

Adaptive Nonlinear Model Predictive Control with Suboptimality and Stability Guarantees

Pontus Giselsson

Department of Automatic Control LTH

Lund University

Box 118, SE-221 00 Lund, Sweden

`pontusg@control.lth.se`

Abstract—Theory for Adaptive Nonlinear Model Predictive Control is developed based on the relaxed dynamic programming inequality. The adaptivity in the controller lies in the choice of control horizon. The control horizon is chosen such that a variation of the relaxed dynamic programming inequality holds for all time steps along the closed loop trajectory. This provides guarantees for asymptotic stability and closed loop suboptimality above a certain pre-specified level.

I. INTRODUCTION

Model Predictive Control (MPC) is recognized as a high performing control structure for complex systems. The foundation of MPC is to, between every applied control action, solve a finite horizon optimization problem that predicts the plant future and minimizes a certain cost functional based on the predictions. When new measurements become available to the controller, another optimization takes place. Over the last decades many successful applications with MPC has been implemented [2], [11].

Many researchers have presented different methods to ensure asymptotic stability for Model Predictive Controllers with fixed control horizon. Most of these methods utilize either a certain cost or constraint on the final state in the optimization horizon to ensure stability. This final cost or constraint is designed such that the optimal value function of the MPC cost is a Lyapunov function to the system, see [9] for a survey of different methods. A couple of MPC schemes with variable control horizon has been presented [10], [13], [14] and [15]. The main idea behind these methods is to vary the control horizon such that the final state of the horizon reaches a certain terminal set. In [10], this terminal set is predefined and in the series of papers [13], [14] and [15], the terminal set depends on the current state.

In the last couple of years, results has been established for asymptotic stability of Model Predictive Controllers without terminal cost or constraints, cf. [1], [4]. These papers show, under different types of controllability and detectability conditions that the system is stable for sufficiently large control horizons. However, bounds on the required horizon are not given. In [7] suboptimality bounds for Model Predictive Controllers are developed based on relaxed dynamic programming which is developed in [8]. Asymptotic stability is a direct consequence of these suboptimality bounds with the value function of the MPC cost as Lyapunov function.

This work has been extended in [5] where the suboptimality bounds together with certain controllability assumptions on the running cost allow for finding a minimal stabilizing control horizon for the class of systems satisfying the controllability assumptions. These controllability assumptions can, however, be a cumbersome task to verify for a given nonlinear system. Further, the control horizon required may vary for different parts of the state space. In an attempt utilize this, an adaptive MPC scheme, based on the same relaxed dynamic programming inequality, is presented in [12]. The control horizon is adapted, on-line, such that the relaxed dynamic inequality holds in every time step. The work in [12] ensures asymptotic stability of the system, and that the step to step suboptimality is above a certain pre-specified level. However, the pre-specified level of suboptimality is not guaranteed for the closed loop system from initial state to the zero-set of running cost.

In this work we consider Adaptive Model Predictive Control with cost functionals containing neither terminal constraints nor terminal cost. The adaptation in the Adaptive MPC lies in the choice of control horizon which is adapted such that an extended version of the relaxed dynamic programming inequality holds in each time step. The extension to the relaxed dynamic programming framework has two properties that differ from the original version. The first differing property is that it allows for time varying control horizons. The second differing property is that it gives less conservative sub-optimality estimates since, using a slack variable, conservatism from previous time steps are used to ease the conditions for future time steps. Ones the conditions of the extended relaxed dynamic programming inequality are satisfied, asymptotic stability is ensured. Further, contrary to what is the case in [12], closed loop performance above a pre-specified level, from initial state to the zero set of the running cost, is guaranteed. This paper has been developed in parallel with [3] in which similar ideas are used in a distributed MPC scheme.

The continuation of this paper is organized as follows. In Section II the MPC idea is presented and some preliminaries stated. In Section III the relaxed dynamic programming inequality is extended and altered to fit our context with varying control horizons. In Section IV the previously stated inequality is used as design tool for the choice of control

horizon in the MPC controller. In Section V two different schemes are presented such that asymptotic stability and suboptimality to a pre-defined degree is obtained. In Section VI a numerical example is presented and in Section VII we conclude the paper.

II. PROBLEM SETUP

The aim of this paper is to develop suboptimal control schemes for nonlinear discrete time systems of the form

$$x(t+1) = f(x(t), u(t)), \quad x(0) = x_0 \quad (1)$$

where $x(t) \in X$ and $u(t) \in U$ for $t \in \mathbb{N}_0$. Given a dynamical system (1) the ultimate objective is to find a feedback control law such that the following infinite horizon cost functional is minimized

$$J_\infty(x_0, u) = \sum_{t=0}^{\infty} \ell(x(t), u(t))$$

where the stage cost $\ell : X \times U \rightarrow \mathbb{R}_0^+$ and u is the sequence of applied control actions. The corresponding optimal value function is denoted

$$V_\infty(x_0) = \min_u J_\infty(x_0, u).$$

To find such a feedback control law the solution of the Hamilton-Jacobi-Bellman equation for the infinite horizon optimal control problem must be found. This is in general not a tractable problem. To circumvent this undesirable property the Model Predictive Control methodology is introduced. The idea behind Model Predictive Control is to truncate the original infinite horizon cost functional at some finite time, solve the resulting optimal control problem with finite horizon, apply the first control action in the open loop solution. This procedure is then repeated for every time step, where the truncated optimization problem is fed with the current state of the system. The iterative nature of the Model Predictive Controller results in a state feedback controller.

In this work we allow for time varying control horizons, leading to the following truncated cost functional in each time step t

$$J_{N(t)}(x(t), u) = \sum_{\tau=0}^{N(t)} \ell(x(t, \tau), u(\tau)). \quad (2)$$

where $x(t, 0) = x(t)$. Throughout this paper, the predicted state trajectory internal to the controller at time t is denoted $x(t, \tau)$, where $\tau = 0, \dots, N(t)$. The closed loop state at time t is denoted $x(t)$. In the MPC controller, the truncated cost functional (2) is minimized for every time step t with the system dynamics as equality constraints:

$$x(t, \tau+1) = f(x(t, \tau), u(\tau)), \quad x(t, 0) = x(t).$$

The corresponding value function is denoted

$$V_{N(t)}(x(t)) = \min_u J_{N(t)}(x(t), u).$$

In each time instant, t , a sequence of control actions, $u(t, \cdot)$, is optimized. Only the first of those actions, $u(t, 0)$, is applied

to the process before the whole optimization process is repeated in the following time step. The closed loop solution state trajectory is denoted

$$x(t+1) = f(x(t), u(t, 0)), \quad x(0) = x_0 \quad (3)$$

for $t \in \mathbb{N}_0$. The resulting infinite horizon cost for the closed loop system is denoted

$$V_\infty^{MPC}(x_0) = \sum_{t=0}^{\infty} \ell(x(t), u(t, 0))$$

The objective of this work is to create an adaptive MPC scheme in which the control horizon may vary between different time instants in order to guarantee a pre-specified bound on the suboptimality, i.e. a bound on the relation between $V_\infty^{MPC}(x_0)$ and $V_\infty(x_0)$.

The work in this paper is based on the relaxed dynamic programming inequality which provides a means to obtain suboptimality bounds for, e.g. an MPC-scheme. Before we are ready to state a proposition about this, we observe that $V_N(x(t)) \geq V_M(x(t))$ if $N \geq M$. This is true since neither terminal constraints nor terminal cost is present in the cost function. Further, due to the construction of the value function, we have that $V_N(x(t)) : X \rightarrow \mathbb{R}_0^+$. The following proposition is a slight variation of [7, Proposition 2.2] or [5, Proposition 2.4].

Proposition 1: Consider a closed loop trajectory $x(\cdot)$ according to (3) and suppose that there exist $\alpha \in (0, 1)$ such that

$$V_N(x(t)) \geq V_N(x(t+1)) + \alpha \ell(x(t), u(t, 0)) \quad (4)$$

holds for all $t \in \mathbb{N}_0$. Then

$$\alpha V_\infty^{MPC}(x(0)) \leq V_\infty(x(0)).$$

Proof. Summation of (4) over $t = 0, \dots, T$ gives

$$\alpha \sum_{t=0}^T \ell(x(t), u(t, 0)) \leq V_N(x(0)) - V_N(x(T+1)).$$

Since $V_N(x(T+1)) \geq 0$ and $V_N(x(0)) \leq V_\infty(x(0))$ we have that

$$\alpha \sum_{t=0}^T \ell(x(t), u(t, 0)) \leq V_\infty(x(0)).$$

The definition of $V_\infty^{MPC}(x(0))$ gives the desired result as $T \rightarrow \infty$. $\square$

Note that the conditions of the proposition are run-time conditions which provide an easy way to estimate the suboptimality of the closed loop trajectory.

Throughout this paper we assume that the optimal infinite horizon cost $V_\infty(x_0)$ is finite. Further, the running cost ℓ is assumed convex and to allow the system to stay in the origin at zero cost, $\ell(0, 0) = 0$. Finally, X is assumed to be control invariant, i.e., for all $x \in X$, $\exists u \in U$ s.t. $f(x, u) \in X$, to avoid feasibility problems.

III. NMPC ANALYSIS TOOLS

In this section the result in Proposition 1 is extended to include the case of time varying control horizons. Further, by introducing a slack variable $s(t)$, we relax the conditions in Proposition 1 which result in less conservative suboptimality bounds. The relaxation to Proposition 1 is given in the following theorem. A similar theorem is presented in [3].

Theorem 1: Consider a closed loop trajectory $x(\cdot)$ according to (3) and suppose that there exist $\alpha \in (0, 1)$ such that

$$V_{N(t)}(x(t)) \geq V_{N(t+1)}(x(t+1)) + \alpha \ell(x(t), u(t, 0)) + s(t) \quad (5)$$

where

$$s(t) = s(t-1) + \alpha \ell(x(t-1), u(t-1, 0)) + V_{N(t)}(x(t)) - V_{N(t-1)}(x(t-1)) \quad (6)$$

and $s(0) = 0$ hold for all $t \in \mathbb{N}_0$. Then

$$\alpha V_{\infty}^{MPC}(x(0)) \leq V_{\infty}(x(0)).$$

Proof. Induction of (6) over t gives at $t = T$

$$\begin{aligned} s(T) &= s(T-1) + \alpha \ell(x(T-1), u(T-1, 0)) + \\ &\quad + V_{N(T)}(x(T)) - V_{N(T-1)}(x(T-1)) \\ &= \dots = \alpha \sum_{t=0}^{T-1} \ell(x(t), u(t, 0)) + \\ &\quad + V_{N(T)}(x(T)) - V_{N(0)}(x(0)). \end{aligned}$$

Insert this into (5) gives

$$\begin{aligned} \alpha \sum_{t=0}^T \ell(x(t), u(t, 0)) &\leq V_{N(0)}(x(0)) - V_{N(T+1)}(x(T+1)) \\ &\leq V_{N(0)}(x(0)) \leq V_{\infty}(x(0)) \end{aligned}$$

This gives, as $T \rightarrow \infty$

$$\alpha V_{\infty}^{MPC}(x(0)) = \alpha \sum_{t=0}^{\infty} \ell(x(t), u(t, 0)) \leq V_{\infty}(x(0)).$$

This completes the proof. $\square$

Remark 1: The difference between Theorem 1 and Proposition 1, besides the fact that we allow for variable time horizons, is the introduced slack variable $s(t)$. This slack variable sums the slack in the inequalities for previous time steps, giving an easier inequality to fulfil in the current time step. Since $s(t) \leq 0$ for all t , the inequality (5) is easier to fulfil for given α than the inequality in Proposition 1.

Our next objective is to prove asymptotic stability under the conditions of Theorem 1. With asymptotic stability in this context, we mean that for the closed loop system $\|x(t)\|_2^2 \rightarrow 0$ as $t \rightarrow \infty$. Before we are ready to state the theorem, which is equivalent to [3, Theorem 2], the following assumption is needed.

Assumption 1: Assume that there exist a $\beta > 0$ such that $\min_u \ell(x, u) \geq \beta \|x\|_2^2$.

Theorem 2: Consider a closed loop trajectory $x(\cdot)$ according to (3) and suppose that Assumption 1 holds and that

$$V_{\infty}^{MPC}(x(0)) \leq M \quad (7)$$

where M is a finite positive real number. Then $\|x(t)\|_2^2 \rightarrow 0$ as $t \rightarrow \infty$.

Proof. We show this by a contradiction argument. We have that

$$V_{\infty}^{MPC}(x(0)) = \sum_{t=0}^{\infty} \ell(x(t), u(t, 0)) \leq M \quad (8)$$

where M is a finite positive real number. Assume that $\|x(t)\|_2^2 \not\rightarrow 0$ as $t \rightarrow \infty$, then there is an $\epsilon > 0$ and a $T \geq 0$ such that $\|x(t)\|_2^2 \geq \epsilon$ for all $t \geq T$. Further

$$\sum_{t=0}^{\infty} \ell(x(t), u(t, 0)) \geq \sum_{t=T}^{\infty} \beta \|x(t)\|_2^2 \geq \beta \epsilon \sum_{t=T}^{\infty} 1 \quad (9)$$

which is unbounded. Thus by contradiction the assertion holds. $\square$

Remark 2: Note that if the conditions of Theorem 1 hold then the conditions of Theorem 2 also hold with $M = V_{\infty}(x(0))/\alpha$ which is finite by assumption.

As previously discussed, the conditions in Theorem 1 are run-time conditions. They can be used as an analysis tool for a NMPC system to, on-line, estimate the performance of the closed loop system. If the conditions fail we cannot deduce anything about stability nor suboptimality in this framework. The objective of the next section is to use the analysis tools in this section as design tools to adaptively choose control horizon $N(t)$ such that the conditions of Theorem 1 hold for all t and for a pre-specified level of suboptimality.

IV. ADAPTIVE MODEL PREDICTIVE CONTROL

The conditions of Theorem 1 cannot directly be used to adaptively choose control horizon in an adaptive MPC context. The reason is that information about the next step value function, $V_{N(t+1)}(x(t+1))$, is not available at time t . The other terms, $V_{N(t)}(x(t))$ and $\ell(x(t), u(t, 0))$ are byproducts from the optimization problem to be solved when calculating the control action, and $s(t)$ contains previously calculated terms. Thus, if an upper bound, $\bar{V}_{N(t+1)}(x(t+1))$, can be calculated at time t such that

$$\bar{V}_{N(t+1)}(x(t+1)) \geq V_{N(t+1)}(x(t+1)) \quad (10)$$

then the conditions of Theorem 1 can be changed to

$$V_{N(t)}(x(t)) \geq \bar{V}_{N(t+1)}(x(t+1)) + \alpha \ell(x(t), u(t, 0)) + s(t)$$

and

$$\begin{aligned} s(t) &= s(t-1) + \alpha \ell(x(t-1), u(t-1, 0)) \\ &\quad + \bar{V}_{N(t)}(x(t)) - V_{N(t-1)}(x(t-1)) \end{aligned}$$

respectively. All terms in these conditions are known at time t which allows for adaptation of the control horizon $N(t)$ such that the conditions hold. The use of upper bounds will give more conservative results than if the actual optimal value function values was used. Most of this conservatism can,

however, be used to ease the conditions for later time steps by incorporating this conservatism in the slack variable $s(t)$ as in the following theorem.

Theorem 3: Consider a closed loop trajectory $x(\cdot)$ according to (3) and suppose that for a pre-specified $\alpha \in (0, 1)$ we can find control horizons, $N(t) \in \mathbb{N}_1$ such that

$$V_{N(t)}(x(t)) \geq \bar{V}_{N(t+1)}(x(t+1)) + \alpha \ell(x(t), u(t, 0)) + s(t) \quad (11)$$

where

$$s(t) = s(t-1) + \alpha \ell(x(t-1), u(t-1, 0)) + \bar{V}_{N(t)}(x(t)) - \bar{V}_{N(t-1)}(x(t-1)) \quad (12)$$

and

$$s(1) = \alpha \ell(x(0), u(0, 0)) + \bar{V}_{N(1)}(x(1)) - V_{N(0)}(x(0)) \quad (13)$$

and $s(0) = 0$ hold for all $t \in \mathbb{N}_0$. Then

$$\alpha V_{\infty}^{MPC}(x(0)) \leq V_{\infty}(x(0))$$

and $\|x(t)\|_2^2 \rightarrow 0$ as $t \rightarrow \infty$.

Proof. Induction of (12) over t gives at $t = T$

$$\begin{aligned} s(T) &= s(T-1) + \alpha \ell(x(T-1), u(T-1, 0)) + \bar{V}_{N(T)}(x(T)) - \bar{V}_{N(T-1)}(x(T-1)) \\ &= \dots = s(1) + \alpha \sum_{t=1}^{T-1} \ell(x(t), u(t, 0)) + \bar{V}_{N(T)}(x(T)) - \bar{V}_{N(1)}(x(1)) \\ &= \alpha \sum_{t=0}^{T-1} \ell(x(t), u(t, 0)) + \bar{V}_{N(T)}(x(T)) - V_{N(0)}(x(0)) \end{aligned}$$

Insert this into (11) gives

$$\begin{aligned} \alpha \sum_{t=0}^T \ell(x(t), u(t, 0)) &\leq V_{N(0)}(x(0)) - \bar{V}_{N(T+1)}(x(T+1)) + \\ &\quad + V_{N(T)}(x(T)) - \bar{V}_{N(T)}(x(T)) \\ &\leq V_{N(0)}(x(0)) - \bar{V}_{N(T+1)}(x(T+1)) \\ &\leq V_{N(0)}(x(0)) \leq V_{\infty}(x(0)) \end{aligned}$$

This gives, as $T \rightarrow \infty$

$$\alpha V_{\infty}^{MPC}(x(0)) = \alpha \sum_{t=0}^{\infty} \ell(x(t), u(t, 0)) \leq V_{\infty}(x(0)),$$

which proves the assertion about suboptimality.

Theorem 2 proves that $\|x(t)\|_2^2 \rightarrow 0$ as $t \rightarrow \infty$ since $V_{\infty}^{MPC}(x(0))$ is finite.

This completes the proof. $\square$

Remark 3: The slack variable $s(t)$ in Theorem 3 consists of two parts, the unused slack from the previous time step:

$$s(t-1) + \alpha \ell(x(t-1), u(t-1, 0)) + \bar{V}_{N(t)}(x(t)) - V_{N(t-1)}(x(t-1))$$

and the conservatism due to the upper bound of the value function in the inequality two time steps back:

$$V_{N(t-1)}(x(t-1)) - \bar{V}_{N(t-1)}(x(t-1)).$$

When these parts are summed up, we get the slack variable as described in Theorem 3.

Remark 4: We ensure that

$$\alpha \sum_{t=0}^{\infty} \ell(x(t), u(t, 0)) \leq V_{\infty}(x(0))$$

by ensuring that

$$\alpha \sum_{t=0}^{\infty} \ell(x(t), u(t, 0)) \leq V_{N(0)}(x(0)).$$

Since we know that $\ell(x(t), u(t, 0)) \rightarrow 0$ as $t \rightarrow \infty$ and by assuming that $\bar{V}_{N(t)}(x(t)) \rightarrow 0$ as $t \rightarrow \infty$ the slack variable converges $s(t) \rightarrow s$ as $t \rightarrow \infty$. Thus, we have the following relationship

$$\alpha \sum_{t=0}^{\infty} \ell(x(t), u(t, 0)) = V_{N(0)}(x(0)) + s$$

where s is a measure on the conservatism with respect to what is needed to ensure the conditions of Theorem 3.

Given the following assumption, a bound on the control horizon needed to ensure the conditions of Theorem 3 is directly given.

Assumption 2: Assume that for a pre-specified $\alpha \in (0, 1)$, a finite number $N_0 \in \mathbb{N}_1$ is known such that

$$V_{N_0}(x) \geq V_{N_0}(f(x, u)) + \alpha \ell(x, u)$$

holds for all $x \in X$ where u is the minimizing control action, i.e. $u = \arg \min V_{N_0-1}(f(x, u)) + \ell(x, u)$.

Remark 5: Consult [5] for literature that address the matter of finding N_0 for given α that satisfies Assumption 2.

Assumption 2 gives that the longest control horizon necessary to satisfy the conditions of Theorem 3 is N_0 since $s(t) \leq 0$ for all $t \in \mathbb{N}_0$.

Two-phase adaptive MPC

An extension to Theorem 3, in which a two-phase strategy is used, is presented here. If sub-optimality of at least α is desired, Assumption 2 can be evaluated to find the corresponding N_0 . Further, for some $N_{\beta} < N_0$ the corresponding $0 < \beta < \alpha$ can be evaluated that satisfies Assumption 2. In the first phase the adaptive MPC runs with α as suboptimality and N_0 as upper bound to the control horizon necessary. At some time $T > 0$, if a certain condition holds, the second phase starts where the adaptive MPC runs with β as suboptimality and N_{β} as upper bound to the control horizon. The following theorem shows that the original performance objective, $\alpha V_{\infty}^{MPC}(x(0)) \leq V_{\infty}(x(0))$, is achieved.

Theorem 4: Assume that Assumption 2 holds for a pre-specified α with horizon N_0 and for some $0 < \beta < \alpha$ with horizon $N_{\beta} < N_0$. Further assume that phase two, with β and N_{β} , starts at some $T > 0$ if

$$\alpha \sum_{t=0}^{T-1} \ell(x(t), u(t, 0)) \leq V_{N(0)}(x(0)) - \frac{\alpha}{\beta} V_{N_{\beta}}(x(T)) \quad (14)$$

where $N(0) \leq N_0$. Then the performance bound $\alpha V_{\infty}^{MPC}(x(0)) \leq V_{\infty}(x(0))$ still hold and for all $t \geq T$ the longest control horizon needed is N_{β} .

Proof. Theorem 3 and Assumption 2 gives that for any switching time $T \geq 0$ one obtains

$$\beta \sum_{t=T}^{\infty} \ell(x(t), u(t, 0)) \leq V_{N_\beta}(x(T)).$$

with control horizon $N \leq N_\beta$ for all $t \geq T$. Further

$$\begin{aligned} \alpha \sum_{t=0}^{T-1} \ell(x(t), u(t, 0)) &\leq V_{N(0)}(x(0)) - \frac{\alpha}{\beta} V_{N_\beta}(x(T)) \\ &\leq V_{N(0)}(x(0)) - \alpha \sum_{t=T}^{\infty} \ell(x(t), u(t, 0)) \end{aligned}$$

which gives

$$\alpha \sum_{t=0}^{\infty} \ell(x(t), u(t, 0)) \leq V_{N(0)}(x(0)) \leq V_\infty(x(0)).$$

This completes the proof. $\square$

The purpose of the this theorem is to reduce the computational complexity of the adaptation scheme by shortening the control horizon needed for the upper bound.

V. ADAPTIVE NMPC SCHEMES

A simple adaptation scheme that finds the control horizon length necessary to satisfy the conditions of Theorem 3 is presented in this section.

A. General adaptation scheme

The following adaptation scheme is used to ensure the conditions of Theorem 3.

Scheme
1) Calculate $V_{N(t)}(x(t))$ 2) Calculate $\bar{V}_{N(t+1)}(f(x(t), u(t, 0)))$ 3) While the conditions of Theorem 3 do not hold (or $N(t) \geq N_0$ [under Assumption 2]) • Set $N(t) \leftarrow N(t) + 1$ • Calculate $V_{N(t)}(x(t))$ • Calculate $\bar{V}_{N(t+1)}(f(x(t), u(t, 0)))$ 4) Apply $u(t, 0)$ 5) Set $N(t+1) \leftarrow \max(N(t) - 1, N_{min})$ 6) Set $t \leftarrow t + 1$ and go to 1)

With this scheme, the largest decrease in control horizon between two consecutive steps is one. To use the proposed scheme for Theorem 4, the condition (14) need to be checked in every sample. If the condition holds for some $T > 0$, start phase two with β sub-optimality degree and N_β as upper bound to the length of the control horizon.

The only term in the conditions, (11), that is not available after the ordinary MPC optimization, is the upper bound to the value function in the next step. This upper bound can, based on different assumptions, be calculated in various ways. We present two schemes which follow the general scheme but differ in their respective ways to calculate the upper bound.

B. Scheme 1

Calculation of the upper bounds in this scheme is based on Assumption 2. The lowest possible upper bound is, of course, $V_{N_0}(f(x(t), u(t, 0)))$. The computational time of calculating that value is the same as calculating $V_{N_0}(x(t))$. This would make the adaptivity superfluous since it would not be possible to reduce the computational complexity compared to using traditional MPC with fixed control horizon, $N = N_0$. However, a feasible control horizon from $f(x(t), u(t, 0))$ over $N - 1$ time steps is available from the calculation of $V_N(x(t))$. An upper bound for the value of the remaining $N_0 - N + 1$ time steps can be obtained by calculating $V_{N_0-N+1}(x(t, N))$. Thus an upper bound is achieved in the following way:

$$\begin{aligned} \bar{V}_{N(t+1)}(x(t+1)) &= V_N(x(t)) - \ell(x(t), u(t, 0)) - \\ &\quad - \ell(x(t, N), u(t, N)) + V_{N_0-N+1}(x(t, N)) \end{aligned}$$

where all terms except the last one are known after computing $V_N(x(t))$.

In the scheme N_{min} needs to be chosen. This value should be chosen such that $N_{min} \geq N_0/2$ to spend more computational effort on the optimization problem needed for control than on the optimization problem needed for the upper bound.

C. Scheme 2

The method to calculate the upper bounds in this second scheme does not rely on any assumptions on the control horizon needed to ensure the conditions in the following time step. To calculate an upper bound to the next step value function, an upper bound to the infinite horizon value function is needed, i.e. $\bar{V}_\infty(f(x(t), u(t, 0)))$. Such an upper bound can be obtained by solving the usual optimal control problem with the additional constraint that the final state in the horizon should be in the origin. The corresponding value function is denoted $V_N^0(x)$. Similar to in the first scheme we get the following upper bound

$$\begin{aligned} \bar{V}_\infty(x(t+1)) &= V_N(x(t)) - \ell(x(t), u(t, 0)) - \\ &\quad - \ell(x(t, N), u(t, N)) + V_N^0(x(t, N)). \end{aligned}$$

The first $N - 1$ steps of the upper bound trajectory is decided by the MPC-optimization. The remaining trajectory is chosen to have control horizon N with final constraints in the origin.

VI. NUMERICAL EXAMPLE

A cart example is presented here to numerically show how the adaptive NMPC scheme performs. The cart moves in two dimensions where each direction of motion is modeled as a discrete time double integrator

$$x(t+1) = \begin{pmatrix} 1 & h & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & h \\ 0 & 0 & 0 & 1 \end{pmatrix} x(t) + \begin{pmatrix} h^2/2 \\ h \\ h^2/2 \\ h \end{pmatrix} u(t).$$

The control signals $u \in \mathbb{R}^2$ and the states $x \in X$ where

$$X = \{x \in \mathbb{R}^4 \mid x_1 \geq -3, x_3 \geq -3, (x_1 + x_3) \geq -1, \\ -30 \leq x_2 \leq 30, -30 \leq x_4 \leq 30\}$$

where the subscript corresponds to its location in the state vector. The stage cost considered is

$$\ell(x, u) = x^T x + u^T u$$

The control objective is to control the system from its initial position, $x(0) = (10, -30, 10, -10)^T$, to the origin while minimizing the stated stage along the closed loop trajectory.

In table I numerical results from simulations based on the schemes presented in the previous section are presented and compared to other schemes. All simulations are performed to achieve the suboptimality bound $\alpha = 0.8$. For this choice of α , Assumption 2 is satisfied for $N_0 = 20$.

MPC scheme comparisons							
Scheme	N_0	β	N_β	N_{min}	$\bar{N}$	α_{calc}	$\bar{c}(10^{-3}s)$
1	20	0.3	10	5	5.6	0.987	0.24
1	20	-	-	10	10.4	0.989	0.72
2	-	-	-	5	6.7	0.993	2.1
const N	20	-	-	-	20	0.994	1.0

TABLE I

RESULTS FROM EXPERIMENTS WITH DIFFERENT NMPC SCHEMES

In the first row, results for scheme 1 are presented, where Assumption 2 is known to hold for $\alpha = 0.8$, $N_0 = 20$ and for $\beta = 0.3$, $N_\beta = 10$. The scheme is based on Theorem 3 and Theorem 4. The second row contain results when running the scheme based on Assumption 2 with $\alpha = 0.8$, $N_0 = 20$. In the third row results for scheme 2, i.e. without assumptions on the control horizon, are presented. The final row contain the results obtained when a fixed horizon controller is used. To guarantee the desired sub-optimality, Assumption 2 is assumed to hold for $\alpha = 0.8$, $N_0 = 20$.

Note that the scheme behind row 1 require most a priori knowledge of the system. The scheme behind row 2 and the one with constant control horizon require the same knowledge. The scheme in row 3 is the scheme that guarantees the desired sub-optimality with least a priori information of the system.

The last three columns of the table present the results. The column with $\bar{N}$ contain the mean length of the control horizon. The mean is calculated using the first 100 time samples. One can note that the mean length is close to the minimum allowed control horizon in all schemes. The column with α_{calc} specifies the performance for each scheme. The performance is very similar for all schemes and very close to optimal performance. The final column might be the most interesting one since it contains the mean calculation time of the optimization problems needed to be solved before the control action can be applied. The scheme with most a priori information has the shortest mean computation time, less than one fourth compared to if constant control horizon was used. The adaptive scheme in row 2 also has shorter

mean computational time than the one with constant horizon. They are based on the same assumptions, which hints that this adaptive MPC approach is more computationally efficient than the traditional fixed horizon approach, at least in this example. The scheme behind the third row is based on less assumptions and more computational time is required to ensure the same performance.

VII. CONCLUSIONS

We have presented theory for Adaptive Model Predictive Control based on the relaxed dynamic programming inequality. Asymptotic stability and closed loop performance above a certain pre-specified level is guaranteed. Numerical examples have shown that the computational complexity of the adaptation scheme is, on average, lower than for a traditional MPC scheme with fixed control horizon that ensures the same performance.

REFERENCES

- [1] Jadbabaie A. and J. Hauser. On the stability of receding horizon control with a general terminal cost. *IEEE Transactions on Automatic Control*, 50:674– 678, 2005.
- [2] Eduardo F. Camacho and Carlos A. Bordons. *Model Predictive Control in the Process Industry*. Springer-Verlag New York, Inc., Secaucus, NJ, USA, 1997.
- [3] Pontus Giselsson and Anders Rantzer. Distributed model predictive control with suboptimality and stability guarantees. In *Proceedings of the 49th Conference on Decision and Control*, 2010. To appear.
- [4] G. Grimm, M.J. Messina, S.E. Tuna, and A.R. Teel. Model predictive control: for want of a local control lyapunov function, all is not lost. *IEEE Transactions on Automatic Control*, 50:546 – 558, 2005.
- [5] Lars Grüne. Analysis and design of unconstrained nonlinear mpc schemes for finite and infinite dimensional systems. *SIAM Journal on Control and Optimization*, 48:1206–1228, 2009.
- [6] Lars Grüne and Jürgen Pannek. Practical nmpe suboptimality estimates along trajectories. *Systems & Control Letters*, 58:161 – 168, 2009.
- [7] Lars Grüne and Anders Rantzer. On the infinite horizon performance of receding horizon controllers. *IEEE Transactions on Automatic Control*, 53:2100–2111, 2008.
- [8] Bo Lincoln and Anders Rantzer. Relaxing dynamic programming. *IEEE Transactions on Automatic Control*, 51:1249–1260, August 2006.
- [9] D. Q. Mayne, J. B. Rawlings, C. V. Rao, and P. O. M. Scokaert. Constrained model predictive control: Stability and optimality. *Automatica*, 36:789 – 814, 2000.
- [10] H. Michalska and D.Q. Mayne. Robust receding horizon control of constrained nonlinear systems. *IEEE Transactions on Automatic Control*, 38:1623 –1633, 1993.
- [11] Manfred Morari and Jay H. Lee. Model predictive control: past, present and future. *Computers and Chemical Engineering*, 23:667–682, 1999.
- [12] Jürgen Pannek. Adaptive nonlinear receding horizon control schemes with guaranteed degree of suboptimality. *System & Control Letters*, 2010. Submitted. Available at <http://www.math.uni-bayreuth.de/~jpannek/>.
- [13] E. Polak and T. H. Yang. Moving horizon control of linear systems with input saturation and plant uncertainty part 1. robustness. *International Journal of Control*, 58:613 – 638, 1993.
- [14] E. Polak and T. H. Yang. Moving horizon control of linear systems with input saturation and plant uncertainty part 2. disturbance rejection and tracking. *International Journal of Control*, 58:639 – 663, 1993.
- [15] T. H. Yang and E. Polak. Moving horizon control of nonlinear systems with input saturation, disturbances and plant uncertainty. *International Journal of Control*, 58:875 – 903, 1993.