

Document Version

Final published version

Licence

CC BY

Citation (APA)

Airaldi, F., De Schutter, B., & Dabiri, A. (2025). Reinforcement Learning With Model Predictive Control for Highway Ramp Metering. *IEEE Transactions on Intelligent Transportation Systems*, 26(5), 5988-6004.
<https://doi.org/10.1109/TITS.2025.3549227>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

In case the licence states "Dutch Copyright Act (Article 25fa)", this publication was made available Green Open Access via the TU Delft Institutional Repository pursuant to Dutch Copyright Act (Article 25fa, the Taverne amendment). This provision does not affect copyright ownership.
Unless copyright is transferred by contract or statute, it remains with the copyright holder.

Sharing and reuse

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

Reinforcement Learning With Model Predictive Control for Highway Ramp Metering

Filippo Airaldi¹, Bart De Schutter², *Fellow, IEEE*, and Azita Dabiri³

Abstract—In the backdrop of an increasingly pressing need for effective urban and highway transportation systems, this work explores the synergy between model-based and learning-based strategies to enhance traffic flow management by use of an innovative approach to the problem of ramp metering control that embeds Reinforcement Learning (RL) techniques within the Model Predictive Control (MPC) framework. The control problem is formulated as an RL task by crafting a suitable stage cost function that is representative of the traffic conditions, variability in the control action, and violations of the constraint on the maximum number of vehicles in queue. An MPC-based RL approach, which leverages the MPC optimal problem as a function approximation for the RL algorithm, is proposed to learn to efficiently control an on-ramp and satisfy its constraints despite uncertainties in the system model and variable demands. Simulations are performed on a benchmark small-scale highway network to compare the proposed methodology against other state-of-the-art control approaches. Results show that, starting from an MPC controller that has an imprecise model and is poorly tuned, the proposed methodology is able to effectively learn to improve the control policy such that congestion in the network is reduced and constraints are satisfied, yielding an improved performance that is superior to the other controllers.

Index Terms—Ramp metering, reinforcement learning, model predictive control.

I. INTRODUCTION

OVER the past decades, the need for urban and highway transportation has become significantly relevant to our society. Extended travel times (primarily spent in traffic jams), impact on stress levels and pollution are but a few examples of the negative effects due to high demands for mobility that traffic engineers are facing nowadays [1], [2]. In particular, the necessity for advanced intelligent traffic control strategies is arising as the construction of new infrastructure and upgrades to the existing one become increasingly unsustainable, both in economical, environmental, and societal terms [3], [4].

Ramp metering (RM) control has been proven to be an effective tool in regulating the traffic flow entering the network at on-ramps [5]. It consists in modulating the flow at an on-ramp via traffic lights with appropriate timings of red,

green, and amber lights. Earlier approaches were fixed-time, i.e., the timings were calculated offline with the help of historic traffic data [6]. More efficient strategies involve real-time measurements of the traffic conditions, enabling online computation of the timings in a traffic-responsive way. In these cases, the underlying algorithms range from simpler feedback strategies [7], [8], [9] to more complex architectures, such as Model Predictive Control (MPC) [10], [11], [12] and Reinforcement Learning (RL) [13], [14]. One of the primary research challenges in RM today is the design of strategies that can control the entering flow in a way that effectively balances congestion in the freeway and the queue length at the on-ramp, two often conflicting goals [15], [16]. Most strategies deployed on actual ramps are either feedback-based local policies, which are inherently myopic, or model-based policies and, as such, are vulnerable to model inaccuracies and thus require precise identification of traffic parameters, as well as an attentive assessment of their efficacy [17]. Conversely, data-driven approaches, which are one of the main focuses of current research, typically necessitate extensive datasets of real or synthetic traffic data for training.

In modern control literature, MPC is an established control paradigm whose appeal can be attributed to its strong theoretical foundations, a diverse array of applications, as well as its remarkable capability to manage multivariate and constrained systems [18]. Its versatility further allows for alternative formulations that can systematically handle disturbances affecting the system, e.g., in a robust [19] or a stochastic [20] manner. Unlike other black-box or purely data-driven approaches, MPC-based control strategies often maintain a high degree of explainability and interpretability, which is especially favourable in settings where constraint satisfaction is critical. These properties are in part thanks to the prediction model that the MPC scheme contains, which allows it to forecast the dynamical evolution of state trajectories along the time horizon, and which is often designed by domain experts.

Nonetheless, it is well known that the performance of such predictive controllers is heavily bound to the quality of the prediction model, which calls for a high level of expertise during the system identification phase. In an attempt to counter this limitation, driven by the recent surge in advancements in Machine Learning (ML) algorithms and facilitated by the growing computational and sensing capabilities of digital systems, there is currently a substantial body of research dedicated to learning-based control methodologies that offer

Received 15 November 2023; revised 21 May 2024, 24 October 2024, and 5 February 2025; accepted 2 March 2025. Date of publication 17 March 2025; date of current version 5 May 2025. This work was supported in part by European Research Council (ERC) through European Union's Horizon 2020 Research and Innovation Programme under Grant 101018826-CLariNet. The Associate Editor for this article was P. Ye. (*Corresponding author: Filippo Airaldi.*)

The authors are with Delft Center for Systems and Control, Delft University of Technology, 2628 CD Delft, The Netherlands (e-mail: f.airaldi@tudelft.nl; b.deschutter@tudelft.nl; a.dabiri@tudelft.nl).

Digital Object Identifier 10.1109/TITS.2025.3549227

ways of leveraging open- or closed-loop data to, e.g., synthesize models and controllers, or enhance existing ones.

Amongst these methodologies, one of the most notable approaches is Reinforcement Learning (RL), a paradigm of ML that has gained substantial attention due to its advancements and successes. It provides a framework that allows to train an agent to interact with an unknown environment and to make decisions to maximise rewards or minimise costs [21]. The popularity in the contemporary ML community tends to favour its model-free variants that rely solely on observed state transitions and stage cost realisations to increase the performance of the control policy. Early RL focused on tabular approaches that enjoy simplicity, some degree of interpretability, and convergence guarantees, but they are only applicable to small and simple environments [22], and are often limited in handling larger state spaces, exhibit slow convergence, and struggle to generalize. As the field progressed, researchers recognized the potential of function approximators and, in particular, of Neural Networks (NNs) [23], leading to the development of Deep RL (DRL). However, while DRL has demonstrated great efficacy [24], its NN-based policies generally lack interpretability, and the integration of prior knowledge (e.g., from domain experts) is still difficult. Thus, providing formal guarantees on properties relevant to a controller, such as stability and constraint satisfaction, in the context of NN-enhanced control architectures is one of today's challenges in the field. Additionally, existing model-free approaches are usually plagued by sample inefficiency and poor real-world performance, since they are often trained offline in approximate simulators and require large amounts of data for meaningful converge [25].

At the same time, on the other side of the coin, the control community is also dedicating efforts in integrating data-driven techniques with control systems in more structured ways, fostering methodologies that aim to learn various components of a controller directly from data [26], [27]. The goals include, e.g., reducing model uncertainties and the impact of disturbances, improving closed-loop performance and battling conservativeness, or promoting robust control policies with guarantees on constraint satisfaction. Notably, learning-based MPC methodologies [28] have been devised to, e.g., leverage Gaussian Process regression to learn unmodelled nonlinear dynamics in the prediction model via state transition observations [29], or use Bayesian Optimisation to automatically tune and optimise the controller hyperparameters [30]. In the literature, MPC has been successfully combined also with RL, e.g., as a safety filter for the learning-based agents (also called *shielding*) [26], [31], as a model reference and baseline control provider [32], [33], and to solve mixed-logical control problems more efficiently [34].

Recently, [35] proposed and justified the use of MPC as function approximation of the optimal policy in model-based RL. Fig. 1 depicts a schematic overview of this architecture. In such a scheme, the MPC controller's optimisation problem acts both as policy provider and value function approximation of the RL task. Concurrently, the learning algorithm (such as Q-learning and stochastic/deterministic policy gradient) is tasked with adjusting the parametrisation of the controller

in an effort to directly or indirectly discover the optimal policy, thus improving closed-loop performance in a data-driven fashion. In this way, despite the presence of mismatches between the prediction model and the real system, the MPC control scheme is able to learn and deliver at convergence the optimal policy and value functions of the underlying RL task, granted the MPC parametrisation is rich enough. Unlike the more traditional pipeline, where a system identification phase (which selects model parameters that match the system's behaviour with observed data by, e.g., minimizing the predicted mean-squared-error) is followed by the controller design (possibly, iteratively), this approach directly tunes the parametrisation to maximize performance, irrespective of the controller's predictive accuracy. Such a rationale is often called Identification for Control, and advocates that the best parametrisation for control shall not yield the best predictions, but rather the best performance on the true system [30], [35], [36]. In contrast to NN-based RL, the key advantages of the approach adopted in this paper are multiple:

- The MPC scheme at its core allows to easily integrate prior information on the system in the form of an expert-based, possibly imperfect, prediction model.
- MPC-based agents are in general more amenable to analysis and certification in terms of stability and constraint satisfaction, a benefit supported by the extensive body of literature on MPC [18]. Although challenging and beyond the scope of the proposed approach, it is possible to ensure that theoretical properties, such as stability, recursive feasibility and constraint guarantees, are preserved during the RL learning process [36], [37]. On the other hand, NNs are notoriously black-box, lack interpretability, and often provide no guarantees prior to convergence.
- MPC controllers can also take constraints on states and actions into account in an explicit and structured way, which NNs are generally incapable of or can do poorly.

Compared to traditional non-learning MPC formulations, the proposed approach enjoys the following benefits:

- The need for extensive open-loop data collection and system identification phase is eliminated, as the MPC parametrisation is automatically adjusted based on closed-loop data to improve performance, rather than prediction accuracy. This bypasses the issue of model inaccuracies and endows the scheme with online adaptation capabilities.
- By appropriately penalising constraint violations, the RL algorithm can guide the learning process to produce an MPC policy that satisfies system constraints.

All combined, these advantages can foster a safer training where the agent may enjoy a more stable and constraint-abiding learning process, opening up in the future the possibility of training directly in the real world [25], and contributing to the suitability of this framework for the RM control problem. In particular, the proposed MPC-RL method automatically learns from data a policy that optimally balances congestion and queue lengths, without requiring further manual tuning. This bypasses the issue of model inaccuracies in traffic predictions, which are often the primary

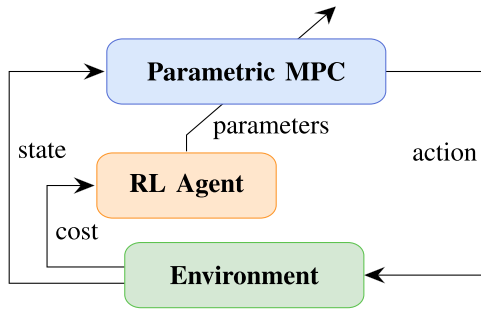


Fig. 1. Diagram of the MPC-based RL architecture.

concern in non-myopic control strategies. It also embeds adaptivity in the RM strategy in the face of stochasticities (e.g., demand forecasts or the behaviour of drivers can vary significantly across different cities or countries, meaning that system identification performed in one region may not generalise well to another). Besides, we stress again that improving the accuracy of the prediction model may not necessarily result in a better closed-loop performance. Hence, this framework allows for improvements to the control policy without focusing on nor needing an accurate prediction model or identification as a separate step in the design process. Additionally, the approach strengthens constraint satisfaction, which is crucial in real-world applications like RM control to avoid critical congested scenarios.

This paper contributes to the state of the art by proposing a novel approach, which embeds the MPC framework within an RL algorithm, to tackle the highway RM control problem. Specifically:

- 1) To the best of the authors' knowledge, an MPC-based RL method inspired by [35] is proposed for the first time to solve the RM control problem. The method combines a differentiable parametrised MPC scheme, which acts as policy provider and function approximation, with an online second-order Q-learning algorithm. The Q-learning algorithm adjusts the MPC parametrisation via gradient updates based only on observed closed-loop data, allowing to automatically learn a congestion-free and constraint-abiding control policy. As explained earlier, this approach overcomes the issue of model inaccuracies and uncertainties in the MPC and in the context of RM.
- 2) The proposed RM strategy is then implemented and tested on a benchmark highway stretch, where it is trained in an online fashion to simultaneously avoid bottlenecks and excessive queue lengths at the on-ramp, incorporating penalties for constraint violations directly into the MPC objective and the RL reward function. Simulation results underscore the effectiveness of the MPC-RL approach. Notably, starting from a poorly tuned controller with an imprecise prediction model, the proposed approach demonstrates quick learning in reducing congestion and satisfying the queue constraint. These results are further validated with a comparison with other state-of-the-art traffic controllers, both model-based and learning-based. Our method empirically outperforms them in performance and sample-efficiency.

The paper is organised as follows. Section II summarises relevant works in MPC, RL, and their combination in RM applications. Background on modelling and controlling traffic networks via MPC is provided in Section III. Section IV presents the proposed methodology, showing how MPC can be combined with RL, as well as explaining the MPC-based RL algorithm itself. The method is then applied to a numerical case study in Section V. Finally, Section VI concludes the paper with some final remarks and future directions.

II. RELATED WORK

Several methods have been proposed to address the RM problem [3]. In this section, we report and discuss recent developments in model-based and model-free approaches, as well as the combination of the two, in comparison to our proposed methodology.

A. MPC Approaches for Ramp Metering Control

A well-studied approach to tackle RM is MPC. First formalized in [38], it was deployed for RM [39] and coordinated traffic control [10], [11], [12], and is still subject of research, e.g., [40] and [41]. As explained in Section I, its predictive capabilities make MPC highly desirable, but at the same time leave it susceptible to model mismatches and uncertainties. When these disturbances are assumed to be modelled, one can resort to robust or stochastic formulations to mitigate their impact, as proposed in, e.g., [42]; however, such assumptions are in general hard to satisfy in practice, and the additional computational burden can be taxing. For these reasons, MPC techniques for traffic control that can adapt to and circumvent model mismatches and uncertainties are still a vibrant research area.

B. Model-Free Approaches for Ramp Metering Control

Perhaps the most established RM approach is ALINEA [7], a local feedback strategy that seeks to maximise freeway flow in a merging area by keeping the bottleneck density near critical value. Many extensions have been proposed to enhance its basic formulation, such as PI-ALINEA, which deals with distant downstream bottlenecks [8], and FF-ALINEA, which anticipates the evolution of a bottleneck density [9]. However, these closed-loop controllers are fixed since they have no means to adapt their gains to changing conditions, and thus suffer from, e.g., uncertainties in traffic parameters and poor tuning. Ways have been devised to automatically account for it, e.g., via Kalman filtering to estimate the traffic parameters [43] or based on historic data [44]. Unfortunately, these methods are still plagued with a fundamentally myopic control strategy, and either assume that parameters evolve according to a known observable model, or that a huge amount of offline recorded data is available and is representative of both current and future traffic scenarios. Moreover, these feedback methods, alongside MPC, often need to be paired with a mechanism for traffic state estimation, enabling them to respond to real-time traffic behaviour [45].

More recently, RL has gained popularity thanks to its ability to automatically learn an optimal policy solely by

interacting with the environment, addressing the issue of uncertain and unknown dynamics. Earlier works adopted tabular Q-learning for the RM problem [46], [47], [48]. However, with function approximations becoming more and more predominant, especially NNs, DRL algorithms have emerged as a relevant alternative to model-based controllers. In [13], different policy gradient DRL algorithms are deployed in freeway control and compared. While results are promising, the issue of state constraints is not addressed, especially on the queue length, which should be contained to avoid spillback [3]. This is a fundamental issue in RL, since these algorithms can often be made aware of the presence of constraints only when violations occur and are appropriately penalised in the reward, whereas MPC can handle constraints explicitly in its formulation. In [14] it is proposed to use a combination of offline historic data and online synthetic data in an iterative way to avoid the RL model from being trapped in an inaccurate training environment. This work underscores the severe impact of high-quality data on the RL learning process, and raises the question of whether research efforts might be more effectively directed towards the development of online model-based learning algorithms that can learn directly from closed-loop interactions with the target environment while addressing constraints more explicitly. DRL has also facilitated the advancement of multi-agent solutions [49], which are especially useful in large-scale highway control problems. However, in addition to the aforementioned issues, multi-agent RL must also address the nonstationarity of its learning targets.

C. Combination of MPC and RL for Ramp Metering

The use of learning-based MPC schemes for traffic management has been investigated in the recent literature. For instance, [50] shows how recurrent NNs can be used to predict traffic behaviour, and [51] proposes to adopt such a network as prediction model in their MPC scheme.

Nonetheless, to the best of the authors' knowledge, the combination of MPC and RL in traffic control has received scarce attention, with [52] and [53] the only relevant works. Therein, an MPC-DRL architecture is proposed that hierarchically combines the two controllers: the high-level MPC controller, which runs at a lower frequency, provides initial optimality while explicitly incorporating constraints, whereas the low-level DRL agent operates at a higher frequency and aims to compensate for model mismatches in the MPC. The MPC and RL control actions are then linearly combined together and fed to the freeway environment as a unique control input.

Given that the final control signal is a summation of an optimisation-based one and a learning-based one, it is anticipated that the resulting policy may not always be able to prevent constraint violations, particularly early in the training process. However, a benefit of [53] compared to the proposed methodology is the higher control frequency at which the DRL policy is run, granting quicker control response. This is especially advantageous in applications with fast dynamics, although traffic systems are typically slow, with sampling times up to tens of seconds. Note that the slower control

frequency of our method is not inherently limiting as, with the improving hardware capabilities and faster solvers in recent years, increasing this frequency is feasible. The two methodologies also differ in computational complexity. While our approach is more demanding during the online training phase due to the need to solve the MPC problem twice per time step, it becomes more efficient at deployment, as it requires only one MPC calculation per time step [35]. In contrast, [53] runs both the MPC and DRL policies at every step, leading to higher computational costs at inference. Moreover, note that if training is performed offline, e.g., via an off-policy algorithm, then the additional computational burden of our method becomes irrelevant to online performance.

D. Comparison With Proposed Methodology

Compared to feedback strategies such as ALINEA, the MPC-based RL framework [35] we propose to use for RM leverages as its core an MPC controller, which, thanks to its predictive capabilities, can yield non-myopic policies. At the same time, its RL module allows to automatically adjust the MPC controller, and thus, it circumvents some of the challenges that conventional MPC typically struggles to address, i.e., model misspecifications and poor identification and tuning. This framework also boasts a higher degree of interpretability, better handling of constraints and integration of prior knowledge compared to other model-free RL methods. Furthermore, in contrast to [53], only one policy provider (the differentiable MPC controller itself) is present. For this reason, as remarked in Section I, this approach has the potential to result in fewer or no constraint violations at all depending on, e.g., the complexity of the underlying model and objective and the presence or lack of uncertainties [36], [37]. Likewise, at deployment, the proposed method relies only on the learnt MPC problem, whereas the method from [53] must additionally compute the output from the NN actor.

III. BACKGROUND

Three preliminary components are necessary to illustrate the proposed MPC-RL control method. First, a modelling framework must be employed to build a representation of the dynamical behaviour of the highway network, and how RM can influence it to prevent, e.g., bottlenecks and spillbacks. Then, we report a classical formulation of MPC for this task that makes use of such a model, and propose some modifications to it that are relevant to this work. Lastly, the essential concepts of RL are briefly introduced.

Notation The set of indices $\{1, \dots, N\} \subseteq \mathbb{N}$ is denoted as $\mathcal{N}$. Vector quantities are in bold. Inequalities on vectors are applied element-wise. To avoid confusion when other subscripts are present, dynamic quantities (e.g., ρ) at time step k can be also referred to as, e.g., $\rho_j(k)$. The i -th index of predicted quantities based on the information available at time step k are denoted as, e.g., $\hat{y}_{i|k}$.

A. METANET Modelling Framework

Based on the traffic phenomena to investigate, different approaches to modelling a highway network exist [54]. In this

paper, the macroscopic second-order METANET framework is employed to obtain a discrete-time dynamical representation of the behaviour of highway traffic under RM control. Proposed as a simulation tool in [55] and [56], it made one of its first appearances in combination with MPC in [12] and [39], and has since been extended to cover several phenomena and traffic measures [57], [58].

In the remainder of this section, the basics of METANET are introduced. In this framework, a network is represented as a directed graph where each segment (a continuous stretch of highway with homogeneous properties) is modelled as an arc, and origins are represented as nodes. Virtual nodes are also employed to separate segments with heterogeneous properties (e.g., due to a lane drop). We assume the network consists of $|\mathcal{S}|$ segments and $|\mathcal{O}|$ origins, where $\mathcal{S}$ and $\mathcal{O}$ are the sets of indices of segments and origins, respectively. At time step k , segment $j \in \mathcal{S}$ is characterized by its traffic density $\rho_j(k)$ and a mean speed $v_j(k)$, which form its state. Conversely, origin $o \in \mathcal{O}$ is characterized by the queue length $w_o(k)$. In the destination-independent variant, the evolution in time of segment states is described by the following equations:

$$\rho_j(k+1) = \rho_j(k) + \frac{T}{L_j \lambda_j} (q_{j-1}(k) - q_j(k) + r_o(k)), \quad (1)$$

$$\begin{aligned} v_j(k+1) = & v_j(k) + \frac{T}{\tau} (V_e(\rho_j(k)) - v_j(k)) \\ & + \frac{T}{L_j} v_j(k) (v_{j-1}(k) - v_j(k)) \\ & - \frac{\eta T}{\tau L_j} \frac{\rho_{j+1}(k) - \rho_j(k)}{\rho_j(k) + \kappa} \\ & - \frac{\mu T r_o(k) v_j(k)}{L_j \lambda_j (\rho_j(k) + \kappa)}, \end{aligned} \quad (2)$$

$$q_j(k) = \lambda_j \rho_j(k) v_j(k), \quad (3)$$

$$V_e(\rho(k)) := v_{\text{free}} \exp \left(-\frac{1}{a} \left(\frac{\rho(k)}{\rho_{\text{crit}}} \right)^a \right), \quad (4)$$

where T is the sampling time, L_j and λ_j are respectively the length of the segment and its number of lanes, (3) describes the traffic flow $q_j(k)$ of segment j at time step k , (4) is the well-known equilibrium speed formula, and $\tau, \eta, \kappa, \mu, a, \rho_{\text{crit}}, v_{\text{free}}$ are model parameters describing different quantities and phenomena. In particular, ρ_{crit} indicates the traffic density at which traffic flow transitions from free-flow conditions to congested conditions. At this density, the flow rate of vehicles through the network is at its maximum, and any increase in the number of vehicles beyond this point results in reduced speeds and increased congestion. The free-flow speed v_{free} describes instead the speed at which vehicles can travel when the road is not congested or disrupted, i.e., it is the maximum speed that vehicles can achieve on a particular segment under low traffic volumes. Lastly, $r_o(k)$ is the incoming flow generated by origin o connected to segment j ; if none is connected, it is null. In these equations, the index $j-1$ refers to the segment upstream of j , and $j+1$ to the segment downstream. For origins, the queue length evolves as

$$w_o(k+1) = w_o(k) + T(d_o^w(k) - r_o(k)), \quad (5)$$

where the origin's demand $d_o^w(k)$ acts as an uncontrollable external input affecting the queue dynamics. This work is mainly concerned with on-ramps; however, different types of models for origins exist, and interested readers are referred to [57]. Having defined the queue length as in (5), the on-ramp flow can be computed as

$$r_o(k) = \tilde{r}_o(k) \min \left\{ d_o^w(k) + \frac{w_o(k)}{T}, C_o, C_o \left(\frac{\rho_{\text{max}} - \rho_{j_o}(k)}{\rho_{\text{max}} - \rho_{\text{crit}}} \right) \right\}, \quad (6)$$

where C_o is the capacity of the origin, $\rho_{j_o}(k)$ is the density of the segment j_o that is connected to this ramp, and ρ_{max} is the maximum density segments can withstand (assumed here independent of index j_o , for sake of simplicity). Lastly, $\tilde{r}_o(k) \in [0, 1]$ is the metering rate control action, which allows modulating the entering ramp flow to prevent, e.g., congestion and spillback. However, to avoid numerical issues in the MPC gradient-based solvers due to the min operator in (6), which can cause the gradient to be zero over a vast region of the state-action space, it is more efficient to control the ramp flow r_o directly. This requires imposing the additional terms in (6) as constraints on the control action, as shown in the next subsection.

Finally, the downstream boundary conditions are here considered to be affected by congestion, i.e., the highway segment j_i that ends in destination $i \in \mathcal{D}$ is subject to the (virtual) downstream density modelled in [11] as

$$\rho_{j_i+1}(k) = \max \left\{ \min \{ \rho_{j_i}(k), \rho_{\text{crit}} \}, d_i^{\rho}(k) \right\}, \quad (7)$$

where $\mathcal{D}$ is the set of destination indices, $d_i^{\rho}(k)$ models the congestion density of destination i , and is regarded as another external disturbance.

In a more concise formulation, the states, actions and external factors at time k can be lumped in vectors as

$$\mathbf{x}_k = [\rho_1(k) \ \dots \ \rho_{|\mathcal{S}|}(k) \ v_1(k) \ \dots \ v_{|\mathcal{S}|}(k) \ w_1(k) \ \dots \ w_{|\mathcal{O}|}(k)]^T, \quad (8)$$

$$\mathbf{u}_k = [r_1(k) \ \dots \ r_{|\mathcal{O}|}(k)]^T, \quad (9)$$

$$\mathbf{d}_k = [d_1^w(k) \ \dots \ d_{|\mathcal{O}|}^w(k) \ d_1^{\rho}(k) \ \dots \ d_{|\mathcal{D}|}^{\rho}(k)]^T, \quad (10)$$

respectively. Then, the METANET framework can be exploited to build the nonlinear dynamics

$$\mathbf{x}_{k+1} = f(\mathbf{x}_k, \mathbf{u}_k, \mathbf{d}_k), \quad (11)$$

where $f : \mathbb{R}^{n_x} \times \mathbb{R}^{n_u} \times \mathbb{R}^{n_d} \rightarrow \mathbb{R}^{n_x}$, with $n_x = 2|\mathcal{S}| + |\mathcal{O}|$, $n_u = |\mathcal{O}|$, $n_d = |\mathcal{O}| + |\mathcal{D}|$. The next section details how these dynamics can be used in designing a model-based controller.

B. MPC for Ramp Metering Control

In what follows, the classical MPC-based paradigm for the RM control problem is presented and summarised.

MPC [18] is a well-known control framework for dynamical systems that is based on the formulation of a finite-horizon optimal control problem, consisting of three fundamental

pillars: the prediction of the evolution of the system's dynamics over a finite window, a suitable objective function to be minimised along this prediction window, and the inclusion of constraints on state and/or action variables. The first is typically achieved by means of a (often approximate) model of the system's behaviour, which is used to simulate the evolution of the states along the window. The second is a function that quantifies the performance of the controller and is often a trade-off between conflicting objectives. Thirdly, MPC allows to systematically integrate into its structure constraints on the states and/or actions, thus yielding policies that are able to take these constraints into account in an optimal way. On the whole, the framework offers quite a degree of flexibility, and it is unsurprising that it lends itself also to RM control problems.

We define the MPC prediction horizon N_p and the control horizon $N_c \leq N_p/M$, where M is a natural number to distinguish the process timescale T from the controller timescale $T_c = T/M$. In other words, the process runs M times faster than the controller, with the MPC being solved only every M time steps (instead of at each time step). Once solved, the first optimal action is then applied for the next M time steps, in a receding horizon fashion. Therefore, in the MPC controller timescale, control action indices are indicated as i_c and can be related to indices i of the dynamics timescale as $i_c(i) = \min\{\lfloor i/M \rfloor, N_c - 1\}$.¹ Additionally, we define a constrained RM control scenario by imposing at time step k , for the on-ramps in the network, the constraint

$$w_o(k) \leq w_{\max}, \quad \forall o \in \mathcal{O}, \quad (12)$$

where $w_{\max}$ represents the maximum queue length the ramps should experience. For the sake of simplicity, it is assumed here that $w_{\max}$ is independent of the index o , and that all queue lengths are subject to this constraint, though it could be readily applied to only a subset of them. As mentioned in Section II, this is usually done to prevent spillback, and similarly constrained scenarios can be found in, e.g., [12]. This constraint is not to be considered hard, but we favour control strategies that are able to satisfy it, if not always, at least for most time steps.

Altogether, given the current state x_k of the network, the MPC controller is given by

$$\min_{\hat{x}_k, \hat{u}_k, \Sigma} \sum_{i=0}^{N_p} L_T(\hat{x}_{i|k}) + \xi \sum_{i=0}^{N_c-1} L_V(\hat{u}_{i|k}) + \sum_{i=0}^{N_p} \zeta_i^\top \sigma_i \quad (13a)$$

$$\text{s.t. } \hat{x}_{0|k} = x_k, \quad (13b)$$

$$\hat{x}_{i+1|k} = f(\hat{x}_{i|k}, \hat{u}_{i_c(i)|k}, \mathbf{d}_k), \quad i = 0, \dots, N_p - 1, \quad (13c)$$

$$h(\hat{x}_{i|k}, \hat{u}_{i_c(i)|k}) \leq 0, \quad i = 0, \dots, N_p, \quad (13d)$$

$$g(\hat{x}_{i|k}, \sigma_i) \leq 0, \quad i = 0, \dots, N_p, \quad (13e)$$

$$\sigma_i \geq 0, \quad i = 0, \dots, N_p, \quad (13f)$$

where the optimisation variables are the actions and states along the MPC control and prediction horizons, and the slack

variables (whose purpose is explained later in the section), respectively in order

$$\hat{X}_k = [\hat{x}_{0|k}^\top \quad \dots \quad \hat{x}_{N_p|k}^\top]^\top \in \mathbb{R}^{n_x(N_p+1)}, \quad (14)$$

$$\hat{U}_k = [\hat{u}_{0|k}^\top \quad \dots \quad \hat{u}_{N_c-1|k}^\top]^\top \in \mathbb{R}^{n_u N_c}, \quad (15)$$

$$\Sigma = [\sigma_0^\top \quad \dots \quad \sigma_{N_p}^\top]^\top \in \mathbb{R}^{|\mathcal{O}|(N_p+1)}, \quad (16)$$

where $\hat{X}_k$ and $\hat{U}_k$ are defined in analogy to (8) and (9) respectively, and σ_i collects the slack variables at the predicted step i corresponding to all on-ramps, i.e., $\sigma_i = [\sigma_i^1 \dots \sigma_i^{|\mathcal{O}|}]^\top$. As aforementioned, once this optimisation problem is solved, we apply the optimal control action $\hat{u}_{0|k}^*$ from time step k to $k + M - 1$, as per the receding horizon approach.

The objective (13a) comprises three terms. The first is the total time spent (TTS), a metric frequently used to quantify efficiency of traffic control, defined as $L_T : \mathbb{R}^{n_x} \rightarrow \mathbb{R}$

$$L_T(x_k) := T \left(\sum_{j \in \mathcal{S}} L_j \lambda_j \rho_j(k) + \sum_{o \in \mathcal{O}} w_o(k) \right). \quad (17)$$

The second term of (13a) is a cost proportional to the variability of the inputs along the control horizon, where $L_V : \mathbb{R}^{n_u} \rightarrow \mathbb{R}$ is

$$L_V(u_k) := \sum_{o \in \mathcal{O}} \left(\frac{r_o(k) - r_o(k-1)}{C_o} \right)^2. \quad (18)$$

Lastly, the third term in (13a) acts as a penalty for the slack variables. The nonnegative weights $\xi \in \mathbb{R}$ and $\zeta_i \in \mathbb{R}^{|\mathcal{O}|}$, $i = 0, \dots, N_p$, govern the trade-off between these three terms and are selected by the designer so as to encode the desired behaviour of the controller [10]. Section IV will show how these parameters (under the new symbols θ_V and Θ_C), together with others, can be adjusted automatically via RL instead of having to be specified manually.

In (13d) and (13e), $h : \mathbb{R}^{n_x} \times \mathbb{R}^{n_u} \rightarrow \mathbb{R}^{4|\mathcal{O}|}$ and $g : \mathbb{R}^{n_x} \times \mathbb{R}^{|\mathcal{O}|} \rightarrow \mathbb{R}^{|\mathcal{O}|}$ correspond to all the inequality constraints to be imposed in the MPC problem. In particular, they are defined as

$$h(x_i, u_{i_c}) := [h_1(x_i, u_{i_c})^\top \quad \dots \quad h_{|\mathcal{O}|}(x_i, u_{i_c})^\top]^\top, \quad (19)$$

$$g(x_i, \sigma_i) := [g_1(x_i, \sigma_i) \quad \dots \quad g_{|\mathcal{O}|}(x_i, \sigma_i)]^\top, \quad (20)$$

where

$$h_o(x_i, u_{i_c}) := \begin{bmatrix} -r_o(i_c) \\ r_o(i_c) - C_o \\ r_o(i_c) - d_o^w(i) - \frac{w_o(i)}{T} \\ r_o(i_c) - C_o \left(\frac{\rho_{\max} - \rho_{j_o}(i)}{\rho_{\max} - \rho_{\text{crit}}} \right) \end{bmatrix}, \quad \forall o \in \mathcal{O}, \quad (21)$$

$$g_o(x_i, \sigma_i) := w_o(i) - w_{\max} - \sigma_i^o, \quad \forall o \in \mathcal{O}, \quad (22)$$

and we have dropped the dependency on i in i_c for readability. The constraints $h_o(x_i, u_{i_c}) \leq 0$ emulate the behaviour of (6) and prevent the control action from assuming invalid values.

¹This definition of $i_c(i)$ entails that the last action is kept constant for the remainder of the prediction horizon, a common choice in literature, though more complex solutions could be employed.

The constraints $g_o(\mathbf{x}_i, \boldsymbol{\sigma}_i) \leq 0$ instead incorporate the queue requirements (12) into the MPC controller as soft constraints, where $\boldsymbol{\sigma}_k \geq 0$ is a slack variable whose purpose is to relax the constraint in order to avoid infeasibility issues, and so it must be appropriately penalised in the objective (13a). We opt for a soft constraint because 1) it is assumed that small violations to this constraint, although undesired, are tolerable, and 2) we need a mechanism to both relax constraints during the learning process and penalise violations in order to inform the agent about policies that violate constraints and ones that do not.

Note that in (13b) we assume that the current process state $\mathbf{x}_k$ is fully measurable. If this is not the case, an observer is required to provide an estimate of it. Likewise, in (13c) we assume the external inputs $\mathbf{d}_k$ to be fully known. If this is not the case, a forecast of the evolution of these quantities along the MPC horizon is required.

C. Reinforcement Learning

To better understand how MPC can be integrated with RL, a brief introduction to the latter is here provided. In what follows, we abide by the canonical RL approach [21]. Given the continuous state s and action $\mathbf{a}$, a discrete-time Markov Decision Process with state transitions $s \xrightarrow{\mathbf{a}} s_+$ and underlying conditional probability density

$$\mathbb{P}[s_+ | s, \mathbf{a}] : \mathbb{R}^{n_s} \times \mathbb{R}^{n_s} \times \mathbb{R}^{n_a} \rightarrow [0, 1] \quad (23)$$

is used to model the discrete-time system dynamics. The performance of a given deterministic policy $\pi_\theta : \mathbb{R}^{n_s} \rightarrow \mathbb{R}^{n_a}$ parametrized in $\theta \in \mathbb{R}^{n_\theta}$ is defined as

$$J(\pi_\theta) := \mathbb{E} \left[\sum_{k=0}^{\infty} \gamma^k L(s_k, \pi_\theta(s_k)) \right], \quad (24)$$

where $L : \mathbb{R}^{n_s} \times \mathbb{R}^{n_a} \rightarrow \mathbb{R}$ is a stage cost function and $\gamma \in (0, 1]$ the discount factor. The goal of the RL algorithm is then to find the optimal policy π_θ^* as

$$\pi_\theta^* = \arg \min_{\theta} J(\pi_\theta). \quad (25)$$

Since it is in general impossible to find and characterise the true unknown optimal value functions and policy V_* , Q_* , π_* , function approximation techniques (such as NN and, as in this paper, MPC) have been employed as a powerful alternative for tackling problem (25) [23]. Depending on the algorithm, these allow formulating the approximations V_θ , Q_θ , and π_θ , which are then used to solve (25) either directly or indirectly via iterative gradient updates of the parametrisation

$$\theta \leftarrow \theta - \alpha \nabla_\theta \sum_{i=1}^m \psi(s_i, \mathbf{a}_i, s_{i+1}, \theta), \quad (26)$$

where $\alpha \in \mathbb{R}_+$ is the learning rate, m is the size of the batch of observations used in the update, and ψ captures the controller's performance and varies from algorithm to algorithm. Direct (or policy-based) methods explicitly parametrise and learn the approximation π_θ of the optimal policy itself. For instance, policy gradient methods attempt to solve (25) directly by moving along the gradient of the policy, i.e.,

$$\mathbb{E}[\psi(s_i, \mathbf{a}_i, s_{i+1}, \theta)] = J(\pi_\theta). \quad (27)$$

On the other hand, indirect (or value-based) methods opt to indirectly derive the optimal policy by estimating and learning the value functions V_θ , Q_θ from the underlying RL task, and implicitly derive the policy from these. A member of this category, Q-learning learns to approximate the action-value by solving

$$\min_{\theta} \mathbb{E} [\|Q_*(s, a) - Q_\theta(s, a)\|^2], \quad (28)$$

with the hope of indirectly recovering the optimal policy from it. In its first-order recursive (i.e., with $m = 1$) formulation, θ is updated via (26) with $\psi(s_i, \mathbf{a}_i, s_{i+1}, \theta) = \delta_i^2$, where δ_i is the well-known temporal difference (TD) error:

$$\delta_i = L(s_i, \mathbf{a}_i) + \gamma V_\theta(s_{i+1}) - Q_\theta(s_i, \mathbf{a}_i). \quad (29)$$

IV. MPC-BASED RL FOR RAMP METERING

As discussed in Section I, the closed-loop performance of the MPC controller (13) is dependent on its components, especially the prediction model. Mismatches in the dynamics can undermine the ability of MPC to reliably predict the system behaviour, thus impacting the controller performance and its ability to avoid constraint violations. The highly nonlinear nature of the METANET model contributes to making this issue even more relevant, since the model calibration requires the application of some nonlinear technique, e.g., nonlinear least-squares [59]. In what follows, we show how these shortcomings can be addressed by leveraging online closed-loop data to learn most of the MPC parameters automatically via RL.

A. Ramp Metering as an RL Task

To appropriately apply an RL algorithm to the RM task, we first need to define a proper stage cost $L(\mathbf{x}_k, \mathbf{u}_k)$ to quantify the performance of a policy in controlling the metering, as per (24). Also referred to as reward in literature (i.e., the negative of the cost), it defines the overall return that the RL algorithm must minimise and it encodes the objective that the designer wants to see minimised by the MPC control policy. Consider

$$L(\mathbf{x}_k, \mathbf{u}_k) := c_T L_T(\mathbf{x}_k) + c_V L_V(\mathbf{u}_k) + c_C L_C(\mathbf{x}_k), \quad (30)$$

where $L_C : \mathbb{R}^{n_x} \rightarrow \mathbb{R}$ is

$$L_C(\mathbf{x}_k) := \sum_{o \in \mathcal{O}} \max \{0, w_o(k) - w_{\max}\}. \quad (31)$$

The first term in (30) penalises the time spent travelling through the network and waiting in queues, at the current time step k . This term is representative of the TTS criterion, which is often the primary target of RM control strategies in literature. The second term penalises variations of the current control action $r_o(k)$ compared to the previous instant $r_o(k-1)$, and attempts to punish policies that are too jerky in the input signal. Lastly, the third contribution acts as the RL counterpart to the MPC soft constraint (22) and penalisation of the corresponding slack variables in the MPC objective (13a): it penalises violations of the queue constraint so that the agent is encouraged to come up with policies that are able to satisfy

this constraint. Scalars c_T , c_V , and c_C are weights that balance each term differently, and it is the designer's task to choose these in such a way to properly embed the learning objective in the RL stage cost.² Next, we delve into how MPC can be used as function approximation to solve the ramp metering RL task.

B. MPC as Function Approximation in RL

While the RL stage cost (30) defines and guides the learning process, a function approximation is needed to estimate the underlying optimal value and policy functions. A recurrent candidate for this rule in the literature is NNs. However, in Section I it was discussed how [35] suggests the employment of a parametric MPC scheme as function approximation in the RL problem, and lies the foundations of its theory. In this paper, we follow this paradigm. In particular, we choose to parameterise the cost function, the prediction model as well as some of the constraints in (13). Then, the RL algorithm is applied to adjust the parametrisation based on observed state transitions and rewards, in order to achieve better closed-loop performance in terms of (24).

Consider the MPC controller with parametrisation θ

$$\begin{aligned} \min_{\hat{x}_k, \hat{u}_k, \Sigma} \quad & \sum_{i=0}^{N_p} \gamma^i \left(\theta_T L_T(\hat{x}_{i|k}) + \theta_C^{\top} \sigma_i \right) \\ & + \theta_V \sum_{i=0}^{N_c-1} \gamma^{iM} L_V(\hat{u}_{i|k}) + \lambda_{\theta}(\hat{x}_{0|k}) \\ & + \sum_{i=1}^{N_p-1} \gamma^i \ell_{\theta}(\mathbf{x}_k) + \gamma^{N_p} \Gamma_{\theta}(\mathbf{x}_{N_p}) \end{aligned} \quad (32a)$$

$$\text{s.t. } \hat{x}_{0|k} = \mathbf{x}_k \quad (32b)$$

$$\hat{x}_{i+1|k} = f_{\theta}(\hat{x}_{i|k}, \hat{u}_{i_c(i)|k}, \mathbf{d}_k), \quad i = 0, \dots, N_p - 1, \quad (32c)$$

$$h_{\theta}(\hat{x}_{i|k}, \hat{u}_{i_c(i)|k}) \leq 0, \quad i = 0, \dots, N_p, \quad (32d)$$

$$g(\hat{x}_{i|k}, \sigma_i) \leq 0, \quad i = 0, \dots, N_p, \quad (32e)$$

$$\sigma_i \geq 0, \quad i = 0, \dots, N_p. \quad (32f)$$

with $\Theta_C = [\theta_C^0 \top \dots \theta_C^{N_p} \top]^{\top} \in \mathbb{R}^{|\mathcal{O}|(N_p+1)}$. Because this controller acts as function approximation and must be able to accurately predict the realisations of the RL stage cost along the prediction horizon, the objective (32a) is devised so as to mimic (30). Furthermore, note the additional parameterised terms $\lambda_{\theta}, \ell_{\theta}, \Gamma_{\theta} : \mathbb{R}^{n_x} \rightarrow \mathbb{R}$. These are the learnable initial, stage, and terminal costs respectively, and serve the purpose of enriching the parametrisation and facilitating generalisation across the state-action space. The initial cost is a linear combination of the state

$$\lambda_{\theta}(\mathbf{x}_k) := \sum_{j \in \mathcal{S}} \left(\theta_{\lambda,j}^{\rho} \frac{\rho_j(k)}{\rho_{\max}} + \theta_{\lambda,j}^v \frac{v_j(k)}{v_{\max}} \right) + \sum_{o \in \mathcal{O}} \theta_{\lambda,o}^w \frac{w_o(k)}{w_{\max}}, \quad (33)$$

²In our work, the values for c_T , c_V , and c_C were manually adjusted to achieve a reasonable learning process.

where $\rho_{\max}$, $v_{\max}$, $w_{\max}$ are here used as fixed, non-learnable normalization constants in order to improve the stability of the learning process. As explained in [35], since our objective is economic, this initial cost is required in order for the MPC scheme to recover the optimal policy. Conversely, stage and terminal costs are chosen to be quadratic quadratic

$$\begin{aligned} \ell_{\theta}(\mathbf{x}_k) := & \sum_{j \in \mathcal{S}} \theta_{\ell,j}^{\rho} \left(\frac{\rho_j(k) - \rho_{\text{sp}}}{\rho_{\max}} \right)^2 \\ & + \sum_{j \in \mathcal{S}} \theta_{\ell,j}^v \left(\frac{v_j(k) - v_{\text{sp}}}{v_{\max}} \right)^2 \\ & + \sum_{o \in \mathcal{O}} \theta_{\ell,o}^w \left(\frac{w_o(k)}{w_{\max}} \right)^2, \end{aligned} \quad (34)$$

$$\begin{aligned} \Gamma_{\theta}(\mathbf{x}_k) := & \sum_{j \in \mathcal{S}} \theta_{\Gamma,j}^{\rho} \left(\frac{\rho_j(k) - \rho_{\text{sp}}}{\rho_{\max}} \right)^2 \\ & + \sum_{j \in \mathcal{S}} \theta_{\Gamma,j}^v \left(\frac{v_j(k) - v_{\text{sp}}}{v_{\max}} \right)^2 \\ & + \sum_{o \in \mathcal{O}} \theta_{\Gamma,o}^w \left(\frac{w_o(k)}{w_{\max}} \right)^2, \end{aligned} \quad (35)$$

where densities and speeds are encouraged to track the fixed non-learnable setpoints ρ_{sp} and v_{sp} respectively, and the queues to be as small as possible. Regarding the METANET model f_{θ} in (32c), we allow the RL algorithm to adjust two fundamental parameters in the dynamics, i.e., $\tilde{\rho}_{\text{crit}}$ and $\tilde{a}$ (which in general can differ, during the learning process, from their counterparts ρ_{crit} , a of the real dynamics). In this way, the agent is able to partially tune the prediction model based on performance (24) rather than on system identification, implying that these learning parameters might grow to different values compared to the ones belonging to the true underlying dynamics. Finally, also constraints h_{θ} (32d) get parameterised, since $\tilde{\rho}_{\text{crit}}$ appears in it; see (13d) and (21).

Remark 1: In general, a parametrisation that makes use of quadratic terms with fixed setpoints, as in (34) and (35), can lead to suboptimal MPC policies in low-demand, free-flow regimes because it penalises low densities and low speeds as much as it penalises high densities and high speeds. At the same time, the theory at the foundation of MPC-based RL assumes the controller's parametrisation to be rich enough to adequately capture the real value function [35]. For this reason, in this work, we opt for the aforementioned adjustable quadratic cost terms. That being said, one could envision more involved and more performing learnable cost terms, especially in more complex and difficult traffic settings. For instance, asymmetric penalty functions are a valid choice, e.g., a left-sided Huber loss that quadratically penalises high traffic quantities, but penalises low ones only linearly.

Remark 2: The MPC parameters $\theta_T, \theta_V, \Theta_C$ are disjoint from the RL stage cost weights c_T, c_V, c_C of (30). While they encode similar notions, there is in general no reason to expect the values of the former to converge to the values of the latter during learning. Rather, they will converge to those values that maximise the controller's closed-loop performance.

TABLE I
PARAMETRISATION θ OF THE MPC FUNCTION APPROXIMATION (32)

Symbol	Scope	Space
$\tilde{\rho}_{\text{crit}}$	model, cost, constraint	$\mathbb{R}$
$\tilde{a}$	model	$\mathbb{R}$
θ_T	cost - TTS weight	$\mathbb{R}$
θ_V	cost - control variability weight	$\mathbb{R}$
Θ_C	cost - slack weights	$\mathbb{R}^{ \mathcal{O} (N_p+1)}$
$\theta_{\{\rho, v\}}^{\{\lambda, \ell, \Gamma\}}$	{init., stage, term.} cost - $\{\rho, v\}$ weights	$\mathbb{R}^{ \mathcal{S} }$
$\theta_{\{\lambda, \ell, \Gamma\}}^w$	{init., stage, term.} cost - w weights	$\mathbb{R}^{ \mathcal{O} }$

The whole parametrisation is

$$\theta = \begin{bmatrix} \tilde{\rho}_{\text{crit}} & \tilde{a} & \theta_{\{T, V\}} & \Theta_C^\top & \theta_{\{\rho, v, w\}}^{\{\lambda, \ell, \Gamma\}}^\top \end{bmatrix}^\top. \quad (36)$$

Table I shows a summary of the parametrisation, reporting the scope and the space of each term. As reported in Table II in the next section, these real spaces are often bounded in order to avoid undesired consequences, e.g., to preserve convexity of the parametric cost terms, or to prevent parameters that have some physical meaning from taking unrealistic values.

Scheme (32) yields the approximation $V_\theta : \mathbb{R}^{n_x} \rightarrow \mathbb{R}$ of the value function as

$$\begin{aligned} V_\theta(x_k) &= \min_{\hat{x}_k, \hat{u}_k, \Sigma} \quad (32a) \\ \text{s.t. } & (32b) - (32f). \end{aligned} \quad (37)$$

The value function (37) satisfies the fundamental equalities of the Bellman equations [35], so that

$$\begin{aligned} Q_\theta(x_k, u_k) &= \min_{\hat{x}_k, \hat{u}_k, \Sigma} \quad (32a) \\ \text{s.t. } & (32b) - (32f), \end{aligned} \quad (38)$$

$$\begin{aligned} \hat{u}_{0|k} &= u_k. \\ \pi_\theta(x_k) &= \arg \min_u Q_\theta(x_k, u). \end{aligned} \quad (39)$$

In practice, policy (39) is found as the first optimal action $\hat{u}_{0|k}^*$ from the arg min of (37). The goal of this parametrisation is to give the MPC scheme (32) enough degrees of freedom to sufficiently adapt to the RL task at hand. The cardinal point is that θ must be rich enough to allow the scheme to capture the optimal policy π^* [35]. Of course, the choice of parametrisation is not unique, and other solutions may be viable: similarly to NNs, where the topology of the network and the number of neurons and activation functions to used in each layer are arbitrary, here also different choices are possible. In the context of RM, the inclusion of some parameters in θ comes from the knowledge of the system and its control, e.g., the system's high sensitivity to the dynamics parameters a , ρ_{crit} and the cost weights θ_T , θ_V , θ_C . Other elements of the parametrisation, such as the additional quadratic terms, are instead dictated by the framework [35] to ensure a rich approximation scheme. However, unlike NNs, the benefit of this model-based approach is that knowledge about the system can be incorporated in the parametrisation rather trivially. Further discussion on different parametrisations and their impact on the learning outcome can be found in the appendix.

Beyond a proper selection of θ , other factors influence the suboptimality of the MPC-based RL solution. During training, the RL algorithm can only process a finite amount of data, whereas theoretical optimality is achieved only as this amount approaches infinity [21]. The MPC approximation scheme is also inherently suboptimal due to its nonlinearity, converging to global optimality only in the asymptotic limit, though the issue can be alleviated via, e.g., multi-start [18].

C. Second-Order LSTD Q-Learning

In order to adjust the parametrisation θ of (32), a second-order least-squares temporal difference (LSTD) Q-learning algorithm [60] is employed to find the policy π_θ that minimises the closed-loop performance. Q-learning is one of the most well-known indirect, temporal difference algorithms available in RL. In essence, it searches for the parametrisation that best fits the action-value function Q_θ to the observed data, with the aim of recovering via this approximation the unknown optimal Q_* and, indirectly from that, the optimal policy π_* . It has also shown very promising results in the context of MPC [61]. The second-order Newton's method, coupled with an experience replay buffer of the past observed transitions [62] and a reformulation of the Q-fitting problem as a least-squares, ensures faster convergence and higher sample efficiency compared to traditional first-order methods. For more details, see [61].

More concretely, given the approximation Q_θ (irrespective of its nature, e.g., MPC, NN, etc.), Q-learning solves the least-squares Bellman residual problem (28) via the second-order gradient update

$$\theta \leftarrow \theta - \alpha H^{-1} p, \quad (40)$$

where $\alpha \geq 0$ is the learning rate, and the gradient p and Hessian H are found as

$$p = - \sum_{i=1}^m \delta_i \nabla_\theta Q_\theta(s_i, a_i), \quad (41)$$

$$H = \sum_{i=1}^m \nabla_\theta Q_\theta(s_i, a_i) \nabla_\theta Q_\theta^\top(s_i, a_i) - \delta_i \nabla_\theta^2 Q_\theta(s_i, a_i), \quad (42)$$

with m denoting the experience sample batch size and δ_i the TD error (29). The update rule (40) requires differentiation of the MPC scheme (38) with respect to θ to find $\nabla_\theta Q_\theta$ and $\nabla_\theta^2 Q_\theta$. A nonlinear programming sensitivity analysis demonstrates how these quantities can be computed through differentiation of the Lagrangian of (38) [63]. Prior to any update, the computed Hessian matrix is modified to be positive-definite by the addition of a multiple of the identity matrix [64]. Furthermore, to impose lower and upper bounds on the parameters (as reported in Table II), (40) is cast as the following optimisation problem

$$\begin{aligned} \Delta\theta^* &= \arg \min_{\Delta\theta} \frac{1}{2} \Delta\theta^\top H \Delta\theta + \alpha p^\top \Delta\theta \\ \text{s.t. } & \theta_{\text{lb}} \leq \theta + \Delta\theta \leq \theta_{\text{ub}}, \\ & \Delta\theta_{\text{lb}} \leq \Delta\theta \leq \Delta\theta_{\text{ub}}. \end{aligned} \quad (43)$$

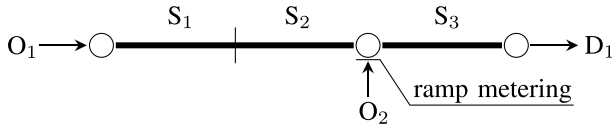


Fig. 2. Structure of the three-segment highway network.

This formulation allows to limit the rate of change of each parameter with $\Delta\theta_{lb}$ and $\Delta\theta_{ub}$. The parametrisation is then updated as $\theta \leftarrow \theta + \Delta\theta^*$. Note that, compared to (32), (43) adds negligible computational complexity as it is a convex quadratic problem and is usually solved at a lower, batch-update frequency (as explained in Section V-A3).

Lastly, as in [36], exploratory behaviour is ensured during learning by adding to the MPC objective (32a) the perturbation term $\mathbf{q}^\top \mathbf{u}_0$, with $\mathbf{q}$ randomly chosen from, e.g., a normal distribution. When properly calibrated, this perturbation on the first action ensures that enough exploration is added on top of the policy in order to prevent the Q-learning agent from getting stuck in very suboptimal local minima.

V. NUMERICAL CASE STUDY

To examine the performance of the proposed method for RM control, we implement and simulate the algorithm when applied to a highway network benchmark.

A. Configuration

1) *Network Environment*: Let us consider a simple highway traffic network taken from [11] and pictured in Fig. 2, which will serve as the RL environment for the task at hand. The network consists of three segments, each 1 km in length and with two lanes. The first segment S_1 is supplied by the uncontrolled mainstream origin O_1 , which simulates incoming traffic from the unmodelled upstream region of the network, and is characterized by its capacity C_1 . Between segments S_2 and S_3 , additional traffic can enter the network via the on-ramp O_2 . This origin has capacity C_2 and offers the only control measure in order to avoid congestion in the network. However, we would also like to keep the queue length on this on-ramp to 50 vehicles or fewer. Lastly, traffic exits the network via the only available destination D_1 with congested outflow.

As described in Section III-A, the network environment is subject to one external input per origin and destination. Section 3 depicts the demands at both origins, starting low and then increasing to near capacity within 20 minutes. With some delay, also congestion at the destination rapidly appears. These profiles are intended to simulate a peak-hour-like traffic (e.g., morning or evening rush), and are devised in such a way to induce a variable degree of congestion in the network, if the on-ramp is not properly controlled. At the start of each training episode, these two-peak profiles are generated randomly to introduce a degree of variability in the task, aimed at enhancing the agent's generalisation to a wider variety of congestion scenarios. Nevertheless, generalisation in RL is challenging [65], and a performance drop is anticipated were the controller to face a new unforeseen scenario, e.g., with a different number of peaks. At the same

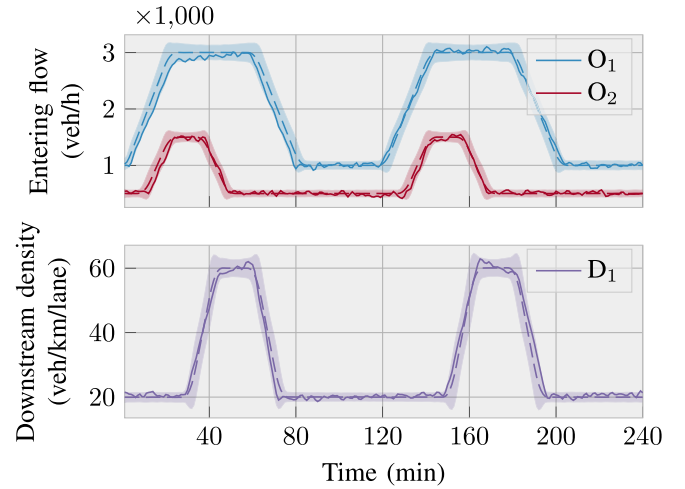


Fig. 3. External inputs affecting the network's dynamics, in the form of the demands at the origins (upper) and the congestion scenario at the destination (lower). The shaded areas represent the 2-standard deviation ranges from which the random demands are sampled, whereas the solid lines represent one random sample for each.

time, our approach exhibits an inherent level of robustness: it incorporates an online learning process that can seamlessly integrate novel training data on the fly, and at its core is a model-based controller, which generally extrapolates better to new situations in comparison to its model-free counterparts.

2) *Model Predictive Control*: The controller (32) is deployed to control the on-ramp flow with the aim of avoiding congestion in the traffic network environment. In line with other similar benchmarks in the literature, e.g., [10], [11], and [12], $M = 6$, meaning that, while the process dynamics are stepped every 10 seconds, the controller is run only once per minute. The prediction horizon is set to $N_p = 24$ (which is in the order of the average travel time through the network), and the control horizon to $N_c = 3$.

To showcase the ability of RL to automatically improve the performance of the MPC controller, the parametrisation θ of the controller is poorly initialised on purpose. In particular, in order to replicate the effects of a poor system identification phase, the parameters of the prediction model are initialised with a 30% error with respect to their true values, i.e., $\tilde{\rho}_{crit} = 0.7\rho_{crit}$, $\tilde{a} = 1.3a$ (both learnable), and $1.3v_{free}$ (fixed). Furthermore, the remaining parametrisation of the objective function of the MPC optimisation problem is left untuned, imitating a suboptimal, low-expertise design process of the controller. Table II reports the initial values of each learnable parameter. The other non-learnable parameters in (32) are set to the same values as in the real process. The queue threshold on O_2 is set to $w_{max} = 50$ veh. Lastly, the setpoints are set to $\rho_{sp} = 0.7\rho_{crit}$, $v_{sp} = 1.3v_{free}$, and the normalization coefficients to $\rho_{max} = 180$ veh/km/lane, $w_{max} = 50$ veh, and $v_{max} = 1.3v_{free}$.

3) *Reinforcement Learning*: The agent is trained in an episodic fashion, with each episode lasting $T_{fin} = 4$ h, i.e., each episode features two peaks in demands and congestion in a row (for reference, see Fig. 3). At the end of each episode, the demands are generated anew randomly, and the state of the network is reset to steady-state in order to avoid unreasonable accumulation of vehicles in queues from past

TABLE II
INITIAL VALUES, BOUNDS AND DIMENSIONS OF THE
PARAMETRISATION θ OF (32)

Symbol	Initial value	Bounds	Dimension
$\bar{\rho}_{\text{crit}}$	23.45	[10, 162]	veh/km/lane
$\bar{a}$	2.4271	[1, 3]	-
θ_T	1	$[10^{-3}, \infty)$	1/veh/h
θ_V	160000	$[10^{-3}, \infty)$	veh ² /h ²
Θ_C	5	$[10^{-3}, \infty)$	1/veh
$\theta_{\{\rho, v, w\}}$	1	$(-\infty, \infty)$	-
$\theta_{\{\rho, v, w\}}$	1	$[10^{-6}, \infty)$	-
$\theta_{\{\ell, \Gamma\}}$			

episodes. With the same frequency, the parametrisation θ is updated, and a new episode is started, till a termination condition of the learning process is met, typically in the form of a maximum number of iterations. Here, we train our agent for 80 episodes. With the aforementioned update frequency, the proposed learning setup is off-policy, since θ remains fixed throughout each episode, ensuring stable data collection.

The stage cost function $L(\mathbf{x}_k, \mathbf{u}_k)$, with which the traffic network feeds a cost signal back to the agent according to (30), has coefficients $c_T = 5$, $c_V = 1600$, and $c_C = 5$.

Based on results obtained from initial simulations, the algorithm's hyperparameters are set as follows. The discount factor is set to $\gamma = 0.98$. The learning rate is initially set to $\alpha = 0.925$, and decayed by a multiplicative factor of 0.925 at the end of each update. The maximum update change of each parameter is set to 30% of its current value, i.e., $\Delta\theta_{\text{ub}} = -\Delta\theta_{\text{lb}} = 0.3\theta$. As aforementioned, a replay buffer is used to store past observations in memory and re-use them. In particular, the buffer size is set up to store transitions from the 10 past episodes and, before performing an update, a batched sample half its size is drawn from this buffer. In turn, half of this batch is dedicated to containing information from the latest two and a half episodes, whereas the remaining half is sampled uniformly at random from even older transitions. All together, these measures contribute to stabilising the learning process.

Lastly, exploration is induced as described in Section IV-C by sampling the cost perturbation from a normal distribution in an ε -greedy fashion, i.e., $\mathbf{q} \sim \mathcal{N}(0, \sigma_q)$ with probability ε ; otherwise, $\mathbf{q} = 0$. The exploration strength is $\sigma_q = 0.025$, and the exploration probability $\varepsilon = 0.5$. Both are decayed by half at the end of each episode.

Table III summarises all the hyper-parameters and non-learnable parameters of the traffic environment and the MPC controller, as found in [11], as well as of the RL algorithm.

4) *Setup*: The simulations were run on a Ubuntu 20.04.6 server equipped with 16 AMD EPYC 7252 (3.1 GHz) processors and 252GB of RAM, and implemented in Python 3.11.4. The nonlinear optimisation problems were formulated and solved with the symbolic framework CasADi [66] and its interface to the IPOPT solver [67]. The source code and simulation results are open and available in the following repository: <https://github.com/FilippoAirdi/mpcrl-for-ramp-metering>.

B. Comparison

Alongside the proposed MPC-based RL controller, we implement and simulate on the same setup three other policies. These include another learning-based approach based on DDPG, a DRL algorithm that has been proven effective in freeway control [13], [53], and two other non-learning solutions: the classical non-learning MPC formulation (13) [11], as well as the well-known local ramp metering PI-ALINEA [7], [8]. To promote a fair comparison, all of the tested controllers suffer from the same inexact knowledge of the traffic parameters.

The relevant hyper-parameters for DDPG are taken from [53]. The PI-ALINEA controller is equipped with a queue management strategy to take into account the queue length constraint [68], and its gains are fine-tuned to this specific task via Bayesian Optimisation.

C. Results

We evaluate the performance of all the aforementioned methods deployed on the traffic network environment and average the results over 15 simulations with different seeds to iron out randomness due to, e.g., exploration. Figures that are shown and discussed next report the average results; however, 95% confidence intervals are shown only for our proposed method to reduce visual clutter.

Fig. 4 shows, for each method, the total RL cost (30) incurred in each episode, i.e., $\sum_{k=1}^{T_{\text{fin}}/T} L(\mathbf{x}_k, \mathbf{u}_k)$, subdivided in its three contributions, and how these evolve during learning. Due to the untuned objective weights and the 30% mismatches in the predictive model's parameters, the initial MPC-based RL controller is poorly performing and cannot reliably avoid congestion and constraint violations, as can be seen from the costs in the first episodes. However, despite these initial shortcomings, it can be noticed that in the next 20 episodes, by exclusively leveraging the observed transitions via Q-learning, the controller is able to achieve a substantial improvement of the Total-Time-Spent cost, i.e., the time spent by vehicles in the network on average, from around 1100 veh h to 700 veh h (or a 35% reduction). At the same time, as constraint violations of w_{max} in O_2 induce the most severe cost realization by several orders of magnitude, one can see that the agent is even quicker in learning to avoid such a constraint-violating behaviour within the first 5 episodes, despite the fact it has no knowledge of the exact traffic dynamics parameters.

This phenomenon can be further appreciated in Fig. 5, which shows the progression of the average queue in the on-ramp O_2 with respect to its constraint, as the parametrised MPC scheme gets better and better at yielding a policy with less or no constraint violation. It must be noted here that, despite the non-stationary nature of the traffic environment, the agent is able to learn a policy that is robust to the variability of the randomly generated demand and congestion scenarios and, therefore, is capable of avoiding all constraint violations.

Further analysis of the results indicates that, while learning to satisfy the constraint on the on-ramp O_2 is directly beneficial to the reduction of cost related to violations,

TABLE III
HYPER- AND OTHER NON-LEARNABLE PARAMETERS

METANET	Symbol Value Dimension	T 10 s	τ 18 s	η 60 km ² /lane	κ 40 veh/km/lane	μ 0.0122 -	$\rho_{\max}$ 180 veh/km/lane	ρ_{crit} 33.5 veh/km/lane	v_{free} 102 km/h	a 1.867 -	$C_{\{1,2\}}$ {3500, 2000} veh/h
MPC	Symbol Value Dimension	M 6 -	N_p 24 -	N_c 3 -	$w_{\max}$ 50 veh	$v_{\max}$ 132.6 km/h	ρ_{sp} 23.45 veh/km/lane	v_{sp} 132.6 km/h			
RL	Symbol Value	c_T 5	c_V 1600	c_C 5	batch size 5 episodes	γ 0.98	α 0.925	α decay 0.925	$\Delta\theta_{\{\text{lb}, \text{ub}\}}$ {-0.3 θ , 0.3 θ }	σ_q 0.025	ε 0.5

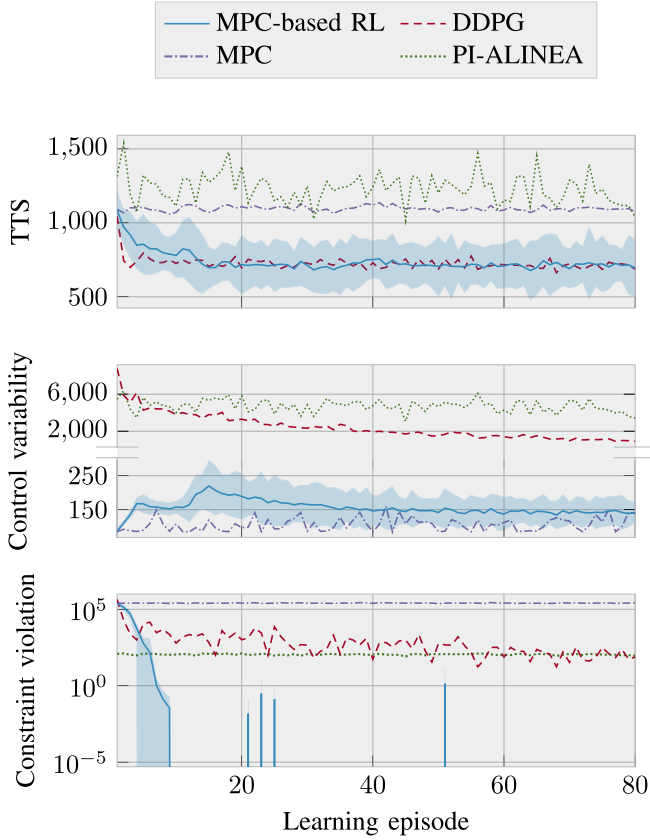


Fig. 4. Evolution of the three contributions to the RL cost (30) during the learning process, namely, from top to bottom, the Total-Time-Spent (TTS), variability of the control action, and violation of the maximum queue constraint on O_2 .

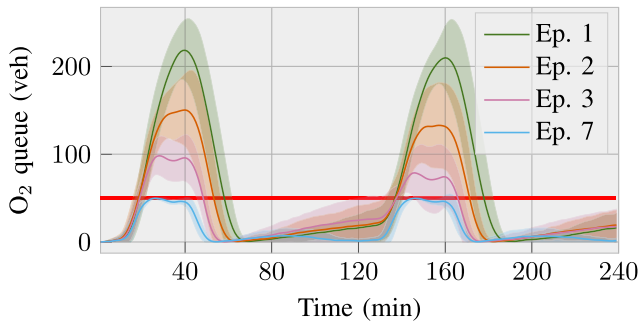


Fig. 5. On average, the policy π_θ gets better at avoiding constraint violations as it learns (the red line represents the threshold $w_{\max}$ imposed on the queue on the on-ramp O_2).

it also indirectly fosters better TTS, since longer queues in O_2 proportionally relate to a higher TTS. Nonetheless, improvements to the violation and TTS contributions seem to occur at the expense of the third contribution, i.e., variability

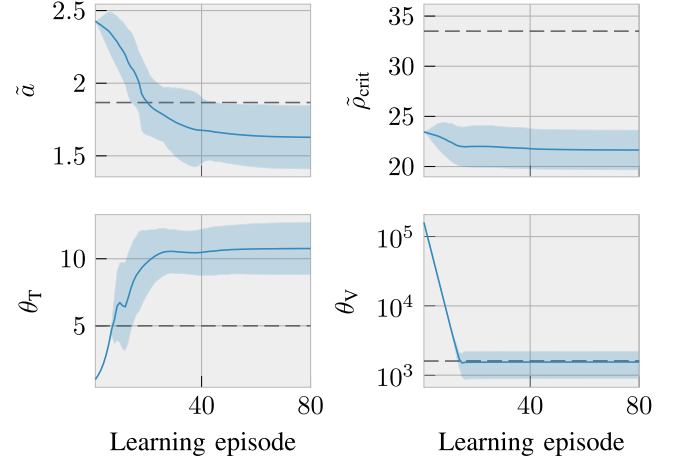


Fig. 6. Evolution of a subset of θ during the learning process (the dashed grey lines represent, in the top plots, the true values of the parameters a and ρ_{crit} , and, in the bottom ones, the constants c_T and c_V).

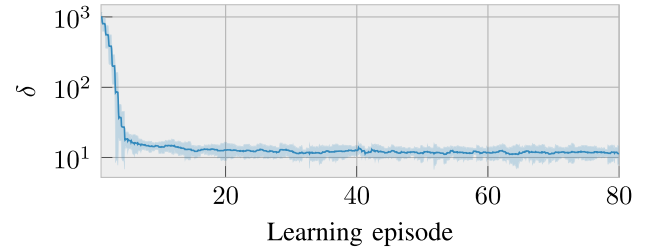


Fig. 7. Moving average of the TD error during the learning process (window length = 239, i.e., number of transitions per episode).

of the control action, which clearly is on the rise during learning for the MPC-based RL agent. However, in the overall context of learning, one must notice that at convergence the sacrifice in control variability (of around 60 points) allows the agent to gain a larger improvement in TTS and to achieve zero constraint violations (which are indeed the largest contributors to the cost of the RL agent). Furthermore, as the learning proceeds past the 16th episode and the other two costs have settled, the agent is also able to achieve further reduction of the cost associated to the control action variability.

The results from the other methods appear to corroborate the usefulness of data-driven adaptability in such a context. The non-learning MPC formulation is stuck at the initial poor performance of its RL-enhanced counterpart, with high TTS and constraint violations. On the contrary, thanks to its queue management strategy, PI-ALINEA is able to curb violations, but at the expense of a much higher control variability (since the strategy frequently requires to switch control action abruptly) and higher TTS. The DDPG agent

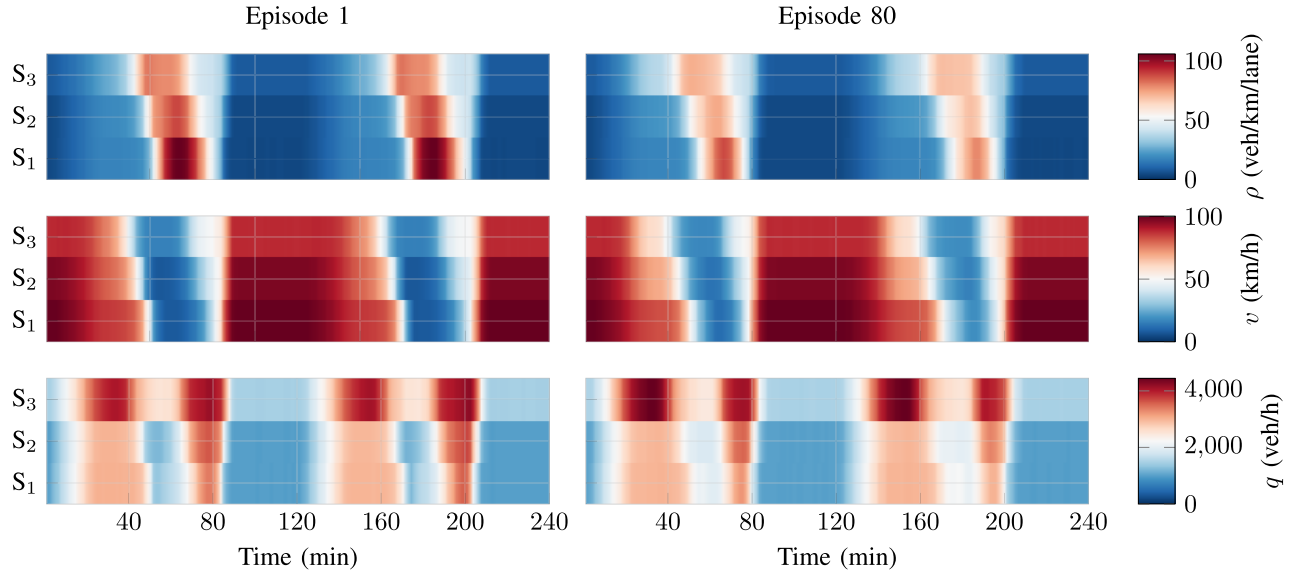


Fig. 8. Differences in traffic quantities of the network's three segments between the first episode and the last episode of the learning process, averaged across the 15 simulation runs.

follows a similar trend to our method in quickly learning to reduce TTS. However, variability and violations are decreasing at a substantially lower rate. This is not unexpected, as the NN approximation of this agent has significantly more learnable parameters (namely, 293 380 vs. 53) and is thus less data-efficient and requires more episodes to establish convergence.

As discussed in Section I, an advantage of the proposed MPC-RL approach over DRL is that its learning process can be more readily inspected and explained via the evolution of its parametrisation θ , reported in Fig. 6. In particular

- the evolution of θ_T and θ_V agrees with the observations made in Fig. 4, that is, the agent learns to favour the TTS cost over the control action variability, and thus increases the weight of the former at the expense of the latter
- the parameter $\tilde{a}$, which the METANET model is known to be very sensitive to, develops at first towards its true value a , but then converges to a slightly lower value, resulting in a less pronounced, less aggressive equilibrium speed V_e (4), i.e., favouring predictions of lower speeds at lower densities, and higher speeds at higher densities
- $\tilde{\rho}_{crit}$ grows in the opposite direction of its true value ρ_{crit} , thus settling for a prediction model that is pessimistic towards congestion, i.e., it tends to overestimate congestion scenarios, as well as resulting in a tighter constraint h_o^1 (21), which in turn restricts the control action further.

Fig. 7 depicts the TD error during the learning process, whose moving average converges to smaller and smaller values as the episodes increase. This is a good indication that the MPC scheme (32), with its parametrisation θ , is able to provide a reliable approximation Q_θ of the true unknown action-value function (and, indirectly, of the optimal policy). Yet, it is important to note that the convergence of the TD error, while consistently diminishing, does not reach an absolute zero. This observation can be attributed to factors such as imperfect parametrisation and the inherent stochasticity present in the environment, which fluctuates

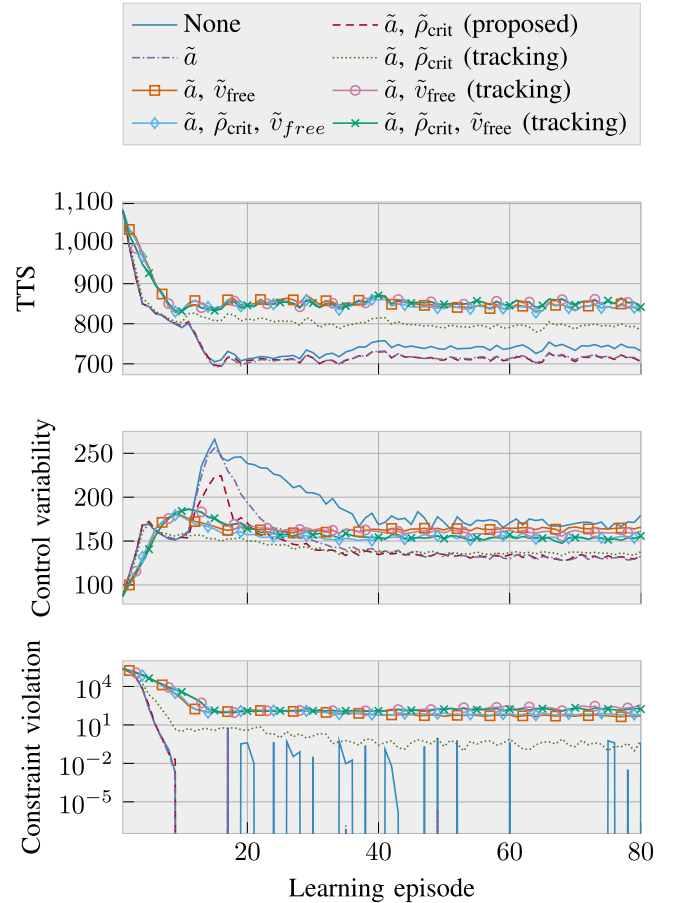


Fig. 9. Evolution of the RL cost contributions for various different parametrizations θ . The legend lists for each entry the dynamics parameters that were allowed to be learnt in the corresponding simulation, including also the parametrisation that we proposed in this work, i.e., (proposed). If either $\tilde{\rho}_{crit}$, $\tilde{v}_{free}$ or both were also employed as tracking setpoints, the entry is marked accordingly, i.e., (tracking). To avoid visual clutter, only the average over the 15 simulation runs is shown for each parametrization.

in its demands, and thus renders the prediction of the value functions more challenging. Besides these issues,

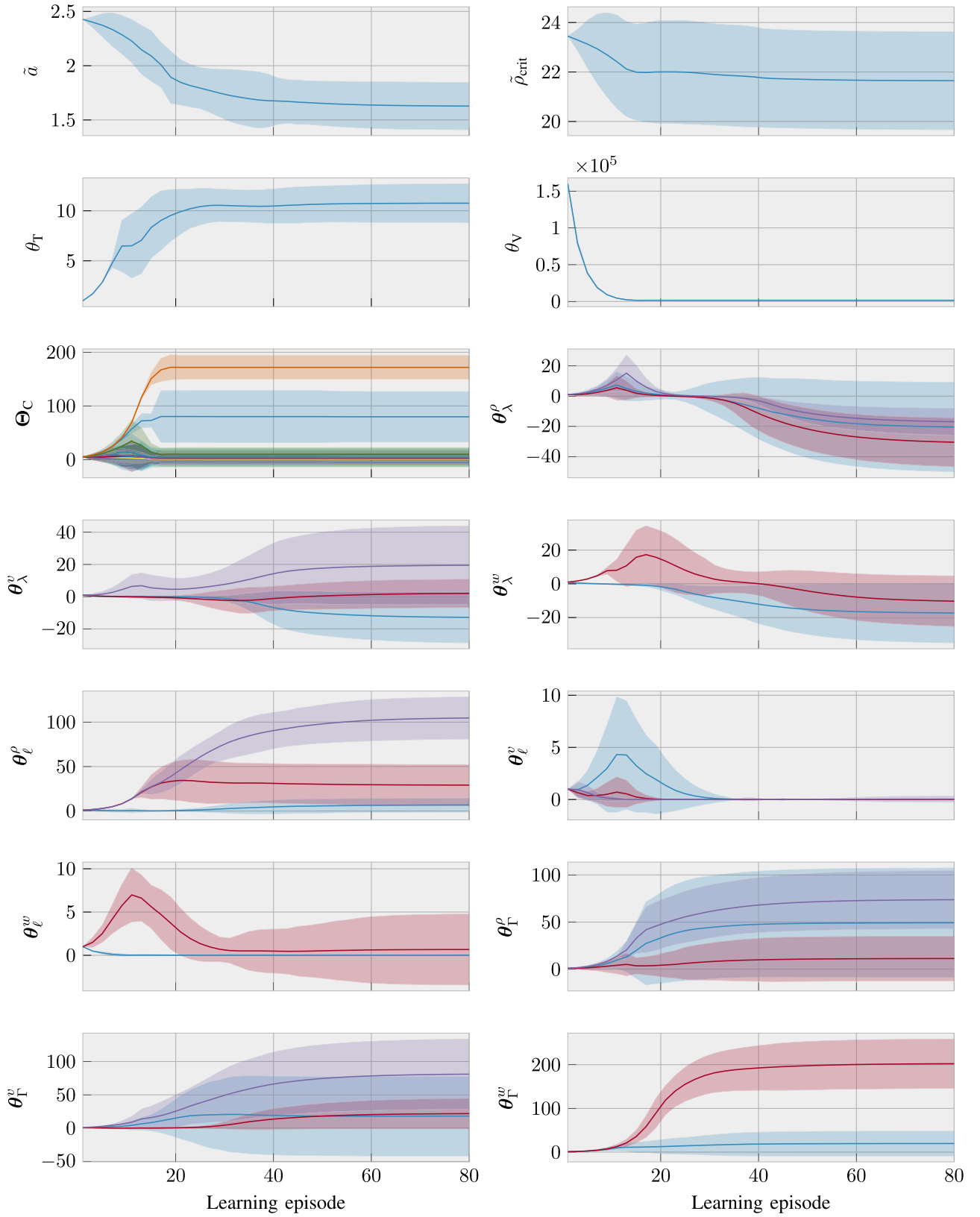


Fig. 10. Average evolution of the parametrization θ during the learning process across the 15 simulation runs. For those parameters that are vector quantities, one colour is associated with each entry.

another challenge related to the TD error is its sparsity. Specifically, the instantaneous TD signal remains low and largely uninformative for most of the episode, only to spike significantly during congestion events. This sparsity can lead

to inefficient updates if not properly addressed. Therefore, we advocate that the use of experience replay and batch updates is essential to smooth the learning process and mitigate the risk of destabilization due to abrupt parameter adjustments.

Finally, Fig. 8 reports the evolution of the traffic quantities of the three segments (i.e., density ρ , speed v , and flow q) that make up the network. Interestingly, it can be noticed how, compared to the initial episode, the final controller is able to better prevent the congestion in the last segment from propagating backwards through the network, and from causing speed drops in the first segment.

VI. CONCLUSION

In this paper, we have proposed a novel learning-based and model-based approach to the ramp metering problem that combines MPC and RL. The two frameworks are integrated in such a way as to promote the strengths of each while countering the disadvantages by exploiting the parametrised MPC scheme as a function approximation of the action-value function and leveraging RL to adjust the parametrisation based on observed data to improve closed-loop performance. Even with wrong model parameters and a poorly tuned initial controller, the proposed methodology shows a remarkable ability in learning to improve traffic control performance and satisfy constraints in an automatic, data-driven fashion.

Future work will focus on 1) the validation of the proposed methodology with microscopic traffic model simulators, 2) the study of formal guarantees on stability and recursive feasibility for the proposed approach during learning and at convergence, 3) more complex parametrisations of the MPC scheme, e.g., with neural networks, in order to better address the nonlinear nature of the METANET framework and to better capture the shape of the true action-value function, as well as 4) extending the current approach by integrating variable speed limits (VSLs) with ramp metering, a coordinated control strategy that has been shown to achieve better results than ramp metering alone, especially in high demand/density regimes.

APPENDIX

We provide here additional insights on our proposed methodology and the numerical results. First, we provide a short investigation on the effects of different parametrisations of the MPC scheme (32). Then, we report the evolution of all the parameters θ during the learning process implemented in Section V.

DIFFERENT MPC PARAMETRISATIONS

Different parametrisations θ of the MPC scheme (32) have been tested in simulations. In particular, we have considered the following variants:

- learning a combination of the dynamics parameters $\tilde{a}$, $\tilde{v}_{\text{free}}$, and $\tilde{\rho}_{\text{crit}}$ (including the empty set, i.e., learning none of these parameters)
- employing the learnable dynamics parameters $\tilde{v}_{\text{free}}$ and $\tilde{\rho}_{\text{crit}}$ also as tracking setpoints of the stage cost (34) and terminal cost (35), i.e., $v_{\text{sp}} = \tilde{v}_{\text{free}}$, $\rho_{\text{sp}} = \tilde{\rho}_{\text{crit}}$.

The rationale behind trying out these various combinations is that, while having a richer parametrisation can potentially help in fitting more complex value functions, a drawback is that the learning task becomes much more complex and more likely to converge to a very suboptimal local minimum. Moreover,

the sensitivity of the RL solution to the parametrisation can vary significantly from parameter to parameter, so learning some of them may end up in a less stable process than others. As is the case with most function approximators, the choice of parametrisation relies mostly on prior and/or expert knowledge (such as, in our case, the sensitivity analysis of the METANET model w.r.t. its parameters) and is mostly an iterative procedure (both for MPC and other function approximators), whose difficulty is attributable to the RL algorithm itself rather than to the proposed control scheme. However, meta-learning methods can be found in the literature to optimise over the selection of this and other hyperparameters via, e.g., Bayesian Optimisation [69] and bilevel optimisation [70].

Fig. 9 shows the influence on the performance of some of the combinations we tested. These tests were carried out with a lower variability of the randomly generated profiles in an effort to filter out the impact of randomness on the performance. One can notice that, when neither the dynamics parameters nor the tracking setpoints are learnt, the resulting controller converges to a solid performance. Adding $\tilde{a}$ to the set of learnable parameters seems to boost the performance even further, which can be attributed to the fact that, among the various parameters, $\tilde{a}$ is the one the METANET model is most sensitive to. Further testing indicates that learning also $\tilde{\rho}_{\text{crit}}$ achieves an additional performance improvement, while learning $\tilde{v}_{\text{free}}$ does not help. The last empirical finding is that making the tracking setpoints learnable, despite increasing the number of degrees of freedom of the parametrisation, does not bring any improvement.

EVOLUTION OF PARAMETRISATION

Interested readers can find in Fig. 10 the evolution, during the whole learning process, of the whole parametrisation θ of the MPC scheme (32), as detailed in Section IV-B and simulated in Section V.

REFERENCES

- [1] O. Johansson, D. Pearce, and D. Maddison, *Blueprint 5: True Costs of Road Transport*. Evanston, IL, USA: Routledge, 2014.
- [2] *Global Status Report on Road Safety 2018*, World Health Org., Geneva, Switzerland, 2018.
- [3] S. Siri, C. Pasquale, S. Saccone, and A. Ferrara, "Freeway traffic control: A survey," *Automatica*, vol. 130, Aug. 2021, Art. no. 109655.
- [4] K. Castañeda, O. Sánchez, R. F. Herrera, and G. Mejía, "Highway planning trends: A bibliometric analysis," *Sustainability*, vol. 14, no. 9, p. 5544, May 2022.
- [5] M. Papageorgiou and A. Kotsialos, "Freeway ramp metering: An overview," *IEEE Trans. Intell. Transp. Syst.*, vol. 3, no. 4, pp. 271–281, Dec. 2002.
- [6] J. A. Wattleworth, "Peak-period analysis and control of a freeway system," *Highw. Res. Rec.*, vol. 157, pp. 1–21, Oct. 1965.
- [7] M. Papageorgiou, H. Hadj-Salem, and J.-M. Blosseville, "ALINEA: A local feedback control law for on-ramp metering," *Transp. Res. Rec.*, vol. 1320, pp. 58–64, Jan. 1991.
- [8] Y. Wang, E. B. Kosmatopoulos, M. Papageorgiou, and I. Papamichail, "Local ramp metering in the presence of a distant downstream bottleneck: Theoretical analysis and simulation study," *IEEE Trans. Intell. Transp. Syst.*, vol. 15, no. 5, pp. 2024–2039, Oct. 2014.
- [9] J. R. D. Frejo and B. De Schutter, "Feed-forward ALINEA: A ramp metering control algorithm for nearby and distant bottlenecks," *IEEE Trans. Intell. Transp. Syst.*, vol. 20, no. 7, pp. 2448–2458, Jul. 2019.

- [10] A. Hegyi, B. De Schutter, H. Hellendoorn, and T. van den Boom, "Optimal coordination of ramp metering and variable speed control—An MPC approach," in *Proc. Amer. Control Conf.*, vol. 5, May 2002, pp. 3600–3605.
- [11] A. Hegyi, "Model predictive control for integrating traffic control measures," Ph.D. dissertation, Delft Center Syst. Control, Delft Univ. Technol., Delft, The Netherlands, 2004.
- [12] A. Hegyi, B. De Schutter, and H. Hellendoorn, "Model predictive control for optimal coordination of ramp metering and variable speed limits," *Transp. Res. C, Emerg. Technol.*, vol. 13, no. 3, pp. 185–209, Jun. 2005.
- [13] C. Wang, Y. Xu, J. Zhang, and B. Ran, "Integrated traffic control for freeway recurrent bottleneck based on deep reinforcement learning," *IEEE Trans. Intell. Transp. Syst.*, vol. 23, no. 9, pp. 15522–15535, Sep. 2022.
- [14] Y. Han, M. Wang, L. Li, C. Roncoli, J. Gao, and P. Liu, "A physics-informed reinforcement learning-based strategy for local and coordinated ramp metering," *Transp. Res. C, Emerg. Technol.*, vol. 137, Apr. 2022, Art. no. 103584.
- [15] K. Shaaban, M. A. Khan, and R. Hamila, "Literature review of advancements in adaptive ramp metering," *Proc. Comput. Sci.*, vol. 83, pp. 203–211, Jan. 2016.
- [16] F. Vrbanić, E. Ivanjko, K. Kušić, and D. Čakija, "Variable speed limit and ramp metering for mixed traffic flows: A review and open questions," *Appl. Sci.*, vol. 11, no. 6, p. 2574, Mar. 2021.
- [17] X. Ma, A. Karimpour, and Y.-J. Wu, "Statistical evaluation of data requirement for ramp metering performance assessment," *Transp. Res. A, Policy Pract.*, vol. 141, pp. 248–261, Nov. 2020.
- [18] J. B. Rawlings, D. Q. Mayne, and M. Diehl, *Model Predictive Control: Theory, Computation, and Design*, 2nd ed., Madison, WI, USA: Nob Hill, 2018.
- [19] A. Bemporad and M. Morari, "Robust model predictive control: A survey," in *Robustness in Identification and Control*, A. Garulli and A. Tesi, Eds., London, U.K.: Springer, 1999, pp. 207–226.
- [20] A. Mesbah, "Stochastic model predictive control: An overview and perspectives for future research," *IEEE Control Syst. Mag.*, vol. 36, no. 6, pp. 30–44, Dec. 2016.
- [21] R. S. Sutton and A. G. Barto, *Reinforcement Learning: An Introduction*. Cambridge, MA, USA: MIT Press, 2018.
- [22] C. J. C. H. Watkins and P. Dayan, "Q-learning," *Mach. Learn.*, vol. 8, pp. 279–292, May 1992.
- [23] L. Busoniu, R. Babuska, B. De Schutter, and D. Ernst, *Reinforcement Learning and Dynamic Programming Using Function Approximators*. Boca Raton, FL, USA: CRC Press, 2017.
- [24] V. Mnih et al., "Playing Atari with deep reinforcement learning," 2013, *arXiv:1312.5602*.
- [25] G. Dulac-Arnold et al., "Challenges of real-world reinforcement learning: Definitions, benchmarks and analysis," *Mach. Learn.*, vol. 110, no. 9, pp. 2419–2468, Sep. 2021.
- [26] L. Hewing, K. P. Wabersich, M. Menner, and M. N. Zeilinger, "Learning-based model predictive control: Toward safe learning in control," *Annu. Rev. Control, Robot., Auton. Syst.*, vol. 3, no. 1, pp. 269–296, May 2020.
- [27] L. Brunke et al., "Safe learning in robotics: From learning-based control to safe reinforcement learning," *Annu. Rev. Control, Robot., Auto. Syst.*, vol. 5, no. 1, pp. 411–444, May 2022.
- [28] A. Mesbah et al., "Fusion of machine learning and MPC under uncertainty: What advances are on the horizon?" in *Proc. Amer. Control Conf. (ACC)*, Jun. 2022, pp. 342–357.
- [29] L. Hewing, J. Kabzan, and M. N. Zeilinger, "Cautious model predictive control using Gaussian process regression," *IEEE Trans. Control Syst. Technol.*, vol. 28, no. 6, pp. 2736–2743, Nov. 2020.
- [30] D. Piga, M. Forgiione, S. Formentin, and A. Bemporad, "Performance-oriented model learning for data-driven MPC design," *IEEE Control Syst. Lett.*, vol. 3, no. 3, pp. 577–582, Jul. 2019.
- [31] K. P. Wabersich, L. Hewing, A. Carron, and M. N. Zeilinger, "Probabilistic model predictive safety certification for learning-based control," *IEEE Trans. Autom. Control*, vol. 67, no. 1, pp. 176–188, Jan. 2022.
- [32] A. Romero, Y. Song, and D. Scaramuzza, "Actor-critic model predictive control," in *Proc. IEEE Int. Conf. Robot. Autom. (ICRA)*, May 2024, pp. 14777–14784.
- [33] Q. Zhang, W. Pan, and V. Reppa, "Model-reference reinforcement learning for collision-free tracking control of autonomous surface vehicles," *IEEE Trans. Intell. Transp. Syst.*, vol. 23, no. 7, pp. 8770–8781, Jul. 2022.
- [34] C. F. O. da Silva, A. Dabiri, and B. De Schutter, "Integrating reinforcement learning and model predictive control with applications to microgrids," 2024, *arXiv:2409.11267*.
- [35] S. Gros and M. Zanon, "Data-driven economic NMPC using reinforcement learning," *IEEE Trans. Autom. Control*, vol. 65, no. 2, pp. 636–648, Feb. 2020.
- [36] M. Zanon and S. Gros, "Safe reinforcement learning using robust MPC," *IEEE Trans. Autom. Control*, vol. 66, no. 8, pp. 3638–3652, Aug. 2021.
- [37] S. Gros and M. Zanon, "Learning for MPC with stability & safety guarantees," *Automatica*, vol. 146, Dec. 2022, Art. no. 110598.
- [38] B. De Schutter and B. De Moor, "Optimal traffic light control for a single intersection," *Eur. J. Control*, vol. 4, no. 3, pp. 260–276, Jan. 1998.
- [39] T. Bellemans, B. De Schutter, and B. De Moor, "Model predictive control for ramp metering of motorway traffic: A case study," *Control Eng. Pract.*, vol. 14, no. 7, pp. 757–767, Jul. 2006.
- [40] Y. Han, M. Ramezani, A. Hegyi, Y. Yuan, and S. Hoogendoorn, "Hierarchical ramp metering in freeways: An aggregated modeling and control approach," *Transp. Res. C, Emerg. Technol.*, vol. 110, pp. 1–19, Jan. 2020.
- [41] U. Todorović, J. R. D. Frejo, and B. De Schutter, "Distributed MPC for large freeway networks using alternating optimization," *IEEE Trans. Intell. Transp. Syst.*, vol. 23, no. 3, pp. 1875–1884, Mar. 2022.
- [42] S. Liu, A. Sadowska, and B. De Schutter, "A scenario-based distributed model predictive control approach for freeway networks," *Transp. Res. C, Emerg. Technol.*, vol. 136, Mar. 2022, Art. no. 103261.
- [43] E. Smaragdakis, M. Papageorgiou, and E. Kosmatopoulos, "A flow-maximizing adaptive local ramp metering strategy," *Transp. Res. B, Methodol.*, vol. 38, no. 3, pp. 251–270, Mar. 2004.
- [44] J. Chen, W. Lin, Z. Yang, J. Li, and P. Cheng, "Adaptive ramp metering control for urban freeway using large-scale data," *IEEE Trans. Veh. Technol.*, vol. 68, no. 10, pp. 9507–9518, Oct. 2019.
- [45] Y. Wang and M. Papageorgiou, "Real-time freeway traffic state estimation based on extended Kalman filter: A general approach," *Transp. Res. B, Methodol.*, vol. 39, no. 2, pp. 141–167, Feb. 2005.
- [46] M. Davarynejad, A. Hegyi, J. Vrancken, and J. van den Berg, "Motorway ramp-metering control with queuing consideration using Q-learning," in *Proc. 14th Int. IEEE Conf. Intell. Transp. Syst. (ITSC)*, Oct. 2011, pp. 1652–1658.
- [47] A. Fares and W. Gomaa, "Freeway ramp-metering control based on reinforcement learning," in *Proc. 11th IEEE Int. Conf. Control Autom. (ICCA)*, Jun. 2014, pp. 1226–1231.
- [48] E. Ivanjko, D. Koltovska Neč oska, M. Gregurić, M. Vujić, G. Jurković, and S. Mandžuka, "Ramp metering control based on the Q-learning algorithm," *Cybern. Inf. Technol.*, vol. 15, no. 5, pp. 88–97, Apr. 2015.
- [49] S. El-Tantawy, B. Abdulhai, and H. Abdelgawad, "Multiagent reinforcement learning for integrated network of adaptive traffic signal controllers (MARLIN-ATSC): Methodology and large-scale application on downtown Toronto," *IEEE Trans. Intell. Transp. Syst.*, vol. 14, no. 3, pp. 1140–1150, Sep. 2013.
- [50] R. L. Abduljabbar, H. Dia, and P.-W. Tsai, "Development and evaluation of bidirectional LSTM freeway traffic forecasting models using simulation data," *Sci. Rep.*, vol. 11, no. 1, p. 23899, Dec. 2021.
- [51] C. Gu, T. Zhou, and C. Wu, "Deep Koopman traffic modeling for freeway ramp metering," *IEEE Trans. Intell. Transp. Syst.*, vol. 24, no. 6, pp. 6001–6013, Jun. 2023.
- [52] W. Remmerswaal, D. Sun, A. Jamshidnejad, and B. De Schutter, "Combined MPC and reinforcement learning for traffic signal control in urban traffic networks," in *Proc. 26th Int. Conf. Syst. Theory, Control Comput. (ICSTCC)*, Oct. 2022, pp. 432–439.
- [53] D. Sun, A. Jamshidnejad, and B. De Schutter, "A novel framework combining MPC and deep reinforcement learning with application to freeway traffic control," *IEEE Trans. Intell. Transp. Syst.*, vol. 25, no. 7, pp. 6756–6769, Jul. 2024.
- [54] M. Treiber and A. Kesting, *Traffic Flow Dynamics: Data, Models and Simulation*. Berlin, Germany: Springer, 2013.
- [55] M. Papageorgiou, J.-M. Blosseville, and H. Hadj-Salem, "Macroscopic modelling of traffic flow on the boulevard Périphérique in Paris," *Transp. Res. B, Methodol.*, vol. 23, no. 1, pp. 29–47, Feb. 1989.
- [56] A. Messmer and M. Papageorgiou, "METANET: A macroscopic simulation program for motorway networks," *Traffic Eng. Control*, vol. 9, nos. 8–9, pp. 466–470, 1990.
- [57] M. Papageorgiou, I. Papamichail, A. Messmer, and Y. Wang, *Traffic Simulation With METANET*. New York, NY, USA: Springer, 2010, pp. 399–430.

- [58] A. Dabiri and B. B. Kulcsár, "Distributed ramp metering—A constrained discharge flow maximization approach," *IEEE Trans. Intell. Transp. Syst.*, vol. 18, no. 9, pp. 2525–2538, Sep. 2017.
- [59] D. Ngoduy and S. P. Hoogendoorn, "An automated calibration procedure for macroscopic traffic flow models," *IFAC Proc. Volumes*, vol. 36, no. 14, pp. 263–268, Aug. 2003.
- [60] M. G. Lagoudakis, R. Parr, and M. L. Littman, "Least-squares methods in reinforcement learning for control," in *Methods and Applications of Artificial Intelligence* (Methods and Applications of Artificial Intelligence). Berlin, Germany: Springer, 2002, pp. 249–260.
- [61] H. N. Esfahani, A. B. Kordabad, and S. Gros, "Approximate robust NMPC using reinforcement learning," in *Proc. Eur. Control Conf. (ECC)*, Jun. 2021, pp. 132–137.
- [62] L.-J. Lin, "Self-improving reactive agents based on reinforcement learning, planning and teaching," *Mach. Learn.*, vol. 8, no. 3, pp. 293–321, 1992.
- [63] C. Büskens and H. Maurer, "Sensitivity analysis and real-time optimization of parametric nonlinear programming problems," in *Online Optimization of Large Scale Systems*, M. Grötschel, S. O. Krumke, and J. Rambau, Eds., Berlin, Germany: Springer, 2001, pp. 3–16.
- [64] J. Nocedal and S. J. Wright, *Numerical Optimization*, 2nd ed., Berlin, Germany: Springer, 2006.
- [65] D. Ghosh, J. Rahme, A. Kumar, A. Zhang, R. P. Adams, and S. Levine, "Why generalization in RL is difficult: Epistemic POMDPs and implicit partial observability," in *Proc. Adv. Neural Inf. Process. Syst.*, vol. 34, M. Ranzato, A. Beygelzimer, Y. Dauphin, P. Liang, and J. W. Vaughan, Eds., Red Hook, NY, USA: Curran Associates, 2021, pp. 25502–25515.
- [66] J. A. E. Andersson, J. Gillis, G. Horn, J. B. Rawlings, and M. Diehl, "CasADi: A software framework for nonlinear optimization and optimal control," *Math. Program. Comput.*, vol. 11, no. 1, pp. 1–36, Mar. 2019.
- [67] A. Wächter and L. T. Biegler, "On the implementation of an interior-point filter line-search algorithm for large-scale nonlinear programming," *Math. Program.*, vol. 106, no. 1, pp. 25–57, 2006.
- [68] A. D. Spiliopoulou, D. Manolis, I. Papamichail, M. Papageorgiou, J. Wu, and X. Jin, "Queue management techniques for metered freeway on-ramps," *Transp. Res. Rec.*, vol. 2178, no. 1, pp. 40–48, Jan. 2010.
- [69] J. Wu, X.-Y. Chen, H. Zhang, L.-D. Xiong, H. Lei, and S.-H. Deng, "Hyperparameter optimization for machine learning models based on Bayesian optimization," *J. Electron. Sci. Technol.*, vol. 17, pp. 26–40, Mar. 2019.
- [70] L. Franceschi, P. Frasconi, S. Salzo, R. Grazzi, and M. Pontil, "Bilevel programming for hyperparameter optimization and meta-learning," in *Proc. 35th Int. Conf. Mach. Learn.*, vol. 80, J. Dy and A. Krause, Eds., Jul. 2018, pp. 1568–1577.



Filippo Airaldo received the B.Sc. and M.Sc. degrees from the Polytechnic University of Turin, Italy, in 2017 and 2019, respectively. He is currently pursuing the Ph.D. degree with Delft Center for Systems and Control, Delft University of Technology, The Netherlands.

His research interests include model predictive control, reinforcement learning, and other machine learning techniques, and in particular in their combination with learning-based control.



Bart De Schutter (Fellow, IEEE) received the Ph.D. degree (summa cum laude) in applied sciences from KU Leuven, Belgium, in 1996.

He is currently a Full Professor with Delft Center for Systems and Control, Delft University of Technology, The Netherlands. His research interests include learning and optimization-based control of large-scale and hybrid systems, multilevel, and distributed control, with applications in intelligent transportation and smart energy systems.

Prof. De Schutter is a Senior Editor of IEEE

TRANSACTIONS ON INTELLIGENT TRANSPORTATION SYSTEMS.



Azita Dabiri received the Ph.D. degree from the Automatic Control Group, Chalmers University of Technology, in 2016. She was a Post-Doctoral Researcher with the Department of Transport and Planning, Delft University of Technology (TU Delft), from 2017 to 2019. In 2019, she received an ERCIM Fellowship and a Marie Curie Individual Fellowship, which allowed her to perform research with the Norwegian University of Technology (NTNU), Norway, as a Post-Doctoral Researcher, from 2019 to 2020, before joining Delft Center for

Systems and Control, TU Delft, as an Assistant Professor. Her research interests include the area of integration of model-based and learning-based control and its applications in transportation networks.