

FLOW CONTROL IN CONNECTION-ORIENTED NETWORKS: A TIME-VARYING SAMPLING PERIOD SYSTEM CASE STUDY

PRZEMYSŁAW IGNACIUK AND ANDRZEJ BARTOSZEWICZ

In this paper congestion control problem in connection-oriented communication network with multiple data sources is addressed. In the considered network the feedback necessary for the flow regulation is provided by means of management units, which are sent by each source once every M data packets. The management units, carrying the information about the current network state, return to their origin round trip time (RTT) after they were sent. Since the source rate is adjusted only at the instant of the control units arrival, the period between the transfer speed modifications depends on the flow rate RTT earlier, and consequently varies with time. A new, nonlinear algorithm combining the Smith principle with the proportional controller with saturation is proposed. Conditions for data loss elimination and full resource utilisation are formulated and strictly proved with explicit consideration of irregularities in the feedback information availability. Subsequently, the algorithm robustness with respect to imprecise propagation time estimation is demonstrated. Finally, a modified strategy implementing the feed-forward compensation is proposed. The strategy not only eliminates packet loss and guarantees the maximum resource utilisation, but also decreases the influence of the available bandwidth on the queue length. In this way the data transfer delay jitter is reduced, which helps to obtain the desirable Quality of Service (QoS) in the network.

Keywords: congestion control, connection-oriented networks, sampled data systems, variable sampling period

AMS Subject Classification: 93A30, 93C57, 94A20

1. INTRODUCTION

Modern telecommunication networks are dynamic systems which require real-time control schemes for the flow regulation. As it is convenient to send one control unit every M data packets (and in this way place a direct limit on the amount of the exchanged management information with respect to the user traffic), the control data in such networks is usually available at irregularly spaced time instants. In this paper we present an efficient solution to the flow control problem in the multi-source connection-oriented networks, which provide feedback information aperiodically. On the contrary to the results published in the past, the variable (input rate dependent)

sampling period is explicitly taken into account in the algorithm design and its properties derivation.

The difficulty of the flow control in the networks mentioned above, apart from the dynamically changing period of control signal availability, is mainly caused by long propagation delays in the system. If congestion occurs at a specific node, information about this condition must be conveyed to all the sources transmitting data through that node. Transferring this information involves feedback propagation delays. After the information has been received by a particular source, it can be used to adjust the rate of this source. However, the adjusted flow rate will start to affect the congested node only after forward propagation delay.

The flow rate control in wide area networks, the example of which are Bandwidth on Demand (BoD) satellite or Asynchronous Transfer Mode (ATM) terrestrial communication systems, has recently been studied in several papers. A valuable survey of earlier congestion control mechanisms is given in [8]. Furthermore, Izmailov [6] considered a single connection controlled by a linear regulator whose output signal is generated according to the several states of the buffer measured at different time instants. The asymptotic stability, nonoscillatory system behaviour and locally optimal rate of convergence have been proved. Chong et al. proposed and thoroughly studied the performance of a simple queue length based flow control algorithm with a dynamic queue threshold adjustment [3]. Kulkarni and Li [9] modeled the data transfer fluctuations (caused by changes in propagation delay and time-varying source activity periods) with random variables and studied their influence on the system performance under the control of parsimonious (binary) and multi-valued feedback mechanisms. Lengliz and Kamoun [11] introduced a proportional plus derivative (PD) controller, which is computationally efficient and can be easily implemented in connection-oriented networks. Imer et al. [5] gave a brief, excellent tutorial exposition of the congestion control problem and presented new stochastic and deterministic control algorithms. Another interesting approach to the problem of the flow rate control in communication networks has been proposed by Quet et al. In paper [16], the authors considered a single bottleneck multi-source network and applied minimization of an H-infinity norm to the design of a flow rate controller. The proposed controller guarantees stability robustness to uncertain and time-varying propagation delays in various channels. Adaptive control strategies for flow regulation in time-delay systems have been proposed by Laberteaux et al. [10]. Their strategies reduce convergence time and improve queue length management. Also, a neural network controller for wide area networks has recently been proposed. Jagannathan and Talluri [7] showed that their neural network controller can guarantee stability of the closed loop system and the desired QoS.

Due to the significant propagation delays, which are critical for the closed loop performance, several researchers applied the Smith principle [18] to control the flow of data in communication networks [1, 2, 4, 12–15]. In paper [12], Mascolo considered a single connection congestion control problem in a general packet switching network. He used the deterministic fluid model approximation of packet flow and exploited transfer functions to describe the network dynamics. The designed continuous time controller was applied to the Available Bit Rate (ABR) traffic control in ATM net-

work and compared with the Explicit Rate Indication for Congestion Avoidance (ERICA) standard. The same author extended the idea of the Smith prediction to control the network supporting multiple data flows with different propagation delays in paper [13]. The proposed control algorithm guarantees no cell loss, full and equal network utilisation, and ensures exponential convergence of queue levels to stationary values without oscillations or overshoots. Gómez-Stern et al. further studied the flow control using the Smith principle [4]. They proposed a continuous time proportional-integral (PI) controller which helps reduce the average queue level and its sensitivity to the available bandwidth. On the other hand, the application of the Smith principle for satellite networks was considered in [15]. In that work, similarly as in [4], the saturation issues in the system with proportional continuous time controller were handled using anti-wind up techniques. In recent paper [14] Mascolo demonstrated that also the TCP flow control mechanism implements the Smith predictor to handle the congestion. The result presented in [14] was supplemented with the analysis of the performance of the Smith predictor based solutions as compared with the traditional proportional-integral-derivative (PID) controllers. It was shown that in the time delay systems the stability requirements significantly limit the dynamics of the PID-based schemes and the Smith principle provides faster reaction to the varying networking conditions. A nonlinear algorithm exploiting the idea of the Smith prediction for the flow regulation in time-delay systems was suggested in [2]. The described continuous time control mechanism guarantees congestion alleviating features and full resource usage even though the propagation delays in the multi-source network can be determined only with limited degree of accuracy. Moreover, since in real networks the feedback information is usually available only at discrete time instants, discrete-time linear controllers were also developed for telecommunication systems [1, 17].

In this paper, the flow control in connection-oriented communication networks is considered. Our approach is similar to that introduced in [1, 4, 12, 13, 15], however as opposed to those papers, we propose a nonlinear control strategy. Moreover, in contrast to [2, 4, 12, 13, 15], where continuous time control schemes were elaborated and [1], where discrete-time controller with constant sampling period was proposed, in this paper, an algorithm, which directly takes into account irregularities in the feedback information availability, is designed. Since the proposed solution does not rely on the continuous feedback information availability nor maintaining synchronisation of constant sampling period (which is a serious challenge in multi-source systems), it is more scalable and requires less control effort than other flow regulation schemes presented earlier in literature. In fact, to our best knowledge, the described class of irregularities in the feedback provision, caused by sending one management unit every M data packets, although present in the existing networks and likely to appear in future solutions, yet has never been explicitly considered in the design of the flow controllers. The proposed strategy combines the Smith principle with the proportional controller with saturation. Our approach guarantees full bottleneck node link utilisation and no packet loss in the network. As a result, the need of data retransmission is eliminated and the maximum throughput is achieved. Since the Smith predictor may be sensitive with respect to imprecisely estimated delay

times, we thoroughly study the effect of a possible mismatch between the actual and measured values of RTT and demonstrate that the proposed strategy maintains its favourable features even though the real delays in the network differ from the calculated ones. Finally, in order to decrease the influence of the available bandwidth on the packet queue length, we implement the concept of the feed-forward bandwidth compensation into the control mechanism. The modified approach not only ensures full bottleneck node link utilisation and no data loss in the network, but also helps reduce the data transfer delay variation and thus obtain better QoS in the network. The transmission rates generated by both strategies are nonnegative and bounded. These properties allow for a direct implementation of the proposed control schemes in real network environment.

The remainder of this paper is organised as follows. The model of the network used throughout the paper is thoroughly described in Section 2. Then, nonlinear flow control strategies and properties of the system with unisochronic feedback are discussed. First, in Section 3.1, the algorithm working under the assumption that round trip times of all connections are known exactly is introduced. The stable operation of the controller is guaranteed, even though the sources adapt their rate at irregular (input dependent) time instants. Section 3.2 addresses the robustness issues related to the imprecise propagation delay calculation in the analysed network. Afterwards, in Section 3.3, a modified strategy with feed-forward compensation is presented. It is shown that the nonlinear controller implementing the Smith predictor can assure insensitivity of the steady-state queue length with respect to the available bandwidth. In consequence, the data transmission delay jitter is decreased, which helps to obtain favourable quality of service in the system and makes the network suitable for multimedia traffic. Subsequently, the properties of the proposed strategies are verified by a simulation example presented in Section 4. Finally, Section 5 comprises the conclusions of the paper.

2. NETWORK MODEL

The connection-oriented network considered in this paper consists of data sources, intermediate switches and destinations. The data transmitted by the sources passes through a number of nodes, which operate in the store-and-forward mode without traffic prioritization, to be finally delivered to its destination. However, somewhere on the transmission path a switch is encountered, whose output link cannot handle the incoming flow. Consequently, congestion occurs and packets, which constitute the data stream, accumulate in the buffer allocated for that link. At the same time, we assume that the sources under consideration are not subject to data capacity limitation, i.e. they always have enough data to transmit. Therefore, they will aggressively compete for the network bandwidth and buffer overflow can be prevented only through an appropriate input rate adjustment.

The sources adapt the transmission speed at the instant of the control unit arrival according to the command calculated by the bottleneck switch. The time period between the arrivals of consecutive management units depends on the emission rate RTT earlier, which, in turn, changes with the variations of the network

The queue length at any instant of time depends on the data arrival speed and on consumed bandwidth h . Therefore, for any $t \geq 0$ the length of the queue at the node may be expressed as

$$x(t) = \sum_{j=1}^n \int_0^t a_j(\tau - T_{fj}) d\tau - \int_0^t h(\tau) d\tau. \quad (2)$$

The j th source rate is determined by the controller placed at the bottleneck switch. Let us denote by $b_j(t)$ the rate calculated by the controller and sent for the j th source at the instant of the management unit passing through the node. Assuming that the sources begin transmission at the time $t = 0$ at the rate established in the connection set-up phase, the following is true

$$\forall_j \forall_{t < 0} a_j(t) = 0 \text{ and } \forall_j \forall_{t \geq 0} a_j(t) = b_j(t - T_{bj}). \quad (3)$$

Since the signal $b_j(t)$ constitutes a vital part of the proposed control scheme, its proper definition will be given together with the description of the flow regulation strategy in the subsequent section.

Suppose that before time instant $t = 0$ the bottleneck buffer was empty, i.e. $x(t < 0) = 0$. Then, as a consequence of (3), no packets arrive at the congested node before $T_{f \min} = \min_{j=1,2,\dots,n}(T_{fj})$ and for any time instant smaller than or equal to $T_{f \min}$ the queue length is zero, i.e. $x(t \leq T_{f \min}) = 0$. Let us denote by $t_{j,k}$ the k th moment of time ($k = 1, 2, \dots$) when the control unit belonging to the j th virtual connection data flow arrives back at the source j . Since the sources adjust the transmission speed only when management unit returns with the network feedback incorporated, then

$$\forall_{t \in [t_{j,k}; t_{j,k+1})} a_j(t) = a_j(t_{j,k}) = b_j(t_{j,k} - T_{bj}) = \text{const}. \quad (4)$$

The first to be transferred by any source is a control unit so that the information about the current network state could be received at the data origin as quickly as possible. As the sources begin transmission at time instant $t = 0$, then for $k = 1$ we have $t_{j,1} = RTT_j$. Furthermore, since the control units are sent every M data packets and not less frequently than the maximum control period T_C one after another, $t_{j,k+1}$ is specified by the following relation

$$t_{j,k+1} = \min(t_{j,k} + \beta_{j,k}, t_{j,k} + T_C) \quad (5)$$

where $\beta_{j,k}$ can be determined from the equation given below

$$\int_{t_{j,k}}^{t_{j,k} + \beta_{j,k}} a_j(\tau - RTT_j) d\tau = M. \quad (6)$$

Definitions (5) and (6) make sense only for nonnegative rates $a_j(t)$. Clearly, any control algorithm should be constructed in such a way that this condition is satisfied for every $a_j(t)$. Finally, let us point out that equations (5) and (6) present the main novelty of the network model used in this paper. These two equations explicitly account for the time-varying, input rate dependent sampling period in the considered system.

3. CONTROL ALGORITHM

In this section we present the control algorithm, which ensures efficient network usage in a multi-source environment, even though the feedback information available for rate adjustment at the sources is received aperiodically. The strategy guarantees that:

- (i) data is not lost due to congestion, which minimises network overhead corresponding to retransmissions;
- (ii) there is always some data ready for transmission in the congested node buffer so that the available bandwidth at the node output link is entirely utilised.

3.1. Principal control scheme

In the sequel we propose a nonlinear flow regulating strategy and demonstrate its basic properties. Rate $b_j(t)$ sent by the controller for each source at the instant of a control unit passing through the bottleneck node is determined by the equations given below

$$\begin{aligned} \forall_{t < T_{fj}} \quad b_j(t) &= 0 \\ \forall_{t \geq T_{fj}} \quad \forall_{t \in [t_{j,k} - T_{bj}; t_{j,k+1} - T_{bj})} \quad b_j(t) &= b_j(t_{j,k} - T_{bj}) = \frac{1}{n} a(t_{j,k} - T_{bj}) \end{aligned} \quad (7)$$

Total rate $a(t)$ is calculated from the following relation

$$a(t) = \begin{cases} 0, & \text{if } W(t) < 0 \\ W(t), & \text{if } 0 \leq W(t) \leq a_{\max} \\ a_{\max}, & \text{if } W(t) > a_{\max} \end{cases} \quad (8)$$

where $a_{\max} > 0$ denotes the upper saturation limit. We define function $W(t)$ as

$$W(t) = K \left[x_d - x(t) - \sum_{j=1}^n \int_{t-RTT_j}^t b_j(\tau) d\tau \right] \quad (9)$$

where $K > 0$ is the controller gain and $x_d > 0$ is the demand queue length. The integral term in equation (9) is responsible for the Smith prediction and it represents the in-flight data. The basic control definition expressed by (9) is similar to those already proposed in papers [1, 4, 12, 13, 15], however the nonlinearity introduced by (8) and explicit consideration of the time-varying control period make our overall strategy essentially different from the previously proposed solutions.

Further in this section we formulate the theorem indicating the amount of memory, which needs to be reserved at the bottleneck node for the data storage so that no packet is discarded irrespective of the available bandwidth changes.

Theorem 1. If sources transmit data according to the conditions formulated by (3)–(9), then the queue length at the bottleneck node does not exceed the value given by the following inequality

$$\forall_{t \geq 0} x(t) \leq x_d + a_{\max} T_C. \quad (10)$$

Proof. First, notice that, as a consequence of the source transfer speed adjustment at discrete moments of time, the total arrival rate at the bottleneck node may also change only at discrete time instants, which will be denoted by θ_m ($m = 1, 2, \dots$). The first such modification coincides with the arrival of the initial packet belonging to the flow with the shortest forward delay, so for $m = 1$ we have $\theta_1 = T_{f \min}$. Interval

$$\alpha_m = \theta_{m+1} - \theta_m \quad (11)$$

between any two consecutive potential changes of the total incoming rate at the congested switch is subject to the constraint $0 \leq \alpha_m \leq T_C$. The zero-length interval reflects the case, when the modification of the transmission speed, which occurred at two or more sources, influences the aggregate rate at the switch at the same moment of time, and the upper bound is the maximum control unit inter-arrival period.

The bottleneck link buffer is empty until the first packets arrive at the switch. Let us denote the queue length at the time instant θ_m by $x_m = x(\theta_m)$. For $m = 1$ and $\theta_1 = T_{f \min}$ we can write $x_1 = x(\theta_1) = 0 < x_d + a_{\max} T_C$. Therefore, the proposition holds for any moment of time $t \leq T_{f \min}$.

Let us consider some $m > 1$ and the queue length at the time instant $t \in [\theta_m, \theta_{m+1})$

$$x(t) = x(\theta_m) + \sum_{j=1}^n \int_{\theta_m}^{\theta_m + \delta} a_j(\tau - T_{fj}) d\tau - \int_{\theta_m}^{\theta_m + \delta} h(\tau) d\tau \quad (12)$$

where $t = \theta_m + \delta$, $\delta \in [0, \alpha_m)$ and α_m is defined by (11). Using (4), we can rewrite (12) as

$$\begin{aligned} x(t) &= x(\theta_m) + \sum_{j=1}^n \int_{\theta_m}^{\theta_m + \delta} b_j(\tau - RTT_j) d\tau - \int_{\theta_m}^{\theta_m + \delta} h(\tau) d\tau \\ &= x(\theta_m) + \sum_{j=1}^n \int_{\theta_m - RTT_j}^{\theta_m + \delta - RTT_j} b_j(\tau) d\tau - \int_{\theta_m}^{\theta_m + \delta} h(\tau) d\tau. \end{aligned} \quad (13)$$

In order to analyze the queue length variations in time interval $[\theta_m, \theta_{m+1})$, we examine the behavior of function W . We will consider two cases: first, the situation when $W(\theta_m) \geq 0$, and, afterwards, the circumstances when $W(\theta_m) < 0$.

Case 1: We analyze the situation when $W(\theta_m) \geq 0$. From the definition of W , we get

$$x(t) \leq x_d + \sum_{j=1}^n \int_{\theta_m - RTT_j}^{\theta_m + \delta - RTT_j} b_j(\tau) d\tau - \int_{\theta_m}^{\theta_m + \delta} h(\tau) d\tau. \quad (14)$$

Since the total arrival rate at the switch does not change in interval $[\theta_m, \theta_{m+1})$, we may evaluate the first integral in (14)

$$x(t) \leq x_d + \delta \sum_{j=1}^n b_j(\theta_m - RTT_j) - \int_{\theta_m}^{\theta_m + \delta} h(\tau) d\tau. \quad (15)$$

The utilised bandwidth is always nonnegative and δ is upper-bounded by T_C . Then, taking into account the fact that $b_j(t) \leq a_{\max}/n$, we may write

$$x(t) \leq x_d + \delta \cdot a_{\max} - 0 \leq x_d + a_{\max} T_C. \quad (16)$$

This ends the first part of the proof.

Case 2: Now, let us examine the situation when $W(\theta_m) < 0$. First we find the last moment $t^* < \theta_m$ when signal W was greater than zero. It should be stressed at this point, that since the control unit emission at the sources and rate generation at the congested node are not synchronized, t^* does not have to coincide with any of the θ_m time instants. According to (7), $W(t \leq T_{f \min}) = K(x_d - 0 - 0 - 0) = Kx_d > 0$. In consequence, the first moment, when signal W may attain a value smaller than zero, is greater than $T_{f \min}$, and instant t^* actually exists. The value of $W(t^*)$ satisfies the following inequality

$$W(t^*) = K \left[x_d - x(t^*) - \sum_{j=1}^n \int_{t^* - RTT_j}^{t^*} b_j(\tau) d\tau \right] > 0. \quad (17)$$

After the term rearrangement, we obtain

$$x(t^*) < x_d - \sum_{j=1}^n \int_{t^* - RTT_j}^{t^*} b_j(\tau) d\tau. \quad (18)$$

The queue length at a time instant $t \in [\theta_m, \theta_{m+1})$ can be expressed as

$$x(t) = x(t^*) + \sum_{j=1}^n \int_{t^* - RTT_j}^{t - RTT_j} b_j(\tau) d\tau - \int_{t^*}^t h(\tau) d\tau. \quad (19)$$

Substituting (18) for $x(t^*)$, we get

$$\begin{aligned} x(t) &< x_d - \sum_{j=1}^n \int_{t^* - RTT_j}^{t^*} b_j(\tau) d\tau + \sum_{j=1}^n \int_{t^* - RTT_j}^{t - RTT_j} b_j(\tau) d\tau - \int_{t^*}^t h(\tau) d\tau \\ &\leq x_d + \sum_{j=1}^n \int_{t^*}^t b_j(\tau) d\tau - \int_{t^*}^t h(\tau) d\tau. \end{aligned} \quad (20)$$

Since the transfer speed generated by the controller in time interval $(t^*, t]$ is equal to 0, and its assignment can be delayed by T_C , integral in (20) can be estimated as

$$\sum_{j=1}^n \int_{t^*}^t b_j(\tau) d\tau \leq \sum_{j=1}^n \int_{t^*}^{t^* + T_C} b_j(\tau) d\tau \leq a_{\max} T_C. \quad (21)$$

Function $h(t)$ is always nonnegative, so we can state that the queue length at time instant t , as given by (20), will be limited by the value given below

$$x(t) < x_d + \sum_{j=1}^n \int_{t^*}^t b_j(\tau) d\tau - \int_{t^*}^t h(\tau) d\tau \leq x_d + a_{\max} T_C. \quad (22)$$

This concludes the proof. $\square$

Full link utilisation, as formulated by (ii), requires the presence of packets in the bottleneck node buffer at any instant of time. In other words, if the queue length is greater than zero, then the total available bandwidth of the congested link is consumed. The theorem presented below shows how the demand queue length should be selected so that entire bandwidth at the output connection is used for the data traffic.

Theorem 2. If sources transmit data according to the conditions formulated by (3)–(9), the maximum rate $a_{\max} > d_{\max}$ and the demand value of the queue length satisfies the following inequality

$$x_d > a_{\max} \left(\sum_{j=1}^n \frac{1}{n} R T T_j + \frac{1}{K} + T_C \right) \quad (23)$$

then for any $t > T_{f \max} + T_C + T_{\max}$, where $T_{f \max} = \max_{j=1,2,\dots,n} (T_{fj})$ and $T_{\max} = (x_d + a_{\max} T_C) / (a_{\max} - d_{\max})$, the queue length is always greater than zero.

Proof. The theorem assumption implies that we deal with time instants $\theta_m > T_{f \max} + T_C + T_{\max}$. Considering some $m > 1$ and the value of signal W at the moment of the node input rate modification θ_m , we may distinguish two cases: the situation when $W(\theta_m) < a_{\max}$, and the circumstances when $W(\theta_m) \geq a_{\max}$.

Case 1: We consider the situation when $W(\theta_m) < a_{\max}$. Directly from the definition of function W , we obtain

$$W(\theta_m) = K \left[x_d - x(\theta_m) - \sum_{j=1}^n \int_{\theta_m - R T T_j}^{\theta_m} b_j(\tau) d\tau \right] < a_{\max}. \quad (24)$$

The maximum rate established by the controller is equal to $a_{\max}$, so $b_j(t) \leq a_{\max}/n$ and

$$x(\theta_m) > x_d - \frac{a_{\max}}{K} - a_{\max} \sum_{j=1}^n \frac{1}{n} R T T_j. \quad (25)$$

Using assumption (23), we obtain

$$\begin{aligned} x(\theta_m) &> a_{\max} \left(\sum_{j=1}^n \frac{1}{n} R T T_j + \frac{1}{K} \right) + a_{\max} T_C \\ &\quad - \frac{a_{\max}}{K} - a_{\max} \sum_{j=1}^n \frac{1}{n} R T T_j = a_{\max} T_C. \end{aligned} \quad (26)$$

Let us examine the queue length at some time instant $t \in [\theta_m, \theta_{m+1})$, as defined by (13). The minimum rate, which can be assigned for each source, is zero and the maximum available bandwidth equals $d_{\max}$. Then, applying (26) to (13), we get

$$x(t) > a_{\max}T_C + 0 - d_{\max}\delta > 0 \quad (27)$$

which completes the first part of the proof.

Case 2: Now, let us study the situation when $W(\theta_m) \geq a_{\max}$. First, we find the last moment $t^* < \theta_m$ when signal W was smaller than $a_{\max}$. It comes from Theorem 1 that the queue length never exceeds the value of $x_d + a_{\max}T_C$. At the same time, packet depletion rate is limited by $d_{\max}$. Thus, the maximum period of time $T_{\max}$, during which the controller may continuously set rate $a_{\max}$ for the sources, is determined by the following identity

$$T_{\max} = (x_d + a_{\max}T_C) / (a_{\max} - d_{\max}). \quad (28)$$

Therefore, instant t^* exists. Since t^* is the last instant, when signal W was smaller than $a_{\max}$ and the actual rate assignment could be delayed by not more than T_C ,

$$t^* \geq \theta_m - (T_{\max} + T_C). \quad (29)$$

From the theorem assumption it also comes that $t^* > T_{f \max} + T_{\max} + T_C - (T_{\max} + T_C) = T_{f \max}$.

The value of $W(t^*)$ satisfies the inequality given below

$$W(t^*) = K \left[x_d - x(t^*) - \sum_{j=1}^n \int_{t^* - RTT_j}^{t^*} b_j(\tau) d\tau \right] < a_{\max}. \quad (30)$$

Following similar reasoning as presented in (24)–(27) we arrive at

$$x(t^*) > x_d - \frac{a_{\max}}{K} - \sum_{j=1}^n \int_{t^* - RTT_j}^{t^*} b_j(\tau) d\tau > 0. \quad (31)$$

The queue length at some time instant $t \in [\theta_m, \theta_{m+1})$ may be expressed as

$$x(t) = x(t^*) + \sum_{j=1}^n \int_{t^*}^t a_j(\tau - T_{fj}) d\tau - \int_{t^*}^t h(\tau) d\tau \quad (32)$$

where $t = \theta_m + \delta$ and $\delta \in [0, \alpha_m)$. The summation term in expression (32) represents the amount of data, which was delivered to the bottleneck node in the time interval $[t^*, t)$. Since $a_j(t) = b_j(t - T_{bj})$, we may rewrite the integral into the form given below

$$x(t) = x(t^*) + \sum_{j=1}^n \int_{t^* - RTT_j}^{t - RTT_j} b_j(\tau) d\tau - \int_{t^*}^t h(\tau) d\tau. \quad (33)$$

Applying result (31) for the queue length at time instant t^* , we get

$$x(t) > x_d - \frac{a_{\max}}{K} - \sum_{j=1}^n \int_{t^* - RTT_j}^{t^*} b_j(\tau) d\tau + \sum_{j=1}^n \int_{t^* - RTT_j}^{t - RTT_j} b_j(\tau) d\tau - \int_{t^*}^t h(\tau) d\tau. \quad (34)$$

Performing algebraic manipulations on the integrals in (34), we get

$$x(t) > x_d - \frac{a_{\max}}{K} - \sum_{j=1}^n \int_{t - RTT_j}^t b_j(\tau) d\tau + \sum_{j=1}^n \int_{t^*}^t b_j(\tau) d\tau - \int_{t^*}^t h(\tau) d\tau. \quad (35)$$

Since the maximum rate assigned for each source is always limited by $a_{\max}/n$, the first summation term in (35) can be estimated as follows

$$-\sum_{j=1}^n \int_{t - RTT_j}^t b_j(\tau) d\tau \geq -a_{\max} \sum_{j=1}^n \frac{1}{n} RTT_j. \quad (36)$$

Transfer speed value $b_j(t)$ assigned for source j at the time instant of the management unit passing at $t^* > T_{f \max}$ or right before t^* would pertain for the duration of $\eta_j \leq T_C$ seconds (until the arrival of the subsequent control unit). Afterwards, the rate for each source was set to $a_{\max}/n$. Thus, the second summation term in (35) satisfies the following inequality

$$\sum_{j=1}^n \int_{t^*}^t b_j(\tau) d\tau \geq \sum_{j=1}^n \int_{t^* + \eta_j}^t b_j(\tau) d\tau \geq a_{\max} [t - (t^* + T_C)]. \quad (37)$$

Applying results (36) and (37) to formula (35), we arrive at

$$x(t) > x_d - \frac{a_{\max}}{K} - a_{\max} \sum_{j=1}^n \frac{1}{n} RTT_j + a_{\max} [t - (t^* + T_C)] - \int_{t^*}^t h(\tau) d\tau. \quad (38)$$

Using the theorem assumption (23), we obtain

$$x(t) > a_{\max} T_C + a_{\max} [t - (t^* + T_C)] - \int_{t^*}^t h(\tau) d\tau = a_{\max} (t - t^*) - \int_{t^*}^t h(\tau) d\tau. \quad (39)$$

The maximum bandwidth, which we may expect to be utilised, is equal to $d_{\max}$. Then, according to the assumption $a_{\max} > d_{\max}$ and the fact that $t > t^*$, we may state that

$$x(t) > a_{\max} (t - t^*) - d_{\max} (t - t^*) = (a_{\max} - d_{\max}) (t - t^*) > 0. \quad (40)$$

This ends the proof of Theorem 2. $\square$

3.2. Robustness analysis

As the Smith predictor may be sensitive to imprecise RTT determination, in this Section we study how possible differences between the real delays existing in the

network and those estimated by the controller in the connection set-up phase, further denoted by $\overline{RTT}_j$, may influence the flow regulation process.

Equations (7) and (8) remain valid. However, in order to take into account discrepancies in the delay parameter values, formula (9) needs to be updated as follows

$$W(t) = K \left[x_d - x(t) - \sum_{j=1}^n \int_{t-\overline{RTT}_j}^t b_j(\tau) d\tau \right] \quad (41)$$

where $\overline{RTT}_j > 0$ is the round trip time of the j th flow as measured by the controller in the connection start-up phase.

Theorems 1 and 2, which were formulated for the network model from the previous section no longer hold. Still though, we can modify the requirements for the buffer space allocation so that the favourable properties of the considered strategy defined by (i) and (ii) will be preserved.

Theorem 3. If sources transmit data according to the conditions formulated by (3)–(9) with signal $W(t)$ defined by (41), then the queue length at the bottleneck node is upper-bounded by the following limit

$$\forall_{t \geq 0} x(t) \leq x_d + a_{\max} T_C + \Delta_{\max} \quad (42)$$

where

$$\Delta_{\max} = \frac{a_{\max}}{n} \sum_{j: RTT_j > \overline{RTT}_j} (RTT_j - \overline{RTT}_j). \quad (43)$$

Proof. The first packets arrive at the switch at $t = T_{f \min}$. Denoting the queue length at the instant θ_m by $x_m = x(\theta_m)$, for $m = 1$ and $\theta_1 = T_{f \min}$ we can write $x_1 = x(\theta_1) = 0 < x_d + a_{\max} T_C + \Delta_{\max}$.

Let us consider the value of signal W at time instant $t = \theta_m$ for some $m > 1$. Two situations are possible: the first occurs when $W(\theta_m) \geq 0$, and the other when $W(\theta_m) < 0$.

Case 1: Investigation of the case

$$W(\theta_m) = K \left[x_d - x(\theta_m) - \sum_{j=1}^n \int_{\theta_m - \overline{RTT}_j}^{\theta_m} b_j(\tau) d\tau \right] \geq 0 \quad (44)$$

leads to the following

$$x(\theta_m) \leq x_d - \sum_{j=1}^n \int_{\theta_m - \overline{RTT}_j}^{\theta_m} b_j(\tau) d\tau \leq x_d. \quad (45)$$

The total arrival rate at the switch is constant in interval $[\theta_m, \theta_{m+1})$. At the same time, the utilised bandwidth is always nonnegative. Since $\delta \leq T_C$ and individual source rate $b_j(t)$ cannot exceed $a_{\max}/n$, we may estimate the queue length at time instant $t = \theta_m + \delta$, given by formula (13), in the following way

$$x(t) \leq x_d + a_{\max} \delta - 0 \leq x_d + a_{\max} T_C + \Delta_{\max}. \quad (46)$$

Case 2: Now, let us examine the circumstances, when $W(\theta_m) < 0$. First, we seek for the last instant $t^* < \theta_m$ when signal W was nonnegative. Since $W(t \leq T_{f \min}) = K(x_d - 0 - 0 - 0) = Kx_d > 0$, the first moment, when W may attain a value smaller than zero is greater than $T_{f \min}$. The value of $W(t^*)$ satisfies the following inequality

$$W(t^*) = K \left[x_d - x(t^*) - \sum_{j=1}^n \int_{t^* - \overline{RTT}_j}^{t^*} b_j(\tau) d\tau \right] > 0. \quad (47)$$

Then, the queue length at the moment t^* can be estimated by the relation given below

$$x(t^*) < x_d - \sum_{j=1}^n \int_{t^* - \overline{RTT}_j}^{t^*} b_j(\tau) d\tau. \quad (48)$$

Substituting (48) for $x(t^*)$ in (19), we obtain

$$\begin{aligned} x(t) &< x_d - \sum_{j=1}^n \int_{t^* - \overline{RTT}_j}^{t^*} b_j(\tau) d\tau + \sum_{j=1}^n \int_{t^* - \overline{RTT}_j}^{t - \overline{RTT}_j} b_j(\tau) d\tau - \int_{t^*}^t h(\tau) d\tau \\ &= x_d - \sum_{j=1}^n \int_{t^* - \overline{RTT}_j}^{t^*} b_j(\tau) d\tau + \sum_{j=1}^n \int_{t^* - \overline{RTT}_j}^{t^* - \overline{RTT}_j} b_j(\tau) d\tau \\ &\quad + \sum_{j=1}^n \int_{t^* - \overline{RTT}_j}^{t - \overline{RTT}_j} b_j(\tau) d\tau - \int_{t^*}^t h(\tau) d\tau. \end{aligned} \quad (49)$$

Let us define a function $\Delta(t)$

$$\Delta(t) = \sum_{j=1}^n \int_{t - \overline{RTT}_j}^{t - \overline{RTT}_j} b_j(\tau) d\tau \quad (50)$$

Its values fall within the range $[-\Delta_{\min}, \Delta_{\max}]$, where $\Delta_{\min}$ and $\Delta_{\max}$ are positive real numbers

$$\begin{aligned} \Delta_{\min} &= -\frac{a_{\max}}{n} \sum_{j: RTT_j < \overline{RTT}_j} (RTT_j - \overline{RTT}_j) \\ \Delta_{\max} &= \frac{a_{\max}}{n} \sum_{j: RTT_j > \overline{RTT}_j} (RTT_j - \overline{RTT}_j). \end{aligned} \quad (51)$$

With this notation we get

$$x(t) < x_d - \sum_{j=1}^n \int_{t^* - \overline{RTT}_j}^{t^*} b_j(\tau) d\tau + \Delta(t^*) + \sum_{j=1}^n \int_{t^* - \overline{RTT}_j}^{t - \overline{RTT}_j} b_j(\tau) d\tau - \int_{t^*}^t h(\tau) d\tau. \quad (52)$$

On the basis of the calculations presented in (20)–(22) and the fact that $\Delta(t) \leq \Delta_{\max}$, we conclude

$$x(t) < x_d + a_{\max} T_C + \Delta_{\max}.$$

This ends the proof. $\square$

In the subsequent part of this section we present the theorem, which shows that if the bottleneck node buffer is selected properly, then all of the bandwidth available at the congested link will be consumed even though the delay times are determined at the controller only with limited accuracy.

Theorem 4. If sources transmit data according to the conditions formulated by (3)–(9), signal $W(t)$ is defined by (41), the maximum rate $a_{\max} > d_{\max}$ and the demand value of the queue length satisfies the following inequality

$$x_d > a_{\max} \left(\sum_{j=1}^n \frac{1}{n} R T T_j + \frac{1}{K} + T_C \right) + \Delta_{\min} \quad (53)$$

then for any $t > T_{f \max} + T_C + T_{\max}$, where $T_{\max} = (x_d + a_{\max} T_C + \Delta_{\max}) / (a_{\max} - d_{\max})$, the queue length is strictly positive.

Proof. Let us consider some $m \geq 1$ and the value of W at the corresponding time instant θ_m . Similarly as previously, we may distinguish two cases: the first occurs when $W(\theta_m) < a_{\max}$, and the other when $W(\theta_m) \geq a_{\max}$.

Case 1: First, let us investigate the situation when $W(\theta_m) < a_{\max}$

$$W(\theta_m) = K \left[x_d - x(\theta_m) - \sum_{j=1}^n \int_{\theta_m - \overline{R T T_j}}^{\theta_m} b_j(\tau) d\tau \right] < a_{\max}. \quad (54)$$

After performing algebraic manipulations, we get

$$\begin{aligned} x(\theta_m) &> x_d - \frac{a_{\max}}{K} - \sum_{j=1}^n \int_{\theta_m - R T T_j}^{\theta_m} b_j(\tau) d\tau + \sum_{j=1}^n \int_{\theta_m - R T T_j}^{\theta_m - \overline{R T T_j}} b_j(\tau) d\tau \\ &= x_d - \frac{a_{\max}}{K} - \sum_{j=1}^n \int_{\theta_m - R T T_j}^{\theta_m} b_j(\tau) d\tau + \Delta(\theta_m). \end{aligned} \quad (55)$$

The maximum rate established by the controller is equal to $a_{\max}$. At the same time, $\Delta(t) \geq -\Delta_{\min}$. So, using the theorem assumption (53), we may write

$$x(\theta_m) > x_d - \frac{a_{\max}}{K} - \frac{a_{\max}}{n} \sum_{j=1}^n R T T_j - \Delta_{\min} \geq a_{\max} T_C. \quad (56)$$

The minimum rate, which can be assigned for each source, is zero and the maximum available bandwidth equals $d_{\max}$. Then, applying (56) to formula (13) for the queue length at time instant $t \in [\theta_m, \theta_{m+1})$, we get $x(t) > a_{\max} T_C - d_{\max} \delta > 0$, which completes the first part of the proof.

Case 2: Now, let us investigate the circumstances when $W(\theta_m) \geq a_{\max}$. First, we find the last moment $t^* < \theta_m$ when signal W was smaller than $a_{\max}$. Theorem 3 implies that the queue length never exceeds the value of $x_d + a_{\max}T_C + \Delta_{\max}$ despite possible lack of precision in RTT 's estimation performed by the controller. On the other hand, packet depletion rate is limited by $d_{\max}$. Thus, the maximum period of time $T_{\max}$, during which the controller may continuously set the rate $a_{\max}$ for the sources, is estimated by the following equation

$$T_{\max} = (x_d + a_{\max}T_C + \Delta_{\max}) / (a_{\max} - d_{\max}). \quad (57)$$

Therefore, instant t^* exists. Since t^* is the last instant, when signal W was smaller than $a_{\max}$ and the actual rate assignment could be delayed by not more than T_C , $t^* \geq \theta_m - (T_{\max} + T_C)$. From the theorem assumption it also comes that $t^* > T_{f\max} + T_{\max} + T_C - (T_{\max} + T_C) = T_{f\max}$. The value of $W(t^*)$ satisfies the inequality given below

$$W(t^*) = K \left[x_d - x(t^*) - \sum_{j=1}^n \int_{t^* - \overline{RTT}_j}^{t^*} b_j(\tau) d\tau \right] < a_{\max}. \quad (58)$$

Following similar reasoning as presented in (55) and (56), we arrive at

$$x(t^*) > x_d - \frac{a_{\max}}{K} - \sum_{j=1}^n \int_{t^* - \overline{RTT}_j}^{t^*} b_j(\tau) d\tau > 0. \quad (59)$$

Applying (59) to formula (19) for the queue length at some time instant $t \in [\theta_m, \theta_{m+1})$, we get

$$x(t) > x_d - \frac{a_{\max}}{K} - \sum_{j=1}^n \int_{t^* - \overline{RTT}_j}^{t^*} b_j(\tau) d\tau + \sum_{j=1}^n \int_{t^* - \overline{RTT}_j}^{t - \overline{RTT}_j} b_j(\tau) d\tau - \int_{t^*}^t h(\tau) d\tau. \quad (60)$$

Similarly as in (55), we may use the concept of the function $\Delta(t)$ and rewrite inequality (60) in the form given below

$$x(t) > x_d - \frac{a_{\max}}{K} - \sum_{j=1}^n \int_{t^* - \overline{RTT}_j}^{t^*} b_j(\tau) d\tau + \Delta(t^*) + \sum_{j=1}^n \int_{t^* - \overline{RTT}_j}^{t - \overline{RTT}_j} b_j(\tau) d\tau - \int_{t^*}^t h(\tau) d\tau. \quad (61)$$

Performing analogical steps to those presented in the proof of Theorem 2, namely (35)–(39), we arrive at

$$x(t) > x_d - \frac{a_{\max}}{K} - a_{\max} \sum_{j=1}^n \frac{1}{n} \overline{RTT}_j + a_{\max} [t - (t^* + T_C)] + \Delta(t^*) - \int_{t^*}^t h(\tau) d\tau. \quad (62)$$

The function $\Delta(t)$ is lower-bounded by $-\Delta_{\min}$, while the maximum bandwidth to be utilised is limited by $d_{\max}$. Then, using assumption (53), we obtain

$$x(t) > \Delta_{\min} + a_{\max} (t - t^*) - \Delta_{\min} - d_{\max} (t - t^*) > 0. \quad (63)$$

This concludes the proof. $\square$

3.3. Modified controller

In this subsection we propose a modified control strategy. The individual transfer speed is still assigned for each source in accordance with (7) and the aggregate rate is limited as in (8). However, this time the function $W(t)$ is defined as

$$W(t) = K \left[x_d - x(t) - \sum_{j=1}^n \int_{t-RTT_j}^t b_j(\tau) d\tau + \lambda \cdot h(t) \sum_{j=1}^n \frac{1}{n} RTT_j \right]. \quad (64)$$

The proposed algorithm combines the Smith principle and feed-forward compensation with the proportional controller with saturation. The integral in (64) represents the Smith prediction, while term $\lambda h(t) \sum_{j=1}^n \frac{1}{n} RTT_j$ is responsible for the bandwidth compensation. The influence of the compensation on the system dynamics is tuned by a nonnegative real constant λ . When $\lambda = 0$ no compensation is applied, and as λ rises, the significance of the feed-forward term increases.

Further in this section we present two theorems, defining the properties of the proposed control scheme.

Theorem 5. If sources transmit data according to the conditions formulated by (3)–(9) with function $W(t)$ determined by (64), then the queue length at the bottleneck node does not exceed the value given below

$$\forall_{t \geq 0} x(t) \leq q_{\max} = x_d + \lambda d_{\max} \sum_{j=1}^n \frac{1}{n} RTT_j + a_{\max} T_C \quad (65)$$

where $q_{\max}$ denotes the maximum queue length.

Proof. First packets arrive at the switch at the $T_{f \min}$ time instant. Similarly as in the previous sections, we denote the queue length at the time θ_m by $x_m = x(\theta_m)$. For $m = 1$ and $\theta_1 = T_{f \min}$ we can write $x_1 = (\theta_1) = 0 < q_{\max}$. Thus, the proposition is valid for any moment of time $t \leq T_{f \min}$.

Let us consider some $m \geq 1$ and the value of the signal $W(t)$ at the time instant $t = \theta_m$. Two cases need to be considered: the situation when $W(\theta_m) \geq 0$, and the other when $W(\theta_m) < 0$.

Case 1: Investigating the circumstances when $W(\theta_m) \geq 0$, we get

$$W(\theta_m) = K \left[x_d - x(\theta_m) - \sum_{j=1}^n \int_{\theta_m - RTT_j}^{\theta_m} b_j(\tau) d\tau + \lambda h(\theta_m) \sum_{j=1}^n \frac{1}{n} RTT_j \right] \geq 0. \quad (66)$$

The utilised bandwidth is limited from above by $d_{\max}$. Therefore, after the term rearrangement, we obtain

$$x(\theta_m) \leq x_d - \sum_{j=1}^n \int_{\theta_m - RTT_j}^{\theta_m} b_j(\tau) d\tau + \lambda h(\theta_m) \sum_{j=1}^n \frac{1}{n} RTT_j \leq x_d + \lambda d_{\max} \sum_{j=1}^n \frac{1}{n} RTT_j. \quad (67)$$

Applying (67) to relation (13) for the queue length at the time $t = \theta_m + \delta$, we conclude

$$x(t) \leq x_d + \lambda d_{\max} \sum_{j=1}^n \frac{1}{n} RTT_j + \sum_{j=1}^n \int_{\theta_m - RTT_j}^{\theta_m + \delta - RTT_j} b_j(\tau) d\tau - \int_{\theta_m}^{\theta_m + \delta} h(\tau) d\tau. \quad (68)$$

In analogy to the steps illustrated in (14)–(16), we get

$$x(t) \leq x_d + \lambda d_{\max} \sum_{j=1}^n \frac{1}{n} RTT_j + \delta a_{\max} - 0 \leq x_d + \lambda d_{\max} \sum_{j=1}^n \frac{1}{n} RTT_j + a_{\max} T_C \quad (69)$$

which ends the first part of the proof.

Case 2: Now, let us examine the situation when $W(\theta_m) < 0$. First, we find the last moment $t < \theta_m$, when the signal $W(t)$ was greater than zero. Indeed, such instant exists since $W(t \leq T_{f \min}) = K(x_d - 0 - 0 + 0) = Kx_d > 0$. Notice also that the first moment when function W may attain a negative value is greater than $T_{f \min}$. The value of $W(t^*)$ satisfies the following inequality

$$W(t^*) = K \left[x_d - x(t^*) - \sum_{j=1}^n \int_{t^* - RTT_j}^{t^*} b_j(\tau) d\tau + \lambda h(t^*) \sum_{j=1}^n \frac{1}{n} RTT_j \right] > 0. \quad (70)$$

Consequently, $x(t^*)$ can be estimated by the relation given below

$$x(t^*) < x_d + \lambda h(t^*) \sum_{j=1}^n \frac{1}{n} RTT_j - \sum_{j=1}^n \int_{t^* - RTT_j}^{t^*} b_j(\tau) d\tau. \quad (71)$$

The expression for the queue length at a time instant $\in [\theta_m, \theta_{m+1})$ is identical to (19). Therefore,

$$\begin{aligned} x(t) < & x_d + \lambda h(t^*) \sum_{j=1}^n \frac{1}{n} RTT_j - \sum_{j=1}^n \int_{t^* - RTT_j}^{t^*} b_j(\tau) d\tau \\ & + \sum_{j=1}^n \int_{t^* - RTT_j}^{t - RTT_j} b_j(\tau) d\tau - \int_{t^*}^t h(\tau) d\tau. \end{aligned} \quad (72)$$

Reasoning similar to the one presented in (20)–(22) brings the conclusion

$$x(t) < x_d + \lambda d_{\max} \sum_{j=1}^n \frac{1}{n} RTT_j + a_{\max} T_C \quad (73)$$

This completes the proof. $\square$

Theorem 6. If sources transmit data according to the conditions formulated by (3)–(9) together with function $W(t)$ defined by (64), the maximum rate $a_{\max} > d_{\max}$ and the demand queue length satisfies the following inequality

$$x_d > a_{\max} \left(\sum_{j=1}^n \frac{1}{n} RTT_j + \frac{1}{K} + T_C \right) \quad (74)$$

then for any $t > T_{f \max} + T_C + T_{\max}$, where $T_{\max} = q_{\max}/(a_{\max} - d_{\max})$ the queue length is strictly positive.

Proof. The minimum utilised bandwidth $h(t)$ is zero. With this remark, the analysed proposition is valid as a direct consequence of Theorem 2. $\square$

The steady-state queue length x_{ss} , i.e. the queue length, when the available bandwidth $d_{ss} > 0$ is constant, can be expressed as follows

$$x_{ss} = x_d - d_{ss}/K - (1 - \lambda) d_{ss} \sum_{j=1}^n \frac{1}{n} RTT_j. \quad (75)$$

When

$$\lambda = 1 + \left(K \sum_{j=1}^n RTT_j/n \right)^{-1} \quad (76)$$

the steady-state queue length $x_{ss} = x_d$, which implies complete insensitivity of x_{ss} to the available bandwidth at the bottleneck link. As the dependency of the queue length on d_{ss} diminishes, the delay jitter for the transferred data is reduced, which helps achieve better QoS in the network and enables provision of a broad class of services connected with video and audio streaming.

4. SIMULATION RESULTS

In order to verify the control strategies proposed in this paper, simulation tests were performed using Matlab–Simulink. First, the model of wide area network with irregular period of feedback information availability was constructed according to the description provided in Section 2. Three connections ($n = 3$) participate in the flow regulation process. They are characterised by the delay parameters summarised in Table.

Table. Delay parameters.

Source	Delay [ms]		
j	T_{fj}	T_{bj}	RTT_j
1	5	15	20
2	10	20	30
3	30	40	70

To check the possibility of application of the proposed control schemes to a real telecommunication system, we adjusted the feedback parameters according to the guidelines of the ATM standard. Consequently, in our model each source sends control units every $M = 32$ equal-size data pieces, but not less frequently than every $T_C = 100$ ms. The maximum available bandwidth $d_{\max}$ was set as 9100 packets/s, which approximately corresponds to 3.7 Mb/s connection, and the upper bound of the overall source rate $a_{\max}$ was adjusted to 10100 packets/s $\approx 1.11d_{\max}$. The bandwidth actually available for the controlled connections $d(t)$ at the bottleneck node is illustrated in Figure 2. We can see from the graph that function d experiences

sudden changes of large amplitude, which reflects the most rigorous networking conditions.

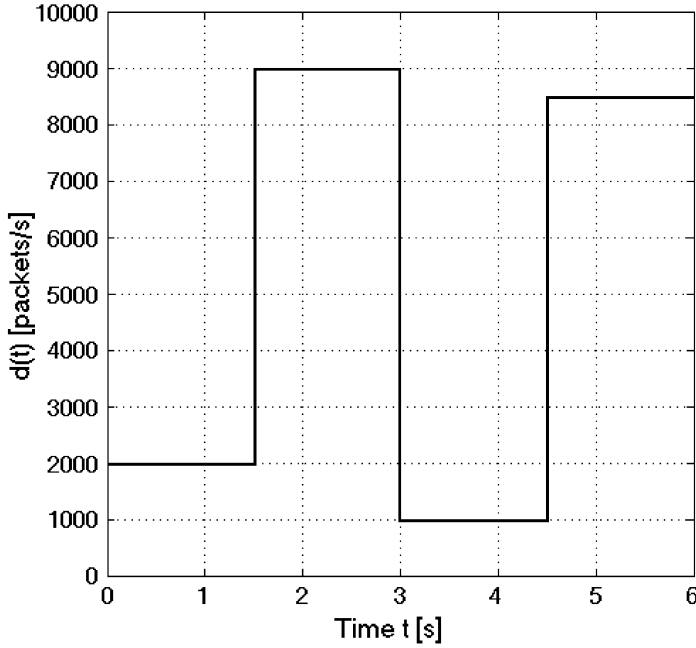


Fig. 2. Available bandwidth.

The system becomes more dynamic with the increase of the controller gain K . Although any value $K > 0$ satisfies Theorems 1–6, we adjusted K to 100 s^{-1} , as it proved to be sufficient for a majority of typical communication scenarios.

Three simulations were run. In the first one it was assumed that the controller has exact knowledge of the real delays existing in the network. Therefore, the Smith predictor parameters matched the true propagation delays: $RTT_1 = 20 \text{ ms}$, $RTT_2 = 30 \text{ ms}$ and $RTT_3 = 70 \text{ ms}$. In order to fulfil the requirements imposed by Theorem 2, the demand queue length $x_d = 1520 > 1515$ packets was chosen. Queue length $x(t)$ resulting from applying the principal regulation scheme described by relations (7)–(9) is shown in Figure 3. As we can see, no data arrives at the bottleneck node before $T_{f \min} = 5 \text{ ms}$. Moreover, the queue length never exceeds the value of $x_d + a_{\max}T_C = 2530$ packets and does not drop to 0. These two properties imply no buffer overflow and full bottleneck link utilisation.

In the second simulation example it is assumed that the controller determines the delay times for the regulated flows with decreased accuracy. Thus, the real RTT s remain unchanged when compared with the first test, but the Smith predictor parameters are modified as follows: $\overline{RTT}_1 = 22 \text{ ms}$, $\overline{RTT}_2 = 34 \text{ ms}$ and $\overline{RTT}_3 = 67 \text{ ms}$. In order to provide full bottleneck node bandwidth usage the demand queue length is adjusted to the value of 1540 packets. The resulting $x(t)$ function is illustrated in Figure 4. Notice that the algorithm preserves its favourable features, i. e. the queue

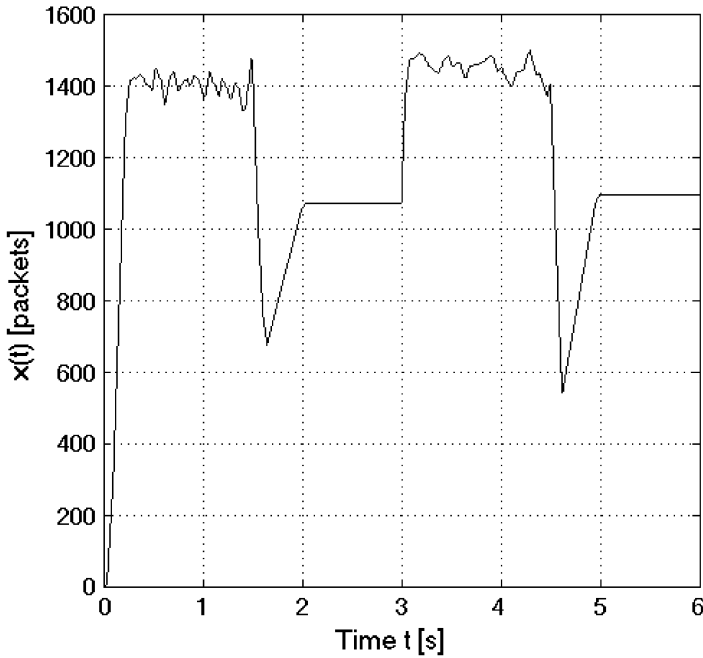


Fig. 3. Queue length – principal control scheme.

length does not grow beyond the level of 2560 packets and $x(t)$ is strictly positive. This incurs no data loss in the examined system and full resource exploitation. However, both examples reveal possible drawback of the proposed strategy with respect to handling the multimedia traffic. Namely, the queue length in steady states is highly dependent on the currently available bandwidth.

In the third simulation scenario it is assumed again that the controller possesses precise knowledge of the propagation delays for the regulated connections. Therefore the Smith predictor parameters become: $RTT_1 = 20$ ms, $RTT_2 = 30$ ms and $RTT_3 = 70$ ms. This time however, rate estimation function W is replaced by that defined in (64). The feed-forward tuning coefficient is selected according to (76), i. e.

$\lambda = 1 + \left(K \sum_{j=1}^n RTT_j / n \right)^{-1} = 1.25$. The queue length evolution in the bottleneck node buffer is presented in Figure 5. The upper curve (a) shows that the packet queue approaches the level of $x_d = 1520$ packets in the steady states. The lower graph (b) illustrates the effect of decreasing the demand queue length in order to limit memory volume and, what is of utmost importance in modern telecommunication systems, to obtain better quality of service (smaller average transfer delay and delay variation). It can be seen from the figure that although the demand queue length is significantly reduced ($x_d = 520$), the packet queue drops to zero only for a very short period of time, bringing negligible deterioration of the link bandwidth usage. Therefore, the proposed strategy allows for a propitious trade-off between QoS and utilisation of the network resources.

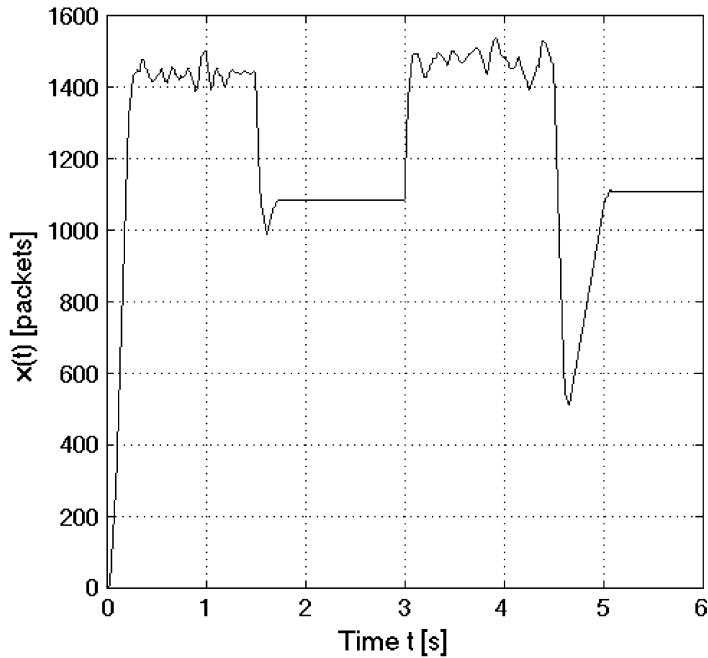


Fig. 4. Queue length – robustness analysis.

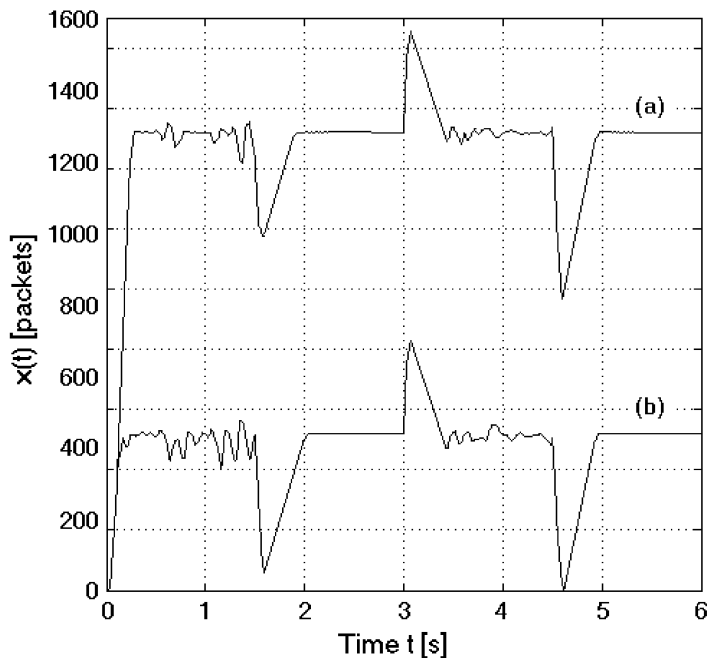


Fig. 5. Queue length – feed-forward compensation.

5. CONCLUSIONS

In this paper new nonlinear algorithms for flow control in connection-oriented telecommunication networks were presented. The proposed algorithms ensure no data loss and full resource utilisation in the multi-source system, where the feedback information necessary for the transfer speed adjustment is accessible at irregular time instants. First, the proper operation of the Smith principle based controller was demonstrated despite possible mismatch occurring between the real propagation delays and those estimated by the algorithm. Following the robustness analysis, a modified flow regulation scheme was proposed. The controller with the extra feed-forward term incorporated not only maintains the favourable features of the principal strategy, but also helps reduce the variation of the steady-state queue length caused by the available bandwidth, thus decreasing the data transport delay jitter. This allows improving of the QoS in the network and facilitates handling of the multimedia traffic. Since the described control schemes do not require constant (as in continuous systems), nor time-synchronised (as in discrete systems with constant sampling period) exchange of the feedback information, they are more scalable and easier in practical implementation than similar solutions proposed earlier in literature.

(Received May 24, 2007.)

REFERENCES

- [1] A. Bartoszewicz and T. Molik: ABR traffic control over multi-source single-bottleneck ATM networks. *J. Appl. Math. Comput. Sci.* *21* (2004), 43–51.
- [2] A. Bartoszewicz: Nonlinear flow control strategies for connection-oriented communication networks. *IEE Control Theory Appl.* *153* (2006), 21–28.
- [3] S. Chong, R. Nagarajan, and Y.T. Wang: First-order rate-based flow control with dynamic queue threshold for high-speed wide-area ATM networks. *Comput. Netw. ISDN Syst.* *29* (1998), 2201–2212.
- [4] F. Gómez-Stern, J.M. Fornés, and F.R. Rubio: Dead-time compensation for ABR traffic control over ATM networks. *Control Engrg. Pract.* *10* (2002), 481–491.
- [5] O.C. Imer, S. Compans, T. Basar, and R. Srikant: Available bit rate congestion control in ATM networks. *IEEE Control Syst. Mag.* *21* (2001), 38–56.
- [6] R. Izmailov: Adaptive feedback control algorithms for large data transfers in high-speed networks. *IEEE Trans. Automat. Control* *40* (1995), 1469–1471.
- [7] S. Jagannathan and J. Talluri: Predictive congestion control of ATM networks: multiple sources/single buffer scenario. *Automatica* *38* (2002), 815–820.
- [8] R. Jain: Congestion control and traffic management in ATM networks: recent advances and a survey. *Comput. Netw. ISDN Syst.* *28* (1996), 1723–1738.
- [9] L. A. Kulkarni and S. Li: Performance analysis of a rate-based feedback control scheme. *IEEE/ACM Trans. Netw.* *6* (1998), 797–810.
- [10] K.P. Laberteaux, C.E. Rohrs and P.J. Antsaklis: A practical controller for explicit rate congestion control. *IEEE Trans. Automat. Control* *47* (2002), 960–978.
- [11] I. Lengliz and F. Kamoun: A rate-based flow control method for ABR service in ATM networks. *Comp. Netw.* *34* (2000), 129–138.
- [12] S. Mascolo: Congestion control in high-speed communication networks using the Smith principle. *Automatica* *35* (1999), 1921–1935.
- [13] S. Mascolo: Smith’s principle for congestion control in high-speed data networks. *IEEE Trans. Automat. Control* *45* (2000), 358–364.

- [14] S. Mascolo: Modeling the Internet congestion control using a Smith controller with input shaping. *Control Engrg. Pract.* *14* (2006), 425–435.
- [15] F. D. Priscoli and A. Pietrabissa: Design of bandwidth-on-demand (BoD) protocol for satellite networks modelled as time-delay systems. *Automatica* *40* (2004), 729–741.
- [16] P. F. Quet, B. Ataslar, A. Iftar, H. Özbay, S. Kalyanaraman, and T. Kang: Rate-based flow controllers for communication networks in the presence of uncertain time-varying multiple time-delays. *Automatica* *38* (2002), 917–928.
- [17] M. L. Sichitiu and P. H. Bauer: Asymptotic stability of congestion control systems with multiple sources. *IEEE Trans. Automat. Control* *51* (2006), 292–298.
- [18] O. J. Smith: *Feedback Control Systems*. McGraw-Hill, New York 1958.

*Przemysław Ignaciuk and Andrzej Bartoszewicz, Institute of Automatic Control, Technical University of Łódź, 18/22 Stefanowskiego St., 90-924 Łódź. Poland.
e-mails: pignaciuk@poczta.onet.pl, andrzej.bartoszewicz@p.lodz.pl*