the end of the finite horizon are in a prescribed terminal region, in which the terminal states are penalized in such a way that the terminal cost bounds the infinite horizon cost of the nonlinear system controlled by a local linear state feedback. Thus, the proposed nonlinear model predictive controller has a *quasi-infinite prediction horizon*, but the input profile to be determined on-line is only of finite nature. If the Jacobian linearization of the nonlinear system to be controlled is stabilizable, the unique positive definite, symmetric solution of an appropriate Lyapunov equation can serve as terminal penalty matrix, and a neighborhood of the origin serving as terminal region can be determined off-line. The closed-loop asymptotic stability is then guaranteed by the feasibility of the open-loop optimal control problem at time $t = 0$. Because the local linear state feedback is only used to calculate the terminal penalty matrix and the terminal region off-line, no switching between controllers is needed in calculating the closed-loop control, no matter whether the states are inside or outside the terminal region.

The paper is structured as follows: Section 2 describes the proposed quasi-infinite horizon nonlinear MPC problem. Section 3 gives some preliminary results about a region of attraction and a performance bound of the nonlinear system controlled by a local linear state feedback. A procedure for determining the terminal region and the terminal penalty matrix off-line is outlined. In Section 4, closed-loop stability is discussed and sufficient stability conditions are given. Simulation results for an unstable system are shown in Section 5.

2 Problem setup

The class of systems to be controlled is described by the following general nonlinear set of ODEs

$$\dot{\mathbf{x}}(t) = \mathbf{f}(\mathbf{x}(t), \mathbf{u}(t)), \quad \mathbf{x}(0) = \mathbf{x}_0 \quad (1)$$

with state vector $\mathbf{x}(t) \in \mathbb{R}^n$, input vector $\mathbf{u}(t) \in \mathbb{R}^m$ and subject to input constraints $\mathbf{u}(t) \in U$. We consider in this paper the state feedback case and thus assume that all states are measurable. It is also assumed that

- A1) $\mathbf{f} : \mathbb{R}^n \times \mathbb{R}^m \rightarrow \mathbb{R}^n$ is twice continuously differentiable and $\mathbf{f}(\mathbf{0}, \mathbf{0}) = \mathbf{0}$. Thus, $\mathbf{0} \in \mathbb{R}^n$ is an equilibrium of the system,
- A2) $U \subset \mathbb{R}^m$ is compact, convex and $\mathbf{0} \in \mathbb{R}^m$ is contained in its interior,
- A3) system (1) has a unique solution for any initial condition $\mathbf{x}_0 \in \mathbb{R}^n$ and any piecewise continuous $\mathbf{u}(\cdot) : [0, \infty) \rightarrow U$.

We shall first introduce some notation: For any vector $\mathbf{x} \in \mathbb{R}^n$, $\|\mathbf{x}\|$ denotes the 2-norm and the weighted norm $\|\mathbf{x}\|_P$ is defined by $\|\mathbf{x}\|_P^2 := \mathbf{x}^T P \mathbf{x}$, where P is an arbitrary Hermitian, positive definite matrix. For any Hermitian matrix A , $\lambda_{\max}(A)$ and $\lambda_{\min}(A)$ denote the largest and the smallest real part of the eigenvalues of the matrix A , and $\|A\|$ stands for the induced 2-norm of A . In order to distinguish clearly between the “real” system and the system model used to predict the future “within” the controller, we denote the internal variables in the controller by a bar ($\bar{\mathbf{x}}, \bar{\mathbf{u}}$) to indicate that the predicted values need not and will not be the same as the actual values.

In the following, we describe the problem setup for the introduced quasi-infinite horizon nonlinear MPC scheme. The open-loop optimal control problem at time t is formulated as

$$\min_{\bar{\mathbf{u}}} J(\mathbf{x}(t), \bar{\mathbf{u}}(\cdot)) \quad (2)$$

with

$$J(\mathbf{x}(t), \bar{\mathbf{u}}(\cdot)) = \int_t^{t+T_p} (\|\bar{\mathbf{x}}(\tau; \mathbf{x}(t), t)\|_Q^2 + \|\bar{\mathbf{u}}(\tau)\|_R^2) d\tau + \|\bar{\mathbf{x}}(t+T_p; \mathbf{x}(t), t)\|_P^2 \quad (3)$$

subject to

$$\dot{\bar{\mathbf{x}}} = \mathbf{f}(\bar{\mathbf{x}}, \bar{\mathbf{u}}), \quad \bar{\mathbf{x}}(t; \mathbf{x}(t), t) = \mathbf{x}(t) \quad (4a)$$

$$\bar{\mathbf{u}}(\tau) \in U \quad (4b)$$

$$\bar{\mathbf{x}}(t+T_p; \mathbf{x}(t), t) \in \Omega, \quad (4c)$$

where $Q \in \mathbb{R}^{n \times n}$ and $R \in \mathbb{R}^{m \times m}$ denote positive definite, symmetric weighting matrices; T_p is a finite horizon; $\bar{\mathbf{x}}(\cdot; \mathbf{x}(t), t)$ is the trajectory of (4a) driven by $\bar{\mathbf{u}}(\cdot) : [t, t+T_p] \rightarrow U$. Note the initial condition in (4a): The system model used to predict the future in the controller is initialized by the actual, measured system states.

The objective functional (3) consists of a finite horizon standard cost to specify the desired control performance and a terminal cost to penalize the states at the end of the finite horizon. The terminal inequality constraint (4c) will force the states at the end of the finite horizon to be in some neighborhood of the origin, called here *terminal region* Ω . As we will see, in the terminal region, the quadratic terminal cost $\|\bar{\mathbf{x}}(t+T_p; \mathbf{x}(t), t)\|_P^2$ bounds the *infinite* horizon cost of the nonlinear system model controlled fictitiously by a local linear state feedback, *i.e.*,

$$\|\bar{\mathbf{x}}(t+T_p; \mathbf{x}(t), t)\|_P^2 \geq \int_{t+T_p}^{\infty} (\|\bar{\mathbf{x}}(\tau; \mathbf{x}(t), t)\|_Q^2 + \|\bar{\mathbf{u}}(\tau)\|_R^2) d\tau + \bar{\mathbf{u}} = K\bar{\mathbf{x}}, \forall \bar{\mathbf{x}}(t+T_p; \mathbf{x}(t), t) \in \Omega. \quad (5)$$

The positive definite and symmetric *terminal penalty matrix* $P \in \mathbb{R}^{n \times n}$ together with the terminal region Ω is determined off-line such that inequality (5) holds true and the input constraints are satisfied. If we substitute (5) into (3), we have

$$J(\mathbf{x}(t), \bar{\mathbf{u}}) \geq \int_t^{\infty} (\|\bar{\mathbf{x}}(\tau; \mathbf{x}(t), t)\|_Q^2 + \|\bar{\mathbf{u}}(\tau)\|_R^2) d\tau,$$

where $\bar{\mathbf{u}}(\tau) = K\bar{\mathbf{x}}(\tau; \mathbf{x}(t), t)$ for $\tau \geq t + T_p$. In this sense, the prediction horizon of the proposed nonlinear predictive controller expands quasi to infinity.

The optimal solution of (2) with (3) subject to (4) at time t , (existence assumed), will be denoted by $\bar{\mathbf{u}}^*(\cdot; \mathbf{x}(t), t) : [t, t + T_p] \rightarrow U$ and the corresponding optimal value by $J^*(\mathbf{x}(t)) := J(\mathbf{x}(t), \bar{\mathbf{u}}^*(\cdot))$. In the MPC implementation, the input profile is applied to the system only until the next measurement becomes available. We assume that this will be the case every δ time-units. So δ denotes the “sampling time” and the closed-loop control is defined by

$$\mathbf{u}^*(\tau) := \bar{\mathbf{u}}^*(\tau; \mathbf{x}(t), t), \tau \in [t, t + \delta]. \quad (6)$$

With the new measurement, the optimization problem will be solved again. Thus, the local linear state feedback controller will never be directly applied to the system. It is only used to determine the terminal penalty matrix P and the terminal region Ω off-line, as described in the next section.

3 Preliminary results

By slight modification of the associated content in [13], we present preliminary results about a region of attraction and a performance bound of the nonlinear system controlled by a local linear state feedback. These results allow us to outline a procedure to systematically determine the terminal region and the terminal penalty matrix, and will be used to prove closed-loop stability of the proposed control scheme. Since the terminal region and the terminal penalty matrix can be calculated *off-line*, variables without a bar will be used in this section.

We consider the Jacobian linearization of (1) at the origin

$$\dot{\mathbf{x}} = A\mathbf{x} + B\mathbf{u}, \quad (7)$$

where $A := \frac{\partial \mathbf{f}}{\partial \mathbf{x}}(\mathbf{0}, \mathbf{0})$ and $B := \frac{\partial \mathbf{f}}{\partial \mathbf{u}}(\mathbf{0}, \mathbf{0})$. If (7) is stabilizable, then a linear state feedback $\mathbf{u} = K\mathbf{x}$ can be determined such that $A_K := A + BK$ is asymptotically stable. Thus, the closed-loop system

$$\dot{\mathbf{x}} = \mathbf{f}(\mathbf{x}, K\mathbf{x}) \quad (8)$$

is locally asymptotically stable. The following lemma can be stated.

Lemma 1 *Suppose that the Jacobian linearization of (1) at the origin is stabilizable. Then,*

(a) *the Lyapunov equation*

$$(A_K + \kappa I)^T P + P(A_K + \kappa I) = -Q^* \quad (9)$$

with $Q^ = Q + K^T R K$ and $\kappa \in [0, -\lambda_{\max}(A_K))$ admits a unique positive definite and symmetric solution P ,*

(b) *there exists a constant $\alpha \in (0, \infty)$ specifying a neighborhood Ω of the origin in the form of*

$$\Omega := \{\mathbf{x} \in \mathbb{R}^n | \mathbf{x}^T P \mathbf{x} \leq \alpha\} \quad (10)$$

such that $K\mathbf{x} \in U$, for all $\mathbf{x} \in \Omega$, and Ω is invariant for (8).

Sketch of the proof: Since Q^* is positive definite and symmetric, (a) is true by the general conditions for the solvability of Lyapunov equation.

Due to $\mathbf{0} \in \mathbb{R}^m$ being contained in the interior of U , we can always find an $\alpha_1 \in (0, \infty)$ such that $K\mathbf{x} \in U$, for all $\mathbf{x} \in \Omega_1 := \{\mathbf{x} \in \mathbb{R}^n | \mathbf{x}^T P \mathbf{x} \leq \alpha_1\}$.

Let $\alpha \in (0, \alpha_1]$ specify a region in the form of (10). Since $\alpha \leq \alpha_1$, the input constraints are satisfied in Ω and need not be concerned.

We differentiate $\mathbf{x}^T P \mathbf{x}$ along a trajectory of (8) and obtain

$$\begin{aligned} \frac{d}{dt} \mathbf{x}(t)^T P \mathbf{x}(t) &= \mathbf{x}(t)^T (A_K^T P + P A_K) \mathbf{x}(t) \\ &\quad + 2\mathbf{x}(t)^T P \phi(\mathbf{x}(t)), \end{aligned} \quad (11)$$

where $\phi(\mathbf{x}) := \mathbf{f}(\mathbf{x}, K\mathbf{x}) - A_K \mathbf{x}$. The term involving $\phi(\mathbf{x})$ in the above inequality is bounded above as follows:

$$\mathbf{x}^T P \phi(\mathbf{x}) \leq \frac{\|P\| \cdot L_\phi}{\lambda_{\min}(P)} \|\mathbf{x}\|_P^2, \quad (12)$$

where $L_\phi := \sup \left\{ \frac{\|\phi(\mathbf{x})\|}{\|\mathbf{x}\|} \mid \mathbf{x} \in \Omega, \mathbf{x} \neq \mathbf{0} \right\}$. Now we choose an $\alpha \in (0, \alpha_1]$ such that in Ω

$$L_\phi \leq \frac{\kappa \cdot \lambda_{\min}(P)}{\|P\|}. \quad (13)$$

Then, inequality (12) leads to

$$\mathbf{x}^T P \phi(\mathbf{x}) \leq \kappa \cdot \mathbf{x}^T P \mathbf{x}. \quad (14)$$

Substituting (14) into (11) yields

$$\frac{d}{dt} \mathbf{x}(t)^T P \mathbf{x}(t) \leq -\mathbf{x}(t)^T Q^* \mathbf{x}(t), \quad (15)$$

where (9) is used.

Since $P > 0$ and $Q^* > 0$, inequality (15) implies that the region Ω defined by (10) is invariant for (8). Moreover, any trajectory of (8) starting in Ω stays in Ω and converges to the origin. $\square$

If we use the notation introduced in Section 2 for internal variables in the controller, integrating (15) from $t + T_p$ to ∞ with initial state $\bar{\mathbf{x}}(t + T_p; \mathbf{x}(t), t) \in \Omega$ yields (5). The solution P of (9) and the region Ω in the form of (10) are able to serve as terminal penalty matrix and terminal region. A procedure can be stated to determine a terminal penalty matrix P and a terminal region Ω (preferably as large as possible) off-line:

Step 1: solve a control problem based on the Jacobian linearization to get a locally stabilizing linear state feedback matrix K ,

- Step 2: choose a constant $\kappa \in [0, -\lambda_{\max}(A_K))$ and solve the Lyapunov equation (9) to get a positive definite and symmetric P ,
- Step 3: find the largest possible $\alpha_1 \in (0, \infty)$ such that $K\mathbf{x} \in U$, for all $\mathbf{x} \in \Omega_1$,
- Step 4: find the largest possible $\alpha \in (0, \alpha_1]$ in (10) such that inequality (13) is satisfied in Ω , or less conservatively, such that the optimal value of the simple optimization problem is nonpositive:

$$\max_{\mathbf{x}} \{ \mathbf{x}^T P \phi(\mathbf{x}) - \kappa \cdot \mathbf{x}^T P \mathbf{x} \mid \mathbf{x}^T P \mathbf{x} \leq \alpha \}. \quad (16)$$

A discussion on the optimization problem (16) and on finding the maximum α_1 in Step 3 can be found in [13].

4 Asymptotic stability result

In this section, we make use of the notation introduced in Section 2 and distinguish between the predicted values in the controller and the actual ones in the “real” system. That is, we use $\bar{\mathbf{x}}(\cdot; \mathbf{x}(t), t)$ to denote the predicted trajectory of the system model (4a) starting from the actual, measured state $\mathbf{x}(t)$ and driven by $\bar{\mathbf{u}}(\cdot)$, when the prediction is made at “real” time t . Now we address the asymptotic stability of the closed-loop system with the proposed quasi-infinite horizon nonlinear model predictive controller (6):

$$\dot{\mathbf{x}}(t) = \mathbf{f}(\mathbf{x}(t), \mathbf{u}^*(t)), \quad \mathbf{x}(0) = \mathbf{x}_0. \quad (17)$$

Theorem 1 *Suppose that*

- (a) *assumptions A1) – A3) are satisfied,*
- (b) *the Jacobian linearization of the nonlinear system (1) is stabilizable,*
- (c) *the open-loop optimal control problem (2) with (3) subject to (4) is feasible at time $t = 0$,*
- (d) *for $\mathbf{x}_0 = \mathbf{0}$, the optimization problem can be globally optimally solved.*

Then, for small sampling time $\delta > 0$ and in the absence of disturbances, the closed-loop system (17) is nominally asymptotically stable. Let $X \subseteq \mathbb{R}^n$ denote the set of all initial states for which assumption (c) is satisfied, then, X gives a region of attraction for the closed-loop system.

Sketch of the proof: According to the principle of MPC, the optimization problem (2) – (4) has to be solved repeatedly, whenever new measurements are available. So closed-loop stability requires first the feasibility of the optimization problem at each time. Following a standard argument also used for example in [7, 13, 16], we can show that, for the nominal system without disturbances and for small sampling time $\delta > 0$, the feasibility of the optimization problem at time $t = 0$ implies its feasibility for all $t > 0$. Thus, for any initial state $\mathbf{x}_0 \in X$, the optimization problem is feasible at each time.

In order to prove closed-loop stability, we secondly show that the optimal value function is strictly decreasing, as long as the closed-loop system is not at the origin. To do this, we start out with the optimal input profile $\bar{\mathbf{u}}^*(\cdot; \mathbf{x}(t), t) : [t, t + T_p] \rightarrow U$ determined at time t and construct a feasible (but not necessarily optimal) solution to the optimization problem at time $t + \delta$. For this feasible input profile we can compute the corresponding value of the objective functional that we denote by $\bar{J}(\mathbf{x}(t + \delta)) := J(\mathbf{x}(t + \delta), \bar{\mathbf{u}}(\cdot))$ and show that

$$\bar{J}(\mathbf{x}(t + \delta)) \leq J^*(\mathbf{x}(t)) - \int_t^{t+\delta} (\|\mathbf{x}(\tau)\|_Q^2 + \|\mathbf{u}(\tau)\|_R^2) d\tau.$$

It is clear that the optimal solution at time $t + \delta$ will be not worse than the feasible one. Thus, we have

$$J^*(\mathbf{x}(t + \delta)) \leq J^*(\mathbf{x}(t)) - \int_t^{t+\delta} (\|\mathbf{x}(\tau)\|_Q^2 + \|\mathbf{u}(\tau)\|_R^2) d\tau.$$

Then, we define a function $V(\mathbf{x})$ for the closed-loop system (17) as $V(\mathbf{x}) := J^*(\mathbf{x})$ and show that it has the properties that $V(\mathbf{0}) = 0$, $V(\mathbf{x}) > 0$ for $\mathbf{x} \neq \mathbf{0}$ and $V(\mathbf{x})$ is continuous at $\mathbf{x} = \mathbf{0}$. Hence, using a standard argument (e.g. [8]) we can prove that the closed-loop system is stable. Furthermore, we can show that $\mathbf{x}(t) \rightarrow \mathbf{0}$ as $t \rightarrow \infty$, without having to use a continuous differentiability assumption on $V(\mathbf{x})$. Thus, the closed-loop system is asymptotically stable. Finally, we can prove by contradiction that X is invariant for the closed-loop system and hence belongs to the region of attraction for the closed-loop system. For a complete proof see [3]. $\square$

Remark 4.1 *When applying this control scheme to practical systems, the numerical optimization employed may not find the globally optimal input profile at every time step, but may for example be caught in a local optimum. Even though optimal performance might be lost this way, stability can be guaranteed nevertheless, as the stability guarantee does not depend on the optimality of the solution found but merely on its feasibility. Indeed, starting with the shifted feasible solution from the previous step, the optimizer can find another feasible solution that is not worse than the shifted one at each step.*

5 Example

As an example for demonstrating the proposed control scheme, we consider the following nonlinear system:

$$\dot{x}_1 = x_2 + u(0.5 + 0.5x_1) \quad (18a)$$

$$\dot{x}_2 = x_1 + u(0.5 - 2x_2). \quad (18b)$$

This system is a modification of the system used in [10] in that it is unstable and its linearized system is stabilizable (not controllable). In addition, it is assumed that the input u suffers the following constraint:

$$U = \{u \in \mathbb{R}^1 \mid -2.0 \leq u \leq 2.0\}. \quad (19)$$

For this unstable constrained system, assumptions A1) – A3) are satisfied. The weighting matrices Q and R in the objective functional (3) are chosen as

$$Q = \begin{pmatrix} 0.5 & 0.0 \\ 0.0 & 0.5 \end{pmatrix}, \quad R = 1.0. \quad (20)$$

In order to find a terminal penalty matrix P and a largest possible terminal region Ω , we follow the procedure described in Section 3. First, solving the linear optimal control problem with the weighting matrices given in (20) for the linearized system, we get a linear locally stabilizing state feedback matrix

$$K = [2.118 \quad 2.118]. \quad (21)$$

The largest eigenvalue of the closed-loop linearized system is $\lambda_{max}(A_K) = -1.0$. Then, we choose a constant $\kappa = 0.95$ which implies that the unique solution of the Lyapunov equation (9)

$$P = \begin{pmatrix} 16.5926 & 11.5926 \\ 11.5926 & 16.5926 \end{pmatrix} \quad (22)$$

is positive definite, symmetric and can be used as terminal penalty matrix. After that, $\alpha_1 = 12.5$ is found to specify a region in which the linear feedback control satisfies the input constraint (19). Finally, we find a region $\Omega = \{\mathbf{x} \in \mathbb{R}^2 | \mathbf{x}^T P \mathbf{x} \leq 0.025\}$ such that inequality (13) is satisfied. However, this region is very small, because of the small value (0.1774) of $\frac{\lambda_{min}(P)}{\|P\|}$. From the simple optimization (16), we can however easily derive a less conservative terminal region:

$$\Omega = \{\mathbf{x} \in \mathbb{R}^2 | \mathbf{x}^T P \mathbf{x} \leq 0.7\}. \quad (23)$$

The open-loop optimal control problem described by (2) – (4) is solved in discrete time with a sampling time of $\delta = 0.1$ time-units and a prediction horizon of $T_p = 1.5$ time-units. A few trajectories corresponding to different initial conditions of the unstable constrained system (18) controlled by the proposed controller with parameters (20), (22) and (23) are shown in Fig. 1. The solid lines represent closed-loop trajectories; the dashed line is the boundary of the terminal region given by (23); the dash-dotted lines denote the predicted open-loop trajectories that are found by solving the optimization problem at time $t = 0$ and are only of finite horizon. It is clearly seen that the finite horizon open-loop trajectories end in the terminal region, as expected to be achieved by the terminal inequality constraint.

It should be emphasized that the linear state feedback K in (21) is *not explicitly used* to calculate the closed-loop control. The closed-loop control is determined by solving on-line the optimization problem given by (2) – (4) repeatedly. For the chosen prediction horizon T_p and sampling time δ , the optimization problem is feasible at each time. Thus, for the nominal system without disturbances, the stability conditions given in Section 4 are all

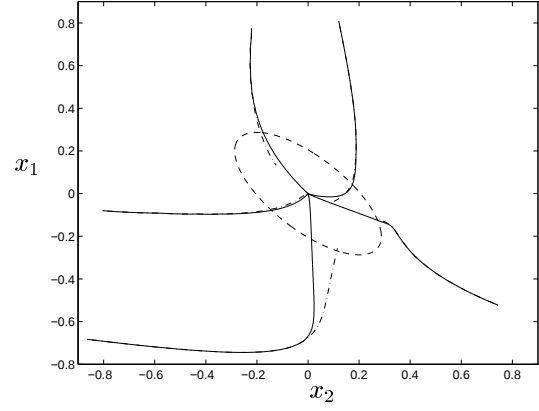


Figure 1: Trajectories from different initial states.

satisfied. The closed-loop trajectories in Fig. 1 are guaranteed to converge to the origin. Time profiles for the closed-loop system (not presented) show that the input constraint (19) is not violated.

The proposed nonlinear MPC scheme has significant computational advantages when compared to other approaches. To show this, we compare the proposed controller (called controller A) to two other predictive controllers (B, C):

- A : P given by (22), terminal inequality constraint $\bar{\mathbf{x}}(t + T_p; \mathbf{x}(t), t) \in \Omega$ with Ω given by (23), $T_p = 1.5$,
- B : $P = 0$, no terminal constraint, $T_p = 3.5$,
- C : $P = 0$, terminal equality constraint $\bar{\mathbf{x}}(t + T_p; \mathbf{x}(t), t) = \mathbf{0}$, $T_p = 3.5$.

Controller A is designed with the proposed method and has guaranteed stability. For controller B, there is no terminal constraint and the terminal states are not penalized. This controller corresponds to the general form of MPC usually used for applications. Closed-loop stability can only be achieved by tuning the horizon T_p . Here, $T_p = 3.5$ time-units is the shortest horizon determined by trial and error such that the closed-loop system is stable, (for $T_p = 1.5$ time-units, the closed-loop system is unstable). For controller C, the well-known terminal equality constraint $\bar{\mathbf{x}}(t + T_p; \mathbf{x}(t), t) = \mathbf{0}$ is used. Hence, a terminal state penalty does not make sense. Closed-loop stability is also guaranteed, if the optimization problem at time $t = 0$ is feasible.

For a total simulation of 10 time-units, the elapsed CPU times are shown in Table 1 for some different initial conditions, where the controllers A, B and C use the same optimization routine and the same integration algorithm with the same numerical parameters (optimality tolerance, feasibility tolerance, integration step, etc.). The symbol “*” in Table 1 indicates that the optimization problem is not feasible at time $t = 0$. It is clearly seen that controller A needs significantly less CPU time than controllers B and C. Here, controller B might be treated somewhat unfairly. In practice, one can use techniques such as blocking or

Table 1: Comparison of elapsed CPU time

Initial state		Elapsed CPU time (s)		
$x_1(0)$	$x_2(0)$	Contr. A	Contr. B	Contr. C
-0.683	-0.864	859	1521	*
-0.523	0.744	818	1492	*
0.808	0.121	615	1638	*
0.774	-0.222	570	1522	5729
0.292	0.228	820	1724	*
-0.08	-0.804	696	1544	5093

confounding to reduce on-line computation time. However, an important drawback of controller B is that stability can only be achieved by tuning parameters such as the prediction horizon. The big difficulty for controller C is the infeasibility of the optimization problem caused by the terminal equality constraint.

6 Conclusions

In this paper we proposed a quasi-infinite horizon nonlinear MPC scheme with guaranteed stability. The setup differs from the standard setup with quadratic objective functional only in that a terminal state penalty term ($\mathbf{x}(t + T_p)^T P \mathbf{x}(t + T_p)$) is added to the finite horizon objective functional and an additional terminal inequality constraint ($\mathbf{x}(t + T_p) \in \Omega$) has to be satisfied. These two terms do however not constitute additional tuning parameters that can be chosen freely, but have to be determined according to the procedure given in the paper. We have proven that asymptotic stability of the closed-loop system is guaranteed by the feasibility of the optimization problem at time $t = 0$, independent of the choice of the performance parameters Q and R in the quadratic objective functional. Thus, a separation between performance and stability issues is achieved in some sense. Terminal state penalty matrix P and terminal region Ω can be determined off-line in a straightforward manner, essentially involving the solution of a linear stabilization problem and a Lyapunov equation.

The main advantage of this scheme is its guaranteed asymptotic stability that does not need to assume that the optimal solution to the minimization problem is found at each time. In addition, the quasi-infinite horizon nonlinear MPC scheme is computationally more attractive than other nonlinear MPC schemes that also guarantee asymptotic stability (terminal equality constraint, infinite horizon). This was also demonstrated with the control of an unstable and constrained system. The computational burden involved with this scheme can be further reduced, when the open-loop system is asymptotically stable [2]. We have furthermore extended this approach to synthesize a *robust* nonlinear predictive controller. A preliminary discussion and first result can be found in [4].

References

- [1] R.R. Bitmead, M. Gevers, and V. Wertz. *Adaptive Optimal Control – The Thinking Man’s GPC*. Prentice Hall, New York, 1990.
- [2] H. Chen and F. Allgöwer. A quasi-infinite horizon nonlinear predictive control scheme for stable systems: Application to a CSTR. accepted for AD-CHEM’97.
- [3] H. Chen and F. Allgöwer. A quasi-infinite horizon nonlinear model predictive control scheme with guaranteed stability. accepted for Automatica, 1997.
- [4] H. Chen, C.W. Scherer, and F. Allgöwer. A game theoretic approach to nonlinear robust receding horizon control of constrained systems. accepted for ACC’97.
- [5] C.R. Cutler and B.L. Ramaker. Dynamic Matrix Control – A computer control algorithm. In *Proc. Joint Automatic Control Conference*, San Francisco, CA, 1980.
- [6] C.E. García, D.M. Prett, and M. Morari. Model Predictive Control: Theory and practice – A survey. *Automatica*, 25:335–347, 1989.
- [7] H. Genceli and M. Nikolaou. Design of robust constrained model-predictive controllers with Volterra series. *AIChE J.*, 41(9):2098–2107, 1995.
- [8] H.K. Khalil. *Nonlinear Systems*. Macmillan Publishing Company, New York, 1992.
- [9] J.H. Lee. Recent advances in model predictive control and other related areas. In *Chemical Process Control – CPC V*. Jan. 1996.
- [10] D.Q. Mayne and H. Michalska. Receding horizon control of nonlinear systems. *IEEE Trans. Automat. Contr.*, AC-35(7):814–824, 1990.
- [11] E.S. Meadows, M.A. Henson, J.W. Eaton, and J.B. Rawlings. Receding horizon control and discontinuous state feedback stabilization. *Int. J. Contr.*, 62(5):1217–1229, 1995.
- [12] H. Michalska and D.Q. Mayne. Receding horizon control of nonlinear systems without differentiability of the optimal value function. *Syst. Contr. Lett.*, 16:123–130, 1991.
- [13] H. Michalska and D.Q. Mayne. Robust receding horizon control of constrained nonlinear systems. *IEEE Trans. Automat. Contr.*, AC-38(11):1623–1633, 1993.
- [14] V. Nevistić and M. Morari. Constrained control of feedback-linearizable systems. In *Proc. 3rd European Control Conference ECC’95*, pages 1726–1731, Rome, 1995.
- [15] S.J. Qin and T.A. Badgwell. An overview of industrial model predictive control technology. In *5th International Conference on Chemical Process Control – CPC V*. Jan. 1996.
- [16] J.B. Rawlings and K.R. Muske. The stability of constrained receding horizon control. *IEEE Trans. Automat. Contr.*, AC-38(10):1512–1516, 1993.
- [17] J. Richalet. Industrial applications of model based predictive control. *Automatica*, 29(5):1251–1274, 1993.
- [18] J. Richalet, A. Rault, J.L. Testud, and J. Papon. Model predictive heuristic control: Application to industrial processes. *Automatica*, 14:413–428, 1978.
- [19] T.H. Yang and E. Polak. Moving horizon control of nonlinear systems with input saturation, disturbances and plant uncertainty. *Int. J. Contr.*, 58(4):875–903, 1993.