



Reliability Enhancement of Redundancy Management in AFDX Networks

Meng Li, Guchuan Zhu, Yvon Savaria, Michaël Lauer

► To cite this version:

Meng Li, Guchuan Zhu, Yvon Savaria, Michaël Lauer. Reliability Enhancement of Redundancy Management in AFDX Networks. IEEE Transactions on Industrial Informatics, 2017, 13 (5), pp.2118-2129. 10.1109/TII.2017.2732345 . hal-01585141

HAL Id: hal-01585141

<https://laas.hal.science/hal-01585141>

Submitted on 11 Sep 2017

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

Reliability Enhancement of Redundancy Management in AFDX Networks

Meng Li, Guchuan Zhu, *Senior Member, IEEE*, Yvon Savaria, *Fellow, IEEE*, and Michaël Lauer

Abstract—AFDX is a safety critical network in which a redundancy management mechanism is employed to enhance the reliability of the network. However, as stated in the ARINC664-P7 standard, there still exists a potential problem, which may fail redundant transmissions due to sequence inversion in the redundant channels. In this paper, we explore this phenomenon and provide its mathematical analysis. It is revealed that the variable jitter and the transmission latency difference between two successive frames are the two main sources of sequence inversion. Thus, two methods are proposed and investigated to mitigate the effects of jitter pessimism, which can eliminate the potential risk. A case study is carried out and the obtained results confirm the validity and applicability of the developed approaches.

Index Terms—Reliability Enhancement, AFDX, Virtual Link, Fault Tolerance.

NOMENCLATURE

$(\cdot)^+$	$\max(\cdot, 0)$.
$L_{\max}$	The maximum frame length.
T_i	The period of VL_i .
σ_i	The maximum frame size plus 20 bytes overhead.
τ_i	A variable delay of VL_i and $0 \leq \tau_i \leq T_i$.
O_i	Time offset of the first frame of VL_i .
$u_{T_i, \tau_i, \sigma_i}$	The staircase arrival curve for VL_i .
J_{e2e}	The end-to-end jitter upper bound.
N_j	The maximum number of VL_j within T_i .
$M_i^{(k)}$	Residual bytes in the worst case when $(k+1)$ th frame arrives taking VL_i as the reference.
$D_{i,j}$	The minimum release time difference between adjacent frames of VL_i and VL_j , where the frame of VL_j is ahead.
$D_{i,j}(q)$	The release time difference between the q th frame of VL_j and the reference frame of VL_i .

I. INTRODUCTION

RELIABILITY is one of the main concerns for safety critical systems (See, e.g., [1]–[4]). A typical example of such systems is avionics communication network for which failures may be catastrophic. Therefore, guaranteeing a reliable communication among avionics systems at every flight phase is critical for aircrafts. To ensure that stringent reliability

requirements are met, certain standards, e.g., ARINC 429, have been developed and successfully deployed since the late 1970s [5]. However, as the amount of electronic components in an aircraft continues to increase, legacy avionics communication protocols are at their limit in terms of performance and design complexity. Among the available technologies for handling the new challenges in avionics systems design, we can find an Ethernet-based technology, namely the Avionics Full Duplex Switched Ethernet (AFDX) [6], which features high speed, low cost, high flexibility, and reduced weight because of less wiring.

Built on the basis of Ethernet technology, the AFDX not only offers a high available bandwidth and a high communication speed, but also provides deterministic performance, which is the most prominent challenge to using such a technology in avionics. In AFDX networks, determinism is enforced mainly through the concept of Virtual Link (VL), which defines a logical unidirectional connection and a bounded data transmission bandwidth. Besides, the allocated bandwidth is reserved by VL's maximum frame size (MFS) and the so-called Bandwidth Allocation Gap (BAG), which defines the minimum time interval between the first bits of two successive frames of a VL at the ingress of the networks. Furthermore, the AFDX network is composed of two independent and redundant networks, which provides the required reliability for ensuring its determinism. Consequently, the unavoidable faults on single paths can be tolerated by the redundancy management mechanism.

Nevertheless, although the AFDX was originally developed for safety critical avionic applications, it has not yet been used in critical systems that require the highest level of reliability, e.g., flight control systems [7]. Much efforts are still required to prove that AFDX networks can achieve the highest reliability requirement of critical functions, i.e., a failure probability of 10^{-9} per flight hour [8]. Specifically, as pointed out in ARINC664-P7 (see Section 3.2.6 in [9]), the redundant transmission mechanism fails if the following two events occur simultaneously: (1) a frame is lost during transmission on one of the redundant networks; (2) the subsequent frame on the network with frame loss arrives earlier to the destination End Systems (ES) than the copy of the lost frame sent through the other network, which is called a sequence inversion in the redundant channels. Obviously, this is a potential risk that could compromise the network reliability. In real avionics communication networks, frame loss, even if observed with a very small probability, is inevitable. Therefore in order to guarantee the reliability of the redundancy management

M. Li, G. Zhu, and Y. Savaria are with the Department of Electrical Engineering, Polytechnique Montréal, Montréal, QC, Canada H3T 1J4 (e-mail: meng.li@polymtl.ca, guchuan.zhu@polymtl.ca; yvon.savaria@polymtl.ca).

M. Lauer is with LAAS-CNRS, Université de Toulouse, 7, avenue du Colonel Roche BP 54200, 31031 Toulouse Cedex 4, France (e-mail: michael.lauer@laas.fr).

mechanism, one must prevent the second condition from occurring. This is a real challenge to system designers, due to the lack of an analytical framework for this problem.

The motivation of this paper is to provide a mathematical analysis of RM failures in the AFDX protocol. It is revealed that the sequence inversion phenomenon, which can result in the invalidity of the redundancy management mechanism, is due to the variable jitter and the transmission latency difference between two successive frames. To tackle this problem, two methods that can contribute to eliminate the sequence inversion problem are proposed. One of these methods is based on local synchronization (LS) [10], [11] and the other exploits the notion of transmission latency difference minimization (TLDM) proposed in this work. This allows enhancing the reliability of RM. We show that these two approaches help mitigating the delay difference between two redundant networks in the worst case. A case study is carried out and the obtained results confirm the validity and the applicability of the developed approaches.

To the best of our knowledge, this is the first work that presents a formal mathematical analysis on the sequence inversion problem. Specifically, the main contributions of this paper are:

- identifying the sources of sequence inversion and providing a mathematical analysis regarding potential failures in RM;
- introducing two approaches that can eliminate potential failures due to frame sequence inversion of the redundant networks.

The aim of the present work focuses mainly on enhancing the determinism and the reliability of AFDX networks to take a step forward towards the application of this promising technology to highly safety-critical avionics systems, such as flight control systems.

The remaining of the paper is organized as follows. Section II introduces the context of AFDX networks. Section III describes some potential failures in redundant AFDX networks and provides a corresponding mathematical analysis. In Section IV two approaches are developed to enhance the reliability of RM. In Section V, a case study is carried out to validate the developed approaches and to evaluate the obtained performance. Finally, some concluding remarks and directions for future research are provided in Section VI.

II. THE CONTEXT OF AFDX NETWORKS AND RELATED WORK

A. Basis of AFDX Networks

An AFDX network is typically composed of three types of elements: ESs, switches and physical links. Each ES is connected to the switches via redundant physical links, denoted by Network A (-A suffix to switches) and Network B (-B suffix to switches) as shown in Fig. 1. Full duplex physical links are adopted to eliminate transmission collisions, which help to ensure deterministic timing performance. In addition, a star topology is applied in switch connections, which makes the network scalable. Usually, it is supposed that the switches have the capability of handling parallel processing. Hence, there

is no interference between the packets forwarded to different outputs.

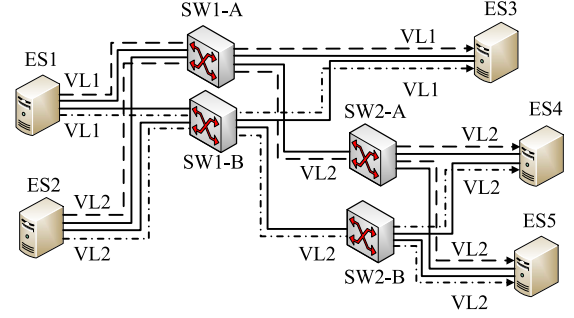


Fig. 1: An example of AFDX network architecture.

The determinism of AFDX networks is enforced mainly through the concept of VL. Specifically, in AFDX networks, only one ES can be the source of a VL and the routing of VLs is statically defined off-line. Furthermore, a VL can be composed of up to four Sub-VLs to improve the bandwidth utilization efficiency [12].

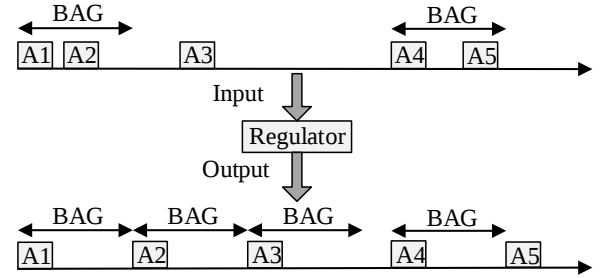


Fig. 2: VL flow regulation.

As shown in Fig. 2, input frames, either periodic or aperiodic, are regulated by the BAG, through which the instantaneous frame rate of a VL is limited. Therefore, the maximum bandwidth allocated to a VL is determined by its MFS and BAG [9]. According to the ARINC664 standard, the MFS should be in the range of 64 to 1518 bytes, including a header of 47 bytes. It also needs to take into account an overhead of 20 bytes (Interframe Gap+Preamble+Start Frame Delimiter) during frame transmission. The BAG should be a power of 2 multiplied by 1 ms within the set {1, 2, 4, 8, 16, 32, 64, 128}(ms).

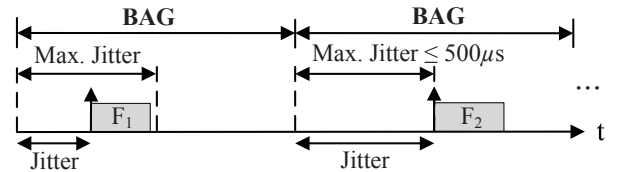


Fig. 3: The jitter of a VL cannot exceed 500 μ s in a source ES [9].

Scheduling in an ES or a switch is performed on a per VL basis, which may introduce jitters due to the congestion of VLs at the outputs. According to the AFDX standard, the jitter of

a VL cannot exceed $500 \mu\text{s}$ in a source ES as shown in Fig. 3. Furthermore, traffic policing is applied in switches to protect the network from babbling-idiot failures [13]. In addition, the technology latency in a switch, which is the time required to process frames that should be less than $100 \mu\text{s}$ irrespective of traffic load [9]. Usually, the technology latency is assumed to be upper bounded during analysis. The characteristics of VLs and the traffic shaping and policing mechanisms are essential for guaranteeing that the end-to-end delay of each frame can be upper bounded.

B. Redundancy Management

As shown in Fig. 1, in an AFDX network, the frames in a VL are transmitted through two redundant and independent paths to achieve a high level of communication reliability. This redundancy mechanism assures a reliable communication against the loss of one complete network (Network A or Network B). For each transmitted frame, a sequence number (SN) is added to enable receivers to reconstruct an ordered stream of frames without duplication. In general, SN ranges from 0 to 255, and it is initially set to 0 and increased by 1 for each consecutive transmission of the same VL. The SN wraps around to 1 following the value of 255. Denoting by i the value of a SN, then the wrap-around operation for SNs can be computed as:

$$i \oplus 1 = (i \bmod 255) + 1. \quad (1)$$

Furthermore, two redundant frames with identical SNs must be received in an interval less than a predefined SkewMax. Otherwise, the latter reception is considered as a new frame. Hence, SkewMax is the upper bound of transmission delay difference for the redundant frames with identical SN.

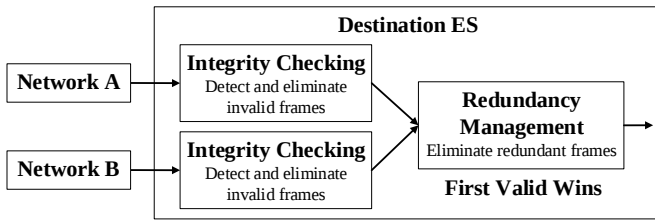


Fig. 4: Redundancy Management in destination ES [9].

As shown in Fig. 4, an integrity checking is independently performed for each network on a per VL basis at the destination ES. At this stage, only well-formed frames, i.e., frames that contain correct Cyclic Redundancy Check (CRC) field and a proper SN, will be forwarded to RM (see Section 3.2.6 in [9] for details). The RM is performed after integrity checking, and hence only the valid frames are processed at this stage. The basic rule used in AFDX redundancy management is the “First Valid Wins” (FVW) policy. In RM, a previous sequence number (PSN) is stored for comparison and the PSN is updated after each valid reception. Normally, if a frame succeeded in integrity checking, its SN will be compared with the stored PSN. If the received SN is increasing compared with the current PSN based on the wrap-around operation, the frame will be forwarded and the PSN is updated accordingly. If the

PSN is equal to the SN, the frame with this SN is regarded as a redundant reception and will be discarded. An exception is when the SkewMax has been exceeded. In this case, the RM will accept any valid frame regardless of its SN. Thus, the SkewMax value for each VL should be carefully assigned.

Although the redundant design in AFDX networks enhances its fault tolerance, there still exist potential situations where the redundancy management mechanism may fail to manage redundant frames, which results in frame losses.

C. Related Work

For safety critical systems, it is essential to guarantee a reliable real-time communication. Thus, the computation of tight and deterministic delay upper bounds is one of the major issues for both communication network design and network certification [14]. Much effort has been dedicated to estimate the upper bounds for data transmission delays in order to guarantee timing behavior of the network based on formal analysis.

Network calculus is a mathematical tool that has been widely applied to performance analysis of communication networks by considering worst case scenarios. It was first introduced by Cruz based on min-plus algebra [15], [16], and then detailed by Le Boudec and Thiran [17]. A principle of “Pay Bursts Only Once” is proposed in [17] based on the property of the convolution of service curves, which contributes to tightening delay bound estimations. Significant improvements in delay upper bound estimation have been achieved by the introduction of a grouping technique, which leads to an approximate gain of up to 40% for a realistic AFDX configuration by considering the effect of serialization stream [18]. A stochastic extension of network calculus has also attracted much interest, and the application of probabilistic bounds in the analysis of AFDX networks can be found in [19]. In the framework of network calculus, traffic inputs are modeled by arrival curves, among which the most popular one is the fluid model. However, staircase models are more accurate for describing the property of packetization, although such models are known to be complex. In [20], a combination of fluid modelling and staircase modelling is introduced to make a trade-off between tighter bounds and computation complexity. It has been reported in [14] that the staircase model can lead to, on average, a gain of 18% for randomly generated configurations. As network calculus is able to deal with both periodic and aperiodic flows, partial synchronization of the source flows can be taken into account during analysis. In [21] periodic flows with known offsets in source ESs are considered to eliminate some pessimistic scenarios, which lead to a reduction of delay upper bounds. In [22], the event-stream model formulated with a staircase model is applied to obtain upper bounds of traffic. Another solution to determine the delay upper bounds is the trajectory approach, which considers the worst-case scenario experienced by a frame along its path [23]–[29]. This technique has been applied in the analysis of AFDX networks in [24], [27] based on the FIFO policy. The grouping technique is also taken into account to mitigate the pessimism. In [26], the source of pessimism in

the computation of upper bounds with the trajectory approach is characterized. Although it has been reported in [30] that the trajectory approach introduces optimism in some corner cases, the problems have been identified and corrected [31]. However, a formal proof for the fix proposed in [31] is still required. Recently, another approach, namely the forward end-to-end delays analysis (FA), has been proposed to obtain the delay upper bounds [32]. Similar to the trajectory approach, this method focuses on one frame and analyzes iteratively the components (ES and Switches) through which the frame passed. In [33], an improvement of the FA approach has been achieved by considering the serialization of frames sharing a common link, i.e., the grouping technique.

Besides the research on end-to-end delay analysis of AFDX networks, some other work focuses on traffic scheduling and redundancy management. In [34], a traffic phase shifting technique is proposed to improve the bandwidth utilization by assigning offsets to periodic traffic flows. However, the improvement is achieved at the expense of increasing end-to-end delay due to the fact that the frames are buffered at switches to wait for their time slots for transmission. A deduplication-aware Deficit Round Robin scheduling scheme is proposed in [35] to offer flexible scheduling and implement fast deduplication. In [36], mixed-criticality traffic scheduling was investigated to enhance resource efficiency. In [13], the frame management in AFDX networks is analyzed and a modified design with a priority queue is proposed. Although it can offer a better data integrity and a higher QoS, the improvement is achieved at the expense of higher latencies. Moreover, it is not a generic solution and its applicability highly depends on the property of applications. In [37], three redundancy management algorithms (RMA1, RMA3, and RMA13) are analyzed using model checking. It is reported that RMA1 and RMA3 have difficulty to handle the redundant networks when one network fails. It seems that RMA13 is the best choice considering safety properties. However, compared with the FVW policy, RMA13 only accepts the frame from the same network as the last frame, which degrades the availability provided by the redundancy mechanism and can cause higher frame loss rates.

Based on the above mentioned works, the QoS of AFDX networks can be improved by using different approaches. However, to the best of our knowledge, there is little work on the rigorous analysis of the sequence inversion problem. The motivation of this paper is then to identify the sources of sequence inversion and to develop solutions for eliminating the potential failures.

III. TRANSMISSION FAILURES IN AFDX NETWORKS

A. Frame Loss Resulting from Sequence Inversion

Although AFDX networks provide a highly reliable communication via redundant networks, the RM may fail in some special cases. Such possible failure cases have been identified in the standard ARINC664-P7 (see, Section 3.2.6 in [9]), which may occur when a frame is lost on the faster network.

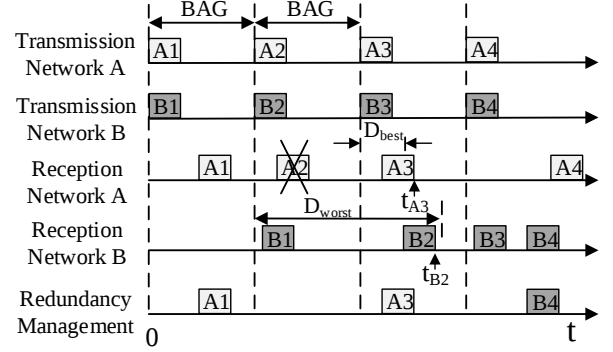


Fig. 5: Impact of a frame lost in a redundant AFDX network due to a transmission failure on the faster network (adapted from [9]).

For example, let us consider two redundant networks, Network A and Network B, that transmit their frames every BAG interval as shown in Fig. 5. Suppose that one frame on the faster network, e.g., A2 on Network A, is lost during transmission, e.g., due to bit errors corrupting the frame contents. To tolerate such failures, redundancy is employed in AFDX networks to increase the network reliability. However, if frame A3 arrives earlier than the frame B2 as shown in this example, a frame loss failure happens in spite of the redundant transmission. This results from the destination ES applying the FVW policy. Essentially, frame loss in the redundant AFDX network is due to frame sequence inversion of Network A and Network B at the destination ES.

B. Mathematical Analysis of the Frame Sequence Inversion

In this section, we provide a detailed mathematical analysis of the frame sequence inversion phenomenon. The analysis is based on three assumptions: (1) the redundant frames are fed to the 2 redundant networks simultaneously at the source ES; (2) Network A and Network B have identical topology and configurations, which include the same set of VLs; (3) the technological latency in both source ESs and switches is upper bounded. Note that these assumptions are used only for the purpose of simplifying the presentation and the relaxation of these assumptions will not introduce any technical difficulty.

Denote by D_{worst} the worst-case delay upper bound experienced by the frames with maximum size in a VL. Let the transmission latency be the transmission time over the physical links. Thus D_{best} , the minimum frame delay, can be taken as the sum of technology latencies and transmission latency, which is determined by the routing of the corresponding VL and the minimum frame size. The difference between D_{worst} and D_{best} is due to the variance of frame size and the jitter caused by the influence of other VLs that share the output ports in source ES or in switches. Based on the assumptions above, the VLs in both networks have the same parameters with respect to D_{worst} and D_{best} . For example, in the case shown in Fig. 5, the delay of A3 cannot be smaller than D_{best} and the delay of B2 cannot exceed D_{worst} . Note that data transmission is considered to be completed when the last bit of the frame is received. Then, the reception is accomplished at t_{A3} for A3 and t_{B2} for B2, respectively. Assume that the

first frame transmission starts at zero, then for the reception of A3 and B2 we have:

$$\begin{cases} t_{A3} \geq 2BAG + D_{\text{best}}, \\ t_{B2} \leq BAG + D_{\text{worst}}. \end{cases} \quad (2)$$

If A3 arrives earlier than B2, then we have $t_{A3} < t_{B2}$. Considering the constraints in (2), we can obtain

$$D_{\text{worst}} - D_{\text{best}} > BAG. \quad (3)$$

Denote by L_{max} and L_{min} the maximum and minimum frame sizes of the VL, respectively. Let C be the transmission rate of the physical links and n be the number of physical links the VL traverses. In that case, we have

$$D_{\text{worst}} - D_{\text{best}} = J_{e2e} + (L_{\text{max}} - L_{\text{min}}) / C \times n,$$

where J_{e2e} represents the end-to-end jitter upper bound induced by its burst and other VLs during data transmission. Note that the order of frames belonging to a VL on each path is maintained by the switches to guarantee no frame sequence inversion. Therefore, order inversion can only occur in destination ES for the frames belonging to different networks, i.e., Network A and Network B.

C. Condition for Avoiding Frame Sequence Inversion

In order to avoid the possible failure due to frame sequence inversion, the transmission delay difference between any two successive frames, that have different SN and come from different paths, should be restricted within a BAG. Note that the transmission delay difference is different from the previously mentioned parameter SkewMax. Given i a natural number, let $D_A(i)$ and $D_B(i)$ be the delay of the i th frame on Network A and Network B, respectively. Denote by $L(i)$ the size of the i th frame. Denote by $J_A(i)$ and $J_B(i)$ the jitters experienced by the frames with index i traversing Network A and Network B, respectively. Then we have:

$$D_A(i) - D_B(i+1) \leq J_A(i) - J_B(i+1) + \frac{n \times (L(i) - L(i+1))^+}{C}, \quad (4)$$

where $(\cdot)^+$ is defined by $\max(\cdot, 0)$. The constraint for $(D_B(i) - D_A(i+1))$ can be obtained similarly as for (4). As the introduced jitter has an upper bound of J_{e2e} and a lower bound of 0, both $(J_A(i) - J_B(i+1))$ and $(J_B(i) - J_A(i+1))$ are upper bounded by J_{e2e} . Note that since the redundant frames of a VL are released by the same source ES, the service latency induced by the source ES can be deducted from the jitter upper bound. In this context, the constant rate service model is applied for source ESs in the following analysis. Furthermore, denoting by $D_{TLD}(i)$ the transmission latency difference between two consecutive frames, then the general expression can be given as

$$D_{TLD}(i) = \frac{n \times (L(i) - L(i+1))^+}{C}. \quad (5)$$

Thus, the condition to avoid the possible failure is given by:

$$J_{e2e} + \max_i \{D_{TLD}(i)\} < BAG. \quad (6)$$

The first part on the left-hand side of this inequality represents the maximum jitter introduced by the VL frame with the maximum size and other VLs during transmission, and the second part denotes the maximum transmission latency difference between two successive frames. Therefore in order to meet the condition (6), it is required to mitigate the pessimism in jitter estimation and to reduce transmission latency difference between two successive frames.

In general, less pessimistic upper bounds can be obtained with more realistic models. Therefore, the staircase model, which is more accurate than the affine model, is employed to achieve better estimations. Relevant research can be found in [14], [20]. The experiment reported in [14] shows that the delay upper bounds can be improved up to 18% on average by using the staircase model instead of the affine arrival curve. However, the condition (6) shows that restricting the jitter upper bound within one BAG still cannot guarantee the elimination of sequence inversion. In the next section, two approaches that can contribute to eliminate the resulting potential failures are proposed.

IV. APPROACHES TO ELIMINATE THE POTENTIAL RISK OF FRAME SEQUENCE INVERSION

As shown in (6), to avoid the SN inversion, the sum of jitter upper bound and transmission latency difference has to be constrained within one BAG. This section addresses the possible solutions for further reducing the jitter and the transmission latency difference. Section IV-A reviews the LS mechanism that considers the release time differences of periodic VLs that can reduce the jitter, and a method to analyse the jitter, whereas Section IV-B presents a method to bound the D_{TLD} term of (5).

A. Local Synchronization

The jitter upper bound estimation is based on the worst case, where it is assumed that frames in all VLs arrive simultaneously. However, this situation will not happen when some applications are executed sequentially on a single processor, which is common in practice. For example, the AFDX ESs are often paired with the ARINC653 operating system (OS). Thus frames of certain VLs are produced with a static and periodic manner on distinct pre-defined time slots. Therefore, LS is a possible solution to mitigate the jitter by exploring the periodic VL characteristics in source ES. Relevant research on LS in ESs can be found in [38] and [21]. It is proposed in [38] to reduce the end-to-end delay by taking into account partition scheduling, which helps to eliminate impossible scenarios (all periodic VLs simultaneously send frames to a scheduler) by introducing a correlation between the release of VLs in each ES. In [21], all VLs are assumed to be periodic and the end-to-end delay upper bounds are improved by taking into account offsets between periodic VLs. In this section, we further develop this idea while leveraging the staircase arrival curve to improve the results based on [38] and [21]. First, we consider a scenario with two VLs, in which the minimum release time difference between adjacent frames is analyzed and a condition to avoid the interference between the two VLs

is deduced. Then the analysis is extended to a general case. Finally, an example is given to illustrate the effect of LS.

A periodic VL, e.g., VL_i , can be characterized by a triplet $\{T_i, \sigma_i, O_i\}$, where T_i is the period of the flow and $T_i = BAG_i$, σ_i is equal to the MFS plus 20 bytes overhead during transmission on physical links, and O_i represents a time offset of the first frame. Then the staircase model is applied in the following analysis. In such a model, the jitter of a frame is caused by the residual bytes left for transmission when the frame arrives at the scheduler.

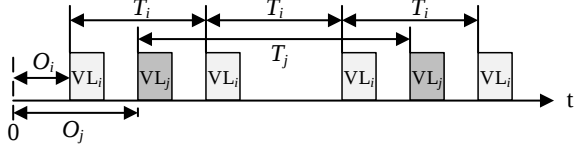


Fig. 6: An example of two periodic VLs with offsets.

1) *A special case with two VLs:* First, we just consider the case with two periodic VLs as shown in Fig. 6, in which VL_i starts earlier than VL_j . Denote by D_{diff} the minimum release time difference between adjacent frames of the two periodic VLs. If D_{diff} is large enough for a frame, either in VL_i or in VL_j , to be transmitted, the two VLs have no interference. Note that it can happen that there is more than one frame from VL_i between two consecutive frames of VL_j . Since T_i is enough for a frame of VL_i to be transmitted, only the adjacent frame ahead of VL_j is taken into account. Furthermore, since the periods are powers of 2, the Greatest Common Divisor (GCD) of the periods corresponds to the operation “min” [39]. Therefore, D_{diff} is given by

$$D_{diff} = \min \left(|O_i - O_j| \bmod \min\{T_i, T_j\}, \min\{T_i, T_j\} - (|O_i - O_j| \bmod \min\{T_i, T_j\}) \right). \quad (7)$$

A sketch of the proof is provided below. Suppose that $f_i(n_i) = O_i + n_i T_i$ and $f_j(n_j) = O_j + n_j T_j$ represent frame starts of VL_i and VL_j , respectively. n_i and n_j are two nonnegative integers. Assume that $T_i \geq T_j$. According to the periodicity, we have that $T_i = k T_j$, where $k = 2^n$ and $n \in \{0, \dots, 7\}$. Then

$$\begin{aligned} |f_i(n_i) - f_j(n_j)| &= |O_i - O_j + n_i T_i - n_j T_j| \\ &= |O_i - O_j + (k n_i - n_j) T_j|. \end{aligned}$$

In this case, the release time difference should be smaller than T_j . Then the minimum of $|f_i(n_i) - f_j(n_j)|$ is either $|O_i - O_j| \bmod T_j$ or $T_j - (|O_i - O_j| \bmod T_j)$. Thus, (7) holds true.

Let $\frac{\max\{\sigma_i, \sigma_j\}}{C}$ represent the upper bound of transmit time for any frame from either VL_i or VL_j . Thus, if the condition $\frac{\max\{\sigma_i, \sigma_j\}}{C} \leq D_{diff}$ is satisfied, the two VLs have no influence on each other, although they share the output port of the same source ES. Once the VLs are delivered from the source ES, they are serialized. If the frame dispatched earlier does not experience any congestion in all switches along its path, it will never interfere with a frame released later. Although jitter may be introduced by a frame dispatched earlier, when congestion happens, it is due to the jitter propagation caused by other VLs.

Obviously, LS contributes to reduce the jitter, as the number of interfering VLs to take into account is diminished.

In addition, if the start time order of the two VLs is fixed, e.g., VL_i always starts ahead of VL_j , then the requirements can be relaxed. In this scenario, VL_j has no influence on VL_i if the condition $\frac{\sigma_j}{C} \leq (\min\{T_i, T_j\} - ((O_j - O_i) \bmod \min\{T_i, T_j\}))$ holds and VL_i has no influence on VL_j if the condition $\frac{\sigma_i}{C} \leq ((O_j - O_i) \bmod \min\{T_i, T_j\})$ can be met. This method can be extended when a set of VLs (> 2) is considered and the corresponding analysis is given as follows.

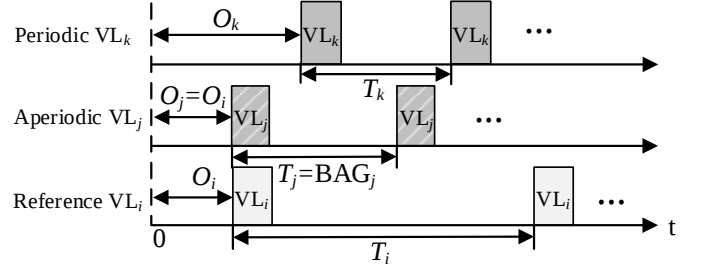


Fig. 7: An example of multiple VLs with offsets.

2) *A general case analysis:* In AFDX networks, all VLs, no matter periodic or aperiodic, are regulated by the BAG. However, unlike periodic flows, arrival of frames on an aperiodic VL can happen at any time as long as they respect the minimal inter-arrival time between each frame, which is the BAG of the VL. This means that an aperiodic VL does not have a fixed offset. Thus, as shown in Fig. 7, to take into account aperiodic VLs in our analysis, the worst case interference the aperiodic VL_j can have on the VL_i is analyzed as follows. If a frame of VL_j arrives just before the first frame considered for VL_i , then the next frame of VL_j arrives as early as its BAG allows. A safe approximation of VL_j is then to consider it as a periodic VL with the same BAG as VL_j and the offset of VL_i . Note that priorities between VLs are not considered in the present analysis, and thus frames cannot preempt each other in an AFDX network.

It is assumed in this paper that all the VLs have an equal priority, and the frames are served with a FIFO policy if contentions occur. Suppose that there is an aggregated flow of VLs, $\mathcal{I} = \{VL_i, VL_j, VL_k, \dots\}$, as shown in Fig. 7 and VL_i is the periodic VL of interest. In a period of T_i , there may be more than one frame belonging to the VLs other than VL_i , e.g., VL_j . Denote by $D_{i,j}$ the minimum release time difference between adjacent frames of VL_i and VL_j , where the frame of VL_j is ahead of that of VL_i . Obviously, $D_{i,j} = 0$, when VL_j is an aperiodic VL. According to (7), $D_{i,j} < \min\{T_i, T_j\}$. Then the number of VL_j within T_i is upper bounded by

$$N_j = \left\lceil \frac{T_i - D_{i,j}}{T_j} \right\rceil. \quad (8)$$

If $T_i \leq T_j$, $N_j = 1$, and it means that there is at most one frame of VL_j within T_i , which is true due to the VL regulation as shown in Fig. 2. If $T_i > T_j$, suppose that $T_i = m T_j$, where $m = 2^n$ and $n \in \{0, \dots, 7\}$. Then the number of VL_j within T_i can be obtained according to the rule of VL regulation.

Thus, there are $m+1$ frames of VL_j within T_i when $D_{i,j}=0$ and there exist m frames of VL_j within T_i when $0 < D_{i,j} < T_j$. Thus, (8) holds true.

For $VL_j \in \mathcal{I}, j \neq i$, let $VL_j(q)$, $q = 1, \dots, N_j$, be the q th frame before a frame of VL_i in the worst case in T_i . For each pair $(VL_i, VL_j(q))$, $q = 1, \dots, N_j$, the release time difference is computed individually. Denote by $\mathbb{D}_j = \{D_{i,j}(q), \dots, D_{i,j}(2), D_{i,j}(1)\}$ a set of release time differences for each $j \neq i$. Specifically, $D_{i,j}(1)=0$, when VL_j is an aperiodic VL. Define D as a sorted vector such that $D = \{D^{(1)}, \dots, D^{(l)}\} = \biguplus_{VL_j \in \mathcal{I}, j \neq i} \mathbb{D}_j$ and $l = \sum_{VL_j \in \mathcal{I}, j \neq i} N_j$ is the total number of frames between two consecutive frames of VL_i in the worst case. Then $D^{(k)} = D_{i,j}(\cdot)$, $1 \leq k \leq l$, and in this case, $\sigma^{(k)} = \sigma_j$, where σ_j is the maximum frame size associated with VL_j .

Let $M_i^{(l)}$ be the residual bytes left for transmission in the worst case when a frame of VL_i arrives. In other words, $M_i^{(l)}$ is the total number of bytes that contributes to the jitter of VL_i in the worst-case scenario. We then have

$$M_i^{(l)} = \left(M_i^{(l-1)} + \sigma^{(l)} - D^{(l)}C \right)^+, \quad l \geq 1. \quad (9)$$

The proof of (9) is given as follows. The recursive computation of $M_i^{(l)}$ starts by setting $M_i^{(0)} = (\sigma_i - (T_i - D^{(1)}C)^+)$. Then for each recursive step, $M_i^{(k)}$, $1 \leq k < l$, is computed. $M_i^{(k)}$ involves three parts: the residual bytes when the previous frame arrives $M_i^{(k-1)}$, the maximum frame size of the previous frame $\sigma^{(k)}$, and the data size that can be transmitted in $D^{(k)} - D^{(k+1)}$ with the rate C . As the residual bytes are nonnegative, it leads to

$$M_i^{(k)} = \left(M_i^{(k-1)} + \sigma^{(k)} - (D^{(k)} - D^{(k+1)})C \right)^+, \quad (10)$$

where $D^{(k)}, D^{(k+1)} \in D$. Finally, the computation stops when $k = l$ and we have $M_i^{(l)} = \left(M_i^{(l-1)} + \sigma^{(l)} - D^{(l)}C \right)^+$.

Let $\alpha_{\mathcal{I}}$ denote the arrival curve of the considered VL aggregate $\mathcal{I}$. Under the staircase model, the arrival curve of a single VL can be expressed by the following function:

$$u_{T,\tau,\sigma}(t) = \left\lfloor \frac{t+\tau}{T} \right\rfloor \sigma + \sigma; \quad t, \tau \geq 0; \quad T, \sigma > 0, \quad (11)$$

where σ is the burst transmission of the VL, τ is a variable delay, $\sigma = L_{\max} + 20$, and $T = \text{BAG}$. Obviously, both periodic and aperiodic VLs in source ESs can be upper bounded by $u_{T,0,\sigma}(t)$ in the worst case, due to the BAG regulation. Suppose that the service rate offered to aggregate $\mathcal{I}$ is C and $C \geq \sum_{VL_j \in \mathcal{I}} \frac{\sigma_j}{T_j}$. As discussed in Section III-C, the constant rate service, $\beta(t) = Ct$, can be applied in source ESs. Based on the above analysis, $M_i^{(l)} + \sigma_i$ is the worst-case backlog of aggregate $\mathcal{I}$, when a frame of VL_i arrives. Obviously, $M_i^{(l)} + \sigma_i$ is upper bounded by the backlog upper bound of aggregate $\mathcal{I}$.

Based on Theorem 1.4.1 of [17], we have:

$$\begin{aligned} M_i^{(l)} + \sigma_i &\leq \sup_{t \geq 0} \left\{ \sum_{VL_j \in \mathcal{I}} u_{T_j,0,\sigma_j}(t) - \beta(t) \right\} \\ &= \sup_{t \geq 0} \left\{ \sum_{VL_j \in \mathcal{I}} \left\lfloor \frac{t}{T_j} \right\rfloor \sigma_j + \sum_{VL_j \in \mathcal{I}} \sigma_j - Ct \right\} \\ &\leq \sup_{t \geq 0} \left\{ \sum_{VL_j \in \mathcal{I}} \frac{t}{T_j} \sigma_j + \sum_{VL_j \in \mathcal{I}} \sigma_j - Ct \right\} \\ &= \sum_{VL_j \in \mathcal{I}} \sigma_j. \end{aligned}$$

Indeed, taking the LS into account may lead to an aggregate arrival curve less conservative than the one obtained by a direct summation of the arrival curves of individual flows. Specifically, when a periodic VL_i is taken as the benchmark, it is assumed that a frame of VL_i is the first one arrived in an arbitrary interval $[s, t]$. Then, for any VL_j in the aggregated flow $\mathcal{I}$, the maximum number of frames arrived in this interval is given by:

$$N_j = \left\lfloor \frac{t-s-\tau_j}{T_j} \right\rfloor + 1,$$

where $\tau_j = T_j - D_{i,j}$ when $i \neq j$, and $\tau_j = 0$ when $i = j$. According to [40], let $R(t)$ be the aggregated flow, then:

$$\begin{aligned} R(t) - R(s) &\leq \sum_{VL_j \in \mathcal{I}} \left(\left\lfloor \frac{t-s-\tau_j}{T_j} \right\rfloor + 1 \right) \sigma_j \\ &= \sum_{VL_j \in \mathcal{I}} u_{T_j,-\tau_j,\sigma_j}(t-s) \\ &:= \tilde{\alpha}_{\mathcal{I}}^i(t-s). \end{aligned}$$

Note that since $\tau_j > 0$ for any $j \neq i$, it is clear from (11) that $u_{T_j,-\tau_j,\sigma_j}(t)$ does not define an arrival curve of VL_j . Furthermore, taking into account the worst-case residual bytes in the transmission of VL_i , the arrival curve for the aggregated flow $\mathcal{I}$, taking VL_i as the benchmark, can be given by:

$$\alpha_{\mathcal{I}}^i(t) = \tilde{\alpha}_{\mathcal{I}}^i(t) + M_i^{(l)} = \sum_{VL_j \in \mathcal{I}} u_{T_j,-\tau_j,\sigma_j}(t) + M_i^{(l)}, \quad t \geq 0, \quad (12)$$

where $\tau_j = T_j - D_{i,j}$ for $i \neq j$ and $\tau_j = 0$ for $i = j$.

Then $\alpha_{\mathcal{I}}^i(t)$ can be used for the end-to-end delay analysis of VL_i combined with the approach presented in [14]. If $M_i^{(l)}$ can be reduced by applying LS, the introduced jitter is mitigated accordingly. Consequently, the end-to-end delay upper bound can be reduced accordingly.

3) *An illustration example:* To illustrate the effect of LS, we consider an example of 3 VLs with $\sigma=1500$ bytes and a BAG of 1 ms. Their offsets are $O_1=0$, $O_2=100 \mu\text{s}$ and $O_3=200 \mu\text{s}$, respectively. Then the set of release time difference is given in Table I.

TABLE I: Time Interval between Frames

VL Pairs (i, j)	1, 2	1, 3	2, 1	2, 3	3, 1	3, 2
$D_{i,j} (\mu\text{s})$	900	800	100	900	200	100

Based on (8), it can be obtained that $l = 2$. By considering LS, the residual bytes when the frame of VL_1 under analysis arrives can be computed by:

$$M_1^{(1)} = \left(M_1^{(0)} + \sigma^{(1)} - (D^{(1)} - D^{(2)})C \right)^+,$$

$$M_1^{(2)} = \left(M_1^{(1)} + \sigma^{(2)} - D^{(2)}C \right)^+ = 0,$$

where $\sigma^{(2)} = \sigma^{(1)} = 1500$ bytes, $D^{(1)} = 900 \mu s$, $D^{(2)} = 800 \mu s$, $M_1^{(0)} = 250$ bytes, and $C = 100$ Mbps. Similarly, we can get $M_2^{(2)} = 250$ bytes and $M_3^{(2)} = 500$ bytes for VL_2 and VL_3 , respectively. In contrast, by applying the conventional approaches without LS, the residual bytes for each VL are $M_1^{(2)} = M_2^{(2)} = M_3^{(2)} = 3\sigma = 4500$ bytes in the worst case. In this case, the residual bytes are significantly reduced with LS.

It is shown in the above analysis that LS contributes to reduce the residual bytes for a periodic VL. Consequently, the jitter for the periodic VL is mitigated, which contributes to avoiding the incidence of frame sequence inversion. The other periodic VLs also benefit from this approach. In fact, the jitters for the other periodic VLs may be further mitigated by properly allocating the offsets of periodic VLs. Moreover, the aperiodic VLs can also benefit from the LS. Unlike the worst case-based analysis where all the periodic and aperiodic VLs are supposed to arrive simultaneously, with the LS mechanism, the number of interfering VLs or the amount of interfering backlog for aperiodic VLs can be reduced. Compared with the approach in [21], our jitter upper bounds are obtained by analyzing the residual number of bytes with respect to LS, instead of using the safe arrival curve of the aggregated flows. Therefore, tighter upper bounds can be achieved. For example, the end-to-end delay upper bound of $v1$ from $e1$ to $e6$ in the case study presented in [21] can be reduced to $96 \mu s$ from $116 \mu s$ due to the fact that the VLs $v1$ and $v2$ have no influence on each other according to our model. It is worth noting that LS can also help to eliminate certain impossible scenarios in switches to further improve jitter estimation as presented in [21]. This feature is taken into account in the case study presented in Section V.

B. Transmission Latency Difference Minimization

It can be seen from (6) that the transmission latency difference between two continuous frames defined in (5) is another factor that may cause sequence inversion. Thus, we consider a scheme aiming at reducing the second term on the left-hand side of the inequality (6) by TLDM.

In traditional delay analysis, much attention has been paid to the MFS, as the minimum length has no effect on the jitter upper bounds. Normally, the default minimum length predetermined by the specification is assigned to each VL. In fact, this makes the transmission latency difference even larger according to (5). In the worst case, the frames with the MFS and the frames with minimum length are delivered alternately as shown in Fig.8. In this scenario, half the received frames experience the worst-case transmission latency, which increases the occurrence probability of the sequence inversion phenomenon.

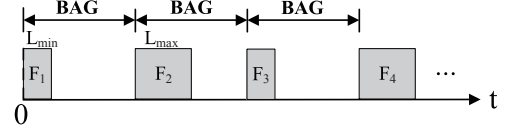


Fig. 8: An example of transmission latency difference in the worst case.

Based on (5), for a predefined VL routing scheme, the transmission latency mitigation can be formulated as an optimization problem aimed at minimizing the maximum size difference between two continuous frames:

$$\min_i \max (L(i) - L(i \oplus 1))^+, \quad (13)$$

where the wrap-around operation $i \oplus 1$ is defined in (1). It can be further simplified as the following problem:

$$\min_i (L_{\max}(i) - L_{\min}(i + 1)). \quad (14)$$

Obviously, the optimal value of (13) and (14) is zero. It can be achieved when every frame in a VL is set to the identical frame size, $L_{\max} = L_{\min}$. However, the configuration for each VL in practice cannot be simply assigned in such a way due to diverse requirements and data source types. In this case, the transmission latency difference can be mitigated by properly selecting the value of $L_{\min}$, and both (13) and (14) are upper bounded by $L_{\max} - L_{\min}$. Even though the optimum of (13) or (14) is not achieved, the TLDM helps control the transmission latency difference by carefully selecting $L_{\min}$. Therefore, this approach contributes to satisfy the inequality (6) so that the sequence inversion can be avoided.

To illustrate how TLDM contributes to reduce the transmission latency difference between two consecutive frames, we consider a case in which a VL has a MFS of 600 bytes and a default minimum length of 64 bytes. The VL traverses 2 switches to reach its destination. Then the transmission latency difference can be up to $\frac{(600-64) \times 8}{C} \times 3 = 128.64 \mu s$, if $C = 100$ Mbps. When $L_{\min}$ is 500 bytes, the upper bound of transmission latency difference can be reduced to $24 \mu s$, less than 20% compared with $128.64 \mu s$. The optimal value of transmission latency difference is zero and it can be achieved with $L_{\min} = 600$ bytes. Since the minimum frame length is not used during the worst case delay analysis, enforcing $L_{\max} = L_{\min}$ does not change the performance of the network in the worst case. This example confirms that the specification of frame size has an impact on the transmission reliability and should be carefully designed.

Design rules allowing improving transmission reliability can be generally given as follows:

- assign identical or similar frame size for all the frames in a VL;
- if the message is too large and needs to be fragmented, assign an equal size to each fragment;
- if Sub-VL aggregation is performed as in [12] to optimize bandwidth utilization, the pre-processing is required first to assort Sub-VLs with similar frame size into a group. Then Sub-VL aggregation strategy is applied to each group to avoid large transmission latency differences.

C. Discussion of LS and TLDM

The LS approach aims at mitigating the impossible interference between VLs by considering the synchronization mechanism in source ESs. However, it also adds latencies by introducing offsets for periodic VLs in general. In this case, a trade-off should be made by considering the practical constraints and design preferences. Furthermore, in the analysis of the LS approach, release jitters given by specific applications for periodic VLs should be taken into account, which is one of our directions for future work. In this case, the release time difference, which leads to the worst-case residual bytes, should be applied to handle this issue.

The TLDM approach focuses on the size difference between two continuous frames, which can be reduced by properly assigning the minimum frame size, $L_{\min}$, of a VL. As this approach does not change the maximum frame size, the worst-case performance will not be affected. However, as the data carried by VLs are generated by different functions, padding data is required when the data size is smaller than $L_{\min}$, which will introduce an overhead for data transmission.

In conclusion, these two methods can be applied separately or in combination to achieve a better performance. It is worth noting that there is no conflict between the grouping technique and LS/TLDM. Indeed, they can be employed jointly as illustrated in the case study in Section V.

V. CASE STUDY

In this section, the proposed approaches are illustrated by a case study with a network shown in Fig. 9, which is adapted from a benchmark configuration reported in [21], [27], [30], [41]–[43] while including more VLs. The VL parameters are specified referring to the realistic cases in the references that are given in Table II and Table III, in which VL1-8 are periodic VLs and each period T is equal to its BAG. As the VL parameters highly depend on the application requirements in realistic AFDX networks, they must be specified on a case-by-case basis.

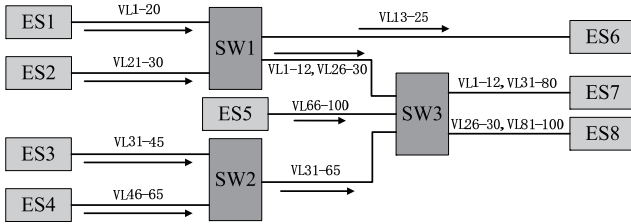


Fig. 9: An example of VL management in source ESs and the end-to-end transmission schematic.

TABLE II: Parameters of periodic VLs

VL	1	2	3	4	5	6	7	8
BAG (ms)	1	4	4	2	4	2	4	64
σ (byte)	620	84	520	820	320	140	1020	520
O (ms)	0.1	0.5	0.5	0.5	0.8	0.8	1.5	1.5
Number of Hops	3	3	3	3	3	3	3	3

TABLE III: Parameters of aperiodic VLs

VL	9-12	13-20	21-25	26-30	31-35	36-40
BAG (ms)	2	1	4	16	2	1
σ (byte)	250	84	620	480	100	84
Number of Hops	3	2	2	3	3	3
VL	41-45	46-55	56-65	66-80	81-90	91-100
BAG (ms)	2	2	4	2	1	8
σ (byte)	260	180	220	84	100	320
Number of Hops	3	3	3	2	2	2

In this case study, we assume that the physical link offers a constant rate $C = 100$ Mbps. Suppose that VL1 is the data flow of interest. First, the end-to-end jitter upper bound obtained from the staircase model is 1.248 ms, in which the grouping technique is also taken into account. As the LS is not applied at this step, VL1 is assumed to be influenced by other VLs that share the same output ports either in source ES or in switches. In addition, the minimum frame size of VL1, $L_{\min}$, is assigned to 64 bytes as default. Since VL1 traverses two switches in its communication path, the transmission latency difference in the worst case can be obtained with $(L_{\max} - L_{\min}) \times 8/C \times n$, where $n = 3$. In this scenario, the maximum transmission latency difference is 128.64 μ s, which is more than 10% of its BAG. Considering a transmission latency difference of 128.64 μ s, the worst-case delay difference can be up to 1.377 ms, which clearly exceeds its BAG, the safe upper bound. In the following analysis, the approaches presented in Section IV are applied step by step to mitigate the delay differences.

The LS focuses on the periodic VL1-8. According to Table II, VL1 is always ahead of VL2-8, then the temporal interval between each pair is calculated based on (7) and listed in Table IV, in which the required transmission time for VL2-8 is also given. We further compute the residual bytes, which may introduce a jitter into VL1. The calculation is based on (9). In this example, $M_1^{(l)} = 0$, where $l=7$. In other words, VL2-8 have sufficient time to be delivered before the arrival of VL1 and hence, they have no impact on VL1 in terms of jitter. The jitter upper bound can be further improved by reducing the number of involved VLs. The obtained result is 0.974 ms, which is less than its BAG. In this case, its burst does not introduce jitter in the worst case when the staircase arrival curve model is employed. Therefore, the end-to-end jitter could be reduced by 0.050 ms, and then the upper bound becomes 0.924 ms.

TABLE IV: Temporal Interval between Frames and the Transmission Time Requirement (in μ s)

VL Pairs (i, j)	1, 2	1, 3	1, 4	1, 5	1, 6	1, 7	1, 8
$D_{i,j}$	600	600	600	300	300	600	600
σ_j/C	6.72	41.6	65.6	25.6	11.2	81.6	41.6

Till now, although a large improvement has been achieved, the requirement cannot be met when considering the fixed transmission latency difference of 128.64 μ s in the worst case.

The sum of the jitter and the latency difference is 1.053 ms, which is very close to the safe upper bound.

Thereafter, the TLDM is applied. As illustrated in Section IV-B, the fixed transmission latency difference can be improved by more than 80% if $L_{\min}$ is 500 bytes, and then the transmission delay difference is 0.948 ms < 1 ms. The optimal result for transmission latency difference is zero, when the VL guarantees that all the frames have an identical frame size. With either of the two configurations, it can be verified that the transmission delay difference will not exceed the BAG and the sequence inversion will never happen for VL1.

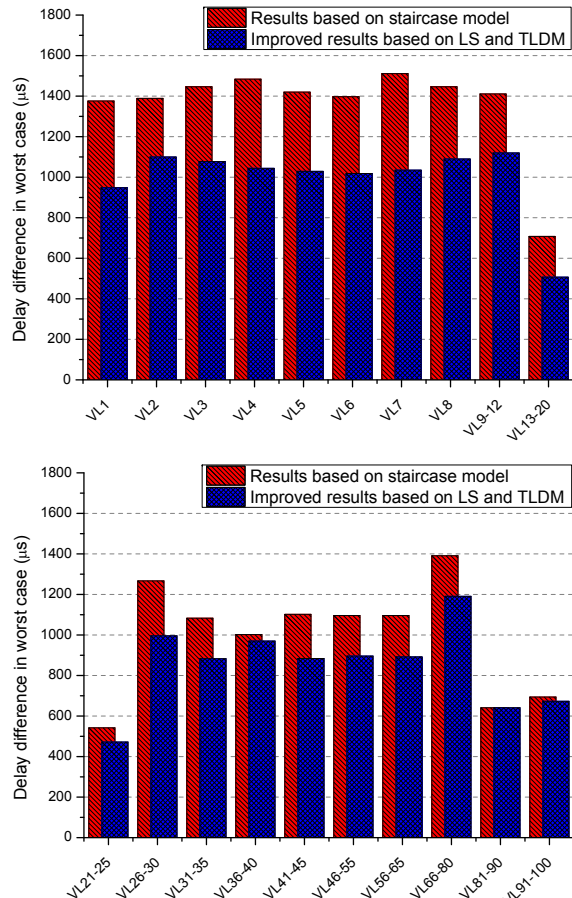


Fig. 10: Delay differences in the worst case for all the VLs.

Finally, the delay differences in the worst case for all other VLs are computed using the staircase model and the approach based on LS and TLDM, respectively. The improvement with TLDM is achieved under the condition that the frame size difference is restricted within 100 bytes. As shown in Fig. 10, there is a potential risk of failures for the redundant transmission of VL1 and VL36-VL40, as the delay differences obtained based on the staircase model are larger than their BAGs (1 ms). When the approaches of LS and TLDM are applied, the delay differences for all the VLs meet the condition (6). Ultimately, the reliability of AFDX networks is enhanced, as the failures due to sequence inversion have been eliminated, according to the analysis presented in Section III-C.

It is worth noting that ultimately, one can assign a VL to each application. Therefore, the constraints on frame size

difference can always be satisfied by adding VLs. In essence, this amounts to a trade-off between the reliability and the bandwidth utilization efficiency. Furthermore, the work presented in this paper is based on network calculus and hence, it can be easily extended to the analysis of networks with different topologies, size, and configurations. Furthermore, the LS is applied in a source ES and TLDM considers one VL at a time. Thus, both of them can be applied in scalable networks without any difficulty, while providing a consistent performance.

VI. CONCLUSION

In this paper, sequence inversion, a potential failure source in the redundant transmission management of AFDX networks is addressed, and a quantitative analysis of this phenomenon is carried out. It has been found that the main reasons for the sequence inversion phenomenon in redundant networks are the jitter and frame size differences. In order to eliminate the resulting potential failures, two approaches are developed. They allow tightening jitter estimation by reducing the number of VLs involved and diminishing transmission latency differences. A case study is carried out to illustrate the proposed approaches. The results confirm that the developed approaches are feasible and effective.

It is worth noting that the focus of the present work was put on enhancing the reliability of AFDX networks. In future work, the degree of automation of the proposed methods will be considered through the development and the implementation of suitable tools. Furthermore, more scheduling policies, e.g., fixed priority scheduling, can be considered to tighten the jitter estimation and then to prevent potential failures.

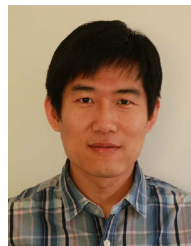
ACKNOWLEDGMENT

The authors thank the Editor-In-Chief, the Associate Editor, and the anonymous reviewers for their constructive comments, which helped to improve the quality and the presentation of this paper. This work was supported in part by CSC, NSERC-CRIAQ CRD project AVIO402, MITACS Acceleration Quebec program, MDEIE-CRIAQ Quebec-China project, and the industrial partners Thales Canada Inc. and Bombardier Aerospace.

REFERENCES

- [1] W. Ni, I. B. Collings, R. Liu, and Z. Chen, "Relay-assisted wireless communication systems in mining vehicle safety applications," *IEEE Trans. Ind. Informat.*, vol. 10, no. 1, pp. 615–627, Feb. 2014.
- [2] A. Kavousi-Fard, M. A. Rostami, and T. Niknam, "Reliability-oriented reconfiguration of vehicle-to-grid networks," *IEEE Trans. Ind. Informat.*, vol. 11, no. 3, pp. 682–691, Jun. 2015.
- [3] F. Dobsław, T. Zhang, and M. Gidlund, "End-to-end reliability-aware scheduling for wireless sensor networks," *IEEE Trans. Ind. Informat.*, pp. 1–10, 2015.
- [4] S. A. Asghari, H. Taheri, H. Pedram, and O. Kaynak, "Software-based control flow checking against transient faults in industrial environments," *IEEE Trans. Ind. Informat.*, vol. 10, no. 1, pp. 481–490, Feb. 2014.
- [5] C. M. Fuchs, "The evolution of avionics networks from ARINC 429 to AFDX," *Innovative Internet Technologies and Mobile Communications (IITM), and Aerospace Networks (AN)*, vol. 65, 2012.

- [6] A. Basu, S. Bensalem, M. Bozga, B. Delahaye, A. Legay, and E. Sifakis, "Verification of an AFDX infrastructure using simulation and probabilities," in *Runtime Verification - First International Conference*, vol. 6418, 2010, pp. 330–344. [Online]. Available: <https://hal.archives-ouvertes.fr/hal-00557717>
- [7] B. W. Harrisand and B. J. Tran, "Fiber optic AFDX for flight control systems," in *Proc. of IEEE AVFOP*, Sep. 2012, pp. 15–17.
- [8] "SAE ARP4761 Guidelines and Methods for Conducting the Safety Assessment Process on Civil Airborne Systems and Equipment," Warrendale, USA, Dec. 1996.
- [9] ARINC 664, *Aircraft Data Network Part 7 Avionics Full-Duplex Switched Ethernet Network*. AERONAUTICAL RADIO INC., 2009.
- [10] M. Grenier, J. Goossens, and N. Navet, "Near-optimal fixed priority preemptive scheduling of offset free systems," in *Proc. of the 14th RTNS*, 2006, pp. 35–42.
- [11] X. Li, "Worst-case delay analysis of real-time switched Ethernet networks with flow local synchronization," Ph.D. dissertation, 2013.
- [12] M. Li, M. Lauer, G. Zhu, and Y. Savaria, "Determinism enhancement of AFDX networks via frame insertion and Sub-Virtual link aggregation," *IEEE Trans. Ind. Informat.*, vol. 10, no. 3, pp. 1684–1695, Aug. 2014.
- [13] M. Anand, S. Vestal, S. Dajani-Brown, and I. Lee, "Formal modeling and analysis of the AFDX frame management design," in *Proc. of the 9th IEEE ISORC*, 2006, pp. 1–7.
- [14] M. Boyer, N. Navet, and M. Fumey, "Experimental assessment of timing verification techniques for AFDX," in *Proc. of ERTS, Toulouse, France*, 2012.
- [15] R. L. Cruz, "A calculus for network delay, part I: Network elements in isolation," *IEEE Trans. Inf. Theory*, vol. 37, no. 1, pp. 114–131, Jan. 1991.
- [16] R. L. Cruz, "A calculus for network delay, part II: Network analysis," *IEEE Trans. Inf. Theory*, vol. 37, no. 1, pp. 132–141, Jan. 1991.
- [17] J.-Y. Le Boudec and P. Thiran, *Network Calculus: A Theory of Deterministic Queuing Systems for the Internet*. Springer, 2001, vol. 2050.
- [18] J. Grien, "Analyse et évaluation de techniques de commutation ethernet pour l'interconnexion des systèmes avioniques," Ph.D. dissertation, Institut National Polytechnique de Toulouse, 2004.
- [19] J.-L. Scharbarg, F. Ridouard, and C. Fraboul, "A probabilistic analysis of end-to-end delays on an AFDX avionic network," *IEEE Trans. Ind. Informat.*, vol. 5, no. 1, pp. 38–49, Feb. 2009.
- [20] M. Boyer, J. Migge, and N. Navet, "An efficient and simple class of functions to model arrival curve of packetised flows," in *Proc. of the 1st International Workshop on Worst-Case Traversal Time*. ACM, 2011, pp. 43–50.
- [21] X. Li, J.-L. Scharbarg, and C. Fraboul, "Improving end-to-end delay upper bounds on an AFDX network by integrating offsets in worst-case analysis," in *Proc. of IEEE ETFA*, Sep. 2010, pp. 1–8.
- [22] J. Rox and R. Ernst, "Formal timing analysis of Full Duplex Switched Based Ethernet network architectures," in *SAE Technical Paper*. SAE International, 04 2010. [Online]. Available: <http://dx.doi.org/10.4271/2010-01-0455>
- [23] S. Martin, P. Minet, and L. George, "End-to-end response time with fixed priority scheduling: trajectory approach versus holistic approach," *International Journal of Communication Systems*, vol. 18, no. 1, pp. 37–56, 2005.
- [24] H. Bauer, J.-L. Scharbarg, and C. Fraboul, "Applying and optimizing trajectory approach for performance evaluation of AFDX avionics network," in *Proc. of IEEE ETFA*, Sep. 2009, pp. 1–8.
- [25] H. Bauer, J.-L. Scharbarg, and C. Fraboul, "Worst-case end-to-end delay analysis of an avionics AFDX network," in *Proc. of DATE*, Mar. 2010, pp. 1220–1224.
- [26] X. Li, J.-L. Scharbarg, and C. Fraboul, "Analysis of the pessimism of the trajectory approach for upper bounding end-to-end delay of sporadic flows sharing a switched Ethernet network," in *Proc. of RTNS*. Citeseer, 2011, pp. 149–158.
- [27] H. Bauer, J.-L. Scharbarg, and C. Fraboul, "Improving the worst-case delay analysis of an AFDX network using an optimized trajectory approach," *IEEE Trans. Ind. Informat.*, vol. 6, no. 4, pp. 521–533, Nov. 2010.
- [28] N. Hu, T. Lv, and N. Huang, "Applying trajectory approach for computing worst-case end-to-end delays on an AFDX network," *Procedia Engineering*, vol. 15, pp. 2555–2560, 2011.
- [29] H. Bauer, J.-L. Scharbarg, and C. Fraboul, "Worst-case backlog evaluation of avionics switched Ethernet networks with the trajectory approach," in *Proc. of the 24th ECRTS*, Jul. 2012, pp. 78–87.
- [30] G. Kemayo, F. Ridouard, H. Bauer, and P. Richard, "Optimistic problems in the trajectory approach in FIFO context," in *Proc. of IEEE 18th Conference on ETFA*, Sep. 2013, pp. 1–8.
- [31] X. Li, O. Cros, and L. George, "The trajectory approach for AFDX FIFO networks revisited and corrected," in *2014 IEEE 20th International Conference on Embedded and Real-Time Computing Systems and Applications*, Aug 2014, pp. 1–10.
- [32] G. Kemayo, F. Ridouard, H. Bauer, and P. Richard, "A forward end-to-end delays analysis for packet switched networks," in *Proc. of the 22nd International Conference on Real-Time Networks and Systems*. ACM, 2014, p. 65.
- [33] G. Kemayo, N. Benammar, F. Ridouard, H. Bauer, and P. Richard, "Improving AFDX end-to-end delays analysis," in *Proc. of IEEE 20th Conference on ETFA*. IEEE, 2015, pp. 1–8.
- [34] R. Mancuso, A. V. Louis, and M. Caccamo, "Using traffic phase shifting to improve AFDX link utilization," in *Proc. of the 15th ACM International Conference on EMSOFT*, Amsterdam, The Netherlands, 2015, pp. 256–265.
- [35] Y. Hua and X. Liu, "Scheduling heterogeneous flows with delay-aware deduplication for avionics applications," *Parallel and Distributed Systems, IEEE Transactions on*, vol. 23, no. 9, pp. 1790–1802, 2012.
- [36] J. Yao, J. Wu, Q. Liu, Z. Xiong, and G. Zhu, "System-level scheduling of mixed-criticality traffics in avionics networks," *IEEE Access*, vol. 4, pp. 5880–5888, 2016.
- [37] J. Täubrich and R. Von Hanxleden, "Formal specification and analysis of AFDX redundancy management algorithms," in *Proc. of International Conference on Computer Safety, Reliability, and Security*. Springer, 2007, pp. 436–450.
- [38] M. Lauer, J. Ermont, C. Pagetti, and F. Boniol, "Analyzing end-to-end functional delays on an IMA platform," in *ISoLA*, vol. 6415 of LNCS. Springer, 2010, pp. 243–257.
- [39] S. P. Dwivedi, "GCD computation of n integers," in *Proc. of RAECs*, March 2014, pp. 1–4.
- [40] L. Zhao, P. Pop, Q. Li, J. Chen, and H. Xiong, "Timing analysis of rate-constrained traffic in TTEthernet using network calculus," *Real-Time Systems*, vol. 53, no. 2, pp. 254–287, 2017.
- [41] M. Adnan, J.-L. Scharbarg, and C. Fraboul, "Minimizing the search space for computing exact worst-case delays of AFDX periodic flows," in *Proc. of IEEE SIES*, Jun. 2011, pp. 294–301.
- [42] T. Hamza, J.-L. Scharbarg, and C. Fraboul, "Priority assignment on an avionics switched Ethernet network (QoS AFDX)," in *Proc. of WFCs*, May 2014, pp. 1–8.
- [43] G. Kemayo, N. Benammar, F. Ridouard, H. Bauer, and P. Richard, "Improving AFDX end-to-end delays analysis," in *Proc. of IEEE ETFA*, Sep. 2015, pp. 1–8.



Meng Li received the B.E. and M.S. degrees in electronic engineering from Beijing University of Aeronautics and Astronautics, Beijing, China, in 2004 and 2007, respectively. He received the Ph.D. degree in electrical engineering from Polytechnique Montréal, Montréal, QC, Canada, in 2016.

Since June 2016, he has been with the Department of Electrical Engineering, Polytechnique Montréal, Montréal, QC, Canada, as a Postdoctoral Fellow. His research interests include fault tolerance, parallel computing, task scheduling, communication networks, and real-time systems.



Guchuan Zhu (M'07-SM'12) received the M.S. degree in electrical engineering from Beijing Institute of Aeronautics and Astronautics, Beijing, China, in 1982; the Ph.D. degree in mathematics and control from École des Mines de Paris, Paris, France, in 1992; and the graduate Diploma degree in computer science from Concordia University, Montreal, QC, Canada, in 1999.

Dr. Zhu joined École Polytechnique de Montréal, Montréal, in 2004, where he is currently a Professor in the Department of Electrical Engineering. His current research interests include control of distributed parameter systems, nonlinear and robust control, and optimization with their applications to microsystems, aerospace systems, communication networks, and smart grid.



Yvon Savaria (S'77, M'86, SM'97, F'08) received the B.Eng. and M.Sc.A. degrees in electrical engineering from École Polytechnique Montreal in 1980 and 1982 respectively. He also received the Ph.D. in electrical engineering in 1985 from McGill University. Since 1985, he has been with Polytechnique Montréal, where he is currently professor in the department of electrical engineering.

He has carried work in several areas related to microelectronic circuits and microsystems such as testing, verification, validation, clocking methods, defect and fault tolerance, effects of radiation on electronics, high-speed interconnects and circuit design techniques, CAD methods, reconfigurable computing and applications of microelectronics to telecommunications, aerospace, image processing, video processing, radar signal processing, and digital signal processing acceleration. He is currently involved in several projects that relate to aircraft embedded systems, green IT, wireless sensor network, virtual network, computational efficiency and application specific architecture design. He holds 16 patents, has published 124 journal papers and 396 conference papers, and he was the thesis advisor of 155 graduate students who completed their studies.

He was program co-chairman of ASAP2006 and the general co-chair of ASAP2007. He has been working as a consultant or was sponsored for carrying research by Bombardier, CNRC, Design Workshop, DREO, Ericsson, Genesis, Gennum, Huawei, Hyperchip, ISR, LTRIM, Miranda, MiroTech, Nortel, Octasic, PMC-Sierra, Technocap, Thales, Tundra and VXP. He is a member of the Regroupement Stratégique en Microélectronique du Québec (RESMIQ), of the Ordre des Ingénieurs du Québec (OIQ), and was a member of CMC Microsystems board since 1999 and chairman of that board from 2008 to 2010. He was awarded in 2001 a Tier 1 Canada Research Chair (www.chairs.gc.ca) on design and architectures of advanced microelectronic systems that he held until June 2015. He also received in 2006 a Synergy Award of the Natural Sciences and Engineering Research Council of Canada.



Michaël Lauer is an associate professor in real-time systems at University of Toulouse. He is conducting research at LAAS-CNRS in the Dependable Computing and Fault Tolerance group (TSF) since 2013. His areas of interest are real-time and resilient systems, adaptive fault-tolerance, and multi/many-cores processors. He received the diploma of Engineering from the Institut National Polytechnique de Toulouse (ENSEEIH) in Telecommunications and Networks. He received the Leopold Escande Award in 2012 for his PhD work on the verification of

real-time requirements in critical embedded systems. He conducted post-doctoral research at Polytechnic School of Montreal on the improvement of the criticality level of AFDX networks and the use of time-triggered architecture in modern avionics. He is contributing to several research projects with industrial partners especially in the avionics and automotive domains.