



# USENIX

THE ADVANCED COMPUTING  
SYSTEMS ASSOCIATION

## A Formal Security Analysis of CAN XL

Zhaozhou Tang, *Georgia Institute of Technology*; Khaled Serag, *Qatar Computing Research Institute*; Z. Berkay Celik, *Purdue University*; Vijay Ganesh, Saman Zonouz, and Raheem Beyah, *Georgia Institute of Technology*

<https://www.usenix.org/conference/usenixsecurity26/presentation/tang-zhaozhou>

This paper is included in the Proceedings of the  
35th USENIX Security Symposium.

August 12–14, 2026 • Baltimore, MD, USA

ISBN 978-1-939133-58-8

Open access to the Proceedings of the  
35th USENIX Security Symposium  
is sponsored by



جامعة الملك عبد الله  
للعلوم والتقنية  
King Abdullah University of  
Science and Technology

# A Formal Security Analysis of CAN XL

Zhaozhou Tang  
*Georgia Institute of Technology*

Khaled Serag  
*Qatar Computing Research Institute*

Z. Berkay Celik  
*Purdue University*

Vijay Ganesh  
*Georgia Institute of Technology*

Saman Zonouz  
*Georgia Institute of Technology*

Raheem Beyah  
*Georgia Institute of Technology*

## Abstract

For decades, the Controller Area Network (CAN) has been the backbone of in-vehicle communication. As modern vehicles integrate cameras, LiDARs, and AI components, classic CAN (CAN CC) faces growing limitations in bandwidth, functionality, and security. To fill these gaps, CAN XL was introduced as the next generation of CAN, aiming to offer longer payloads, higher bandwidth, and enhanced security.

CAN XL makes significant standard-level changes across multiple stack layers. It also introduces several security features, but it is unclear whether these are mere add-on extensions or whether the standard redesign itself tackles CAN's chronic security weakness: the MAC sub-layer. This sub-layer governs frame formats and error handling and has historically enabled many CAN CC attacks. As the industry transitions to CAN XL, the security posture of its yet-unexplored MAC sub-layer must be understood before widespread deployment.

This paper presents the first security analysis of the CAN XL standard, focusing on its MAC sub-layer. We develop a bit-precise CAN XL formal model and release it to facilitate future research. We design a formal analysis workflow guided by CAN XL's field-oriented structure to uncover vulnerabilities. Contrary to expectations, our analysis shows that CAN XL remains vulnerable to all known CAN CC MAC sub-layer issues while introducing seven new vulnerabilities, arguably worsening security. We validate them using commercial CAN XL controllers and demonstrate exploitability via two multi-stage attacks on a testbed simulating real vehicle traffic. Finally, we propose mitigations including formally verifying standard revisions that could prevent several attacks.

## 1 Introduction

The Controller Area Network (CAN) has long been the most widely used in-vehicle network protocol, enabling Electronic Control Units (ECUs) to manage vehicle operations. As modern vehicles integrate advanced features from entertainment to autonomous driving, in-vehicle networks increasingly de-

mand higher bandwidth, performance, functionality, and security. To keep up with these demands, CAN Extra Long (CAN XL) was introduced in May 2024 as a major upgrade of classic CAN (CAN CC). Its standard brings significant changes across protocol layers to offer longer payloads, higher bandwidth, and many more features. It is expected by many to become the primary CAN variant and the main bus for intelligent driver assistance systems in future vehicles [1, 4–6].

In its effort to enhance security, CAN XL introduces several features such as CANsec, tunneling, and virtual networks. However, it remains unclear whether the standard redesign itself happened with security in mind, or whether security is only confined to these add-on extensions. This question is especially acute for the MAC sub-layer—arguably CAN CC's weakest point—which governs bus access, frame formats, and error handling. It is the root cause of a major portion of CAN CC's vulnerabilities [30] and underlies many of the most prominent CAN attacks, such as full network DoS, injecting errors, mapping message sources, and shutting down target nodes [13, 24, 26–29, 33, 35]. CAN XL substantially revises this sub-layer: it introduces a new frame format, overhauls error handling (including new states and mechanisms), and makes parts of error handling configurable (enabled/disabled). To support incremental deployment, it also adds coexistence features for mixed CAN CC/FD/XL networks, increasing protocol complexity. Whether these far-reaching changes fix known MAC sub-layer vulnerabilities, maintain the status quo, or possibly introduce new ones remains an open question.

To close this gap, we conduct a security analysis focusing specifically on CAN XL's MAC sub-layer to examine whether CAN XL's standard changes eliminate any CAN CC vulnerabilities or introduce new ones. Through the analysis, we aim to answer an overarching question: *From a strict MAC sub-layer standpoint, have the changes made by CAN XL improved, maintained, or worsened security?* Conducting this analysis poses several challenges. (C1) Precision: The CAN XL MAC sub-layer operates on bits, not messages. An analysis must accurately capture bit-level protocol operations across more complex protocol fields and error handling mechanism un-

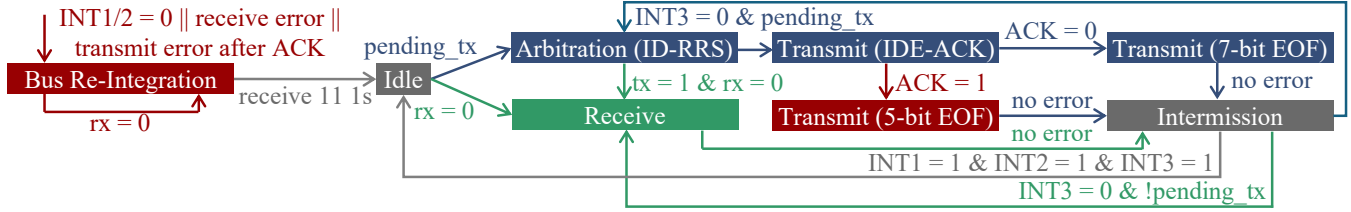


Figure 1: Simplified CAN XL error signalling disabled state diagram.

der all possible attacker actions. (C2) Generation coverage: As CAN XL is backward compatible, the analysis needs to thoroughly account for interactions between CAN XL and prior CAN generations. These include multiple frame formats, different node types, and error states. (C3) Tractability: Due to the previous two complexities, a CAN XL network has a large state space, making it difficult to perform a comprehensive exploration. Ad hoc approaches used by prior CAN CC security analysis, including manual standard reviews and testing [13, 28, 35], cannot handle these challenges.

To address them, we formally analyze the CAN XL MAC sub-layer by formally modeling a CAN XL network in the presence of an attacker. Using the nuXmv model checker [11], we verify properties relevant to new CAN XL features and known CAN CC weaknesses to discover standard-level vulnerabilities. To the best of our knowledge, this is the first formal security analysis of any part of the CAN XL standard. To address (C1), our model is bit-precise to accurately reflect the standard and thoroughly cover attacker actions. To address (C2), we model all CAN generations (CAN CC/FD/XL) to analyze cross-generation interactions. To address (C3), we employ domain-specific modeling techniques, including sequentially decomposing CAN XL according to its field-oriented design and using data and behavior abstractions to over-approximate possible protocol operations. Building on this field-level decomposition, we design a formal analysis workflow that breaks down the global protocol state space into manageable portions and analyzes them separately.

Our analysis yields counterintuitive results: CAN XL preserves all previously known CAN CC vulnerabilities and introduces seven additional ones. Exploiting them, attackers can launch many new attacks, including intercepting messages, identifying node types, new stealthy spoofing attacks, and violating data consistency extensively among nodes. These vulnerabilities also allow attacks that achieve equivalent goals to all common CAN CC attacks, even though some protocol features exploited in CAN CC have been modified in CAN XL. Examples include DoS of the whole network, shutting down target nodes, and mapping message sources. Moreover, through a comparative assessment on attack effectiveness, we find that CAN XL’s new vulnerabilities make several attacks faster and less detectable than similar CAN CC attacks. Thus, we conclude that despite CAN XL’s extensive modifications and attempts to enhance security, *the security of its MAC*

*sub-layer not only fails to improve, but arguably worsens.*

We validate the presence of all discovered vulnerabilities using off-the-shelf CAN XL protocol controllers. To demonstrate the vulnerabilities’ exploitability, we use them to construct and evaluate two multi-stage attacks from reconnaissance to execution. Finally, we propose mitigations, including formally verifying standard revisions to prevent several attacks. To inform relevant stakeholders before CAN XL’s widespread adoption, we reported these vulnerabilities to Bosch, the developer of the CAN XL standard.

In summary, we make the following contributions:

- We present the first formal model of the MAC sub-layer of all CAN generations (CAN CC, CAN FD, and CAN XL) and release it publicly to support future research.
- We formally analyze CAN XL’s MAC sub-layer by designing a formal analysis workflow guided by its field-oriented structure. To the best of our knowledge, this is the first standard security analysis of CAN XL.
- We verify that CAN XL inherits all CAN CC’s MAC sub-layer vulnerabilities and discover seven new standard vulnerabilities introduced uniquely by CAN XL.
- We conduct a comparative security analysis of CAN XL and CAN CC’s MAC sub-layer by assessing the effectiveness of attacks enabled by CAN CC and CAN XL vulnerabilities.
- We validate all discovered vulnerabilities on off-the-shelf protocol controllers. To demonstrate exploitability, we devise and evaluate two multi-stage attacks comprising the discovered vulnerabilities.
- We propose mitigations, including formally verifying standard revisions that could prevent several attacks.

## 2 Background

**CAN XL MAC Sub-Layer Operation.** CAN XL is a multi-master broadcast bus protocol. Its MAC sub-layer can operate in two significantly different error signalling configurations:

*Error Signalling Disabled:* This is newly introduced in CAN XL and can only be used when all nodes in a network support CAN XL. Fig. 1 shows its simplified state diagram. A node may start arbitration when the bus is idle, transmitting an XL frame (Fig. 2) bit by bit. When multiple nodes start arbitration concurrently, a node yields if it sends a 1 but receives a 0 when transmitting the ID. Otherwise, it continues trans-





an error frame.

**Formal Specifications.** In App. A, we show the properties’ formal definitions in computational tree logic (CTL).

## 4.2 Threat Model

We assume a remote attacker has compromised an ECU on the CAN bus and is able to execute arbitrary code. This is in line with the most common threat model assumed by CAN security literature [13, 24, 27, 28], and can be achieved through ECUs’ external connectivity channels, as demonstrated by many real-world attacks [3, 10, 12, 17, 18, 23, 25]. To thoroughly discover vulnerabilities, we consider two possible levels of attacker control over the protocol layers under this model, following a prior systematization [30]. ① HL attackers: The attacker controls the higher layer only and can transmit or receive valid frames, abiding by all link-layer rules such as arbitration and error handling. ② LL attackers: The attacker additionally controls the link layer. Previous works assumed such attackers require strong privileges (e.g., physical access) and considered them less often. Recent research [15, 19, 32] shows this assumption is not needed and LL attackers require only the same privileges as HL attackers. Remote attackers can achieve link-layer control on a significant portion of commercial ECUs using various techniques, such as by programming the ECU’s GPIO to directly access the bus (i.e., bit-banging, which is implemented by several open-source libraries [9, 34]). Attackers can read or write individual bits on the bus when it is idle or busy, but must abide by physical layer rules. Thus, they cannot modify bits on the bus arbitrarily and can only modify the bus state from 1 to 0 in SIC mode. In many attack scenarios, CAN mechanisms such as error handling can maintain key properties. For example, network-wide data consistency may be maintained by making errors visible to all nodes. However, corner cases could still lead to successful attacks, which we aim to find.

## 4.3 Model Construction

We model a CAN XL network in nuXmv, with benign nodes and an attacker attached to the channel (Fig. 4). A network may contain an arbitrary number of nodes. However, all core protocol interactions except for arbitration are between a single transmitter and multiple receivers that behave identically. All receivers have the same protocol state machine and receive the same bits. Therefore, they can be modeled as a single receiver node. In the case of arbitration, its outcome is decided by a single winner and a contender with the closest priority. Therefore, two nodes are sufficient to capture core protocol interactions in both cases, and we adopt in our model.

**MAC Sub-Layer Modeling.** The node models the MAC sub-layer based on the standard. We model it as a finite state machine (FSM) synchronized to a clock whose period equals one bit duration. Nodes read/write a bit from/to the chan-

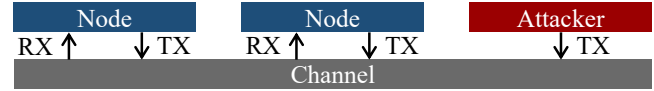


Figure 4: Model architecture.

nel every clock cycle. We make this choice because CAN XL’s lower layers provide mechanisms to ensure nodes are synchronized to the bit stream on the bus [2]. We model it using the nuXmv tool as it natively supports this synchronous FSM formalism. Formally, the MAC sub-layer is an FSM  $M = (S, I, O, S_0, \delta)$ .  $I/O$  is a finite set of inputs/outputs. For CAN XL, they are the bits received from/transmitted to the channel, with possible valuations  $\{0, 1\}$ .  $S$  is a finite set of states. In our case, it tracks a node’s transmit/receive status and current protocol field.  $S_0$  is the set of initial states (idle for CAN XL), and  $\delta : S \times I \rightarrow S \times O$  represents the state transition function. After receiving a bit, depending on its value and the current state, the node transitions to the next state and outputs a bit. CAN XL has different states and transitions depending on the error signalling configurations. Thus, the node model has two variants: error signalling disabled and error signalling enabled. They model the transitions in Fig. 1 and 3 respectively at the bit level, with transitions from intermission back to idle to model unbounded message sequences. An error signalling enabled node is nondeterministically initialized to a node type (Tbl. 1), allowing interactions between any CAN generations to be analyzed. For error signalling enabled, we do not model error states transitions (Sec. 2). Instead, a node is nondeterministically initialized to and stays in an error state. This avoids long traces to reach the error-passive state. It is sufficient for the properties we analyze because error state transitions only occur instantaneously at frame boundaries, and nodes’ behaviors during frame transmission/reception only depend on their current error states.

**Other Protocol Layers.** Since we focus on the MAC sub-layer, we do not assume specific message schedules, payload semantics, or protocols at higher layers. Instead, nodes may have a pending transmission at any arbitrary time, whose payload is nondeterministically chosen (i.e., we consider all possible messages). The physical layer is modeled using a channel that receives the outputs from all nodes and the attacker and sends the resulting bit to nodes based on physical layer rules. We assume it operates in SIC mode (0s override 1s) due to its compatibility with all frame formats, node types, and error signalling options. In Sec. 8 and 9, we discuss other layers’ impacts on MAC sub-layer vulnerabilities.

**Attacker.** We model an LL attacker (Sec. 4.2) as it possesses all abilities of an HL attacker, allowing us to discover vulnerabilities exploitable by either attacker type. We then manually separate the discovered vulnerabilities into those exploitable by LL attackers only or by either type. The attacker model nondeterministically outputs a 1 or 0 every clock cycle. This captures all possible attacker behaviors, including injecting

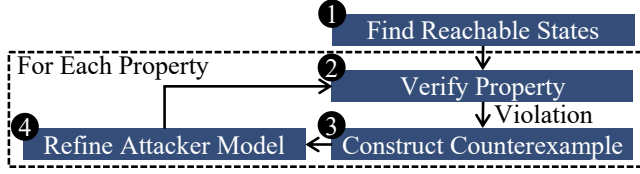


Figure 5: Analysis workflow.

valid messages or bits in any protocol fields.

**Addressing Tractability.** Since CAN XL’s MAC sub-layer operates on bits, it must be modeled at the bit level to accurately preserve behavioral details. As CAN XL frames have long and diverse fields, and we consider all possible messages with a versatile LL attacker that can manipulate any protocol fields, the protocol has a large state space. Complexity grows further in mixed networks, as different node types across CAN generations interact, each transmitting multiple frame formats and having distinct error states. These complexities lead to state explosion and make formal analysis intractable. To address this, we devise two domain-specific modeling techniques based on CAN XL’s structure:

**Sequential Decomposition:** Since frames comprise fields processed sequentially, we decompose the protocol model based on its fields. This breaks down the protocol’s state space into smaller portions without losing precision. It allows us to analyze them separately, which is more tractable than exploring all states as a whole. Formally, we decompose the protocol FSM  $M = (S, I, O, S_0, \delta)$  into  $\{M_0, M_1, \dots, M_n\}$ , where  $M_i = (S_i, I, O, S_{0,i}, F_i, \delta_i)$  is the FSM of the field  $i$ .  $S_i/\delta_i$  denotes the states/transition function restricted to  $i$ , and  $S_{0,i}/F_i$  denotes the starting/reachable ending states of  $i$ . To connect fields, a field’s starting states must equal the union of all its immediate predecessor fields’ ending states:

$$\forall j. S_{0,j} = \bigcup_{i:\text{next}(i,j)} F_i \quad (1)$$

where  $\text{next}(i, j)$  means field  $i$  is  $j$ ’s immediate predecessor. This ensures the decomposition is sound: any partial execution trace of a field is contained in a full protocol execution. As Eq. 1 ensures every starting state of each field is reachable via a trace in one of its immediate predecessor fields, a trace in field  $i$  can be extended by a trace in field  $i - 1$ . Inductively, the trace can be extended to field  $i - 2$ , etc, until the protocol’s initial state, forming a full protocol execution.

Concretely, we decompose CAN XL into: ① a precise model of the common arbitration field shared by all frame formats; for each frame format, ② an abstracted (explained next) model of the variable-value control to CRC fields coded with bit stuffing and ③ a precise model of the fixed-format ACK to intermission fields that do not use bit stuffing; and finally, ④ precise models for special frame formats (e.g., error frames) that may be transmitted from any field.

---

#### Algorithm 1 Find Reachable States

---

```

1:  $F \leftarrow \{\emptyset, \emptyset, \dots, \emptyset\}$ 
2:  $S_0 \leftarrow \{\text{idle}, \emptyset, \dots, \emptyset\}$ 
3: repeat
4:   for each field  $j$  do
5:      $S_0[j] \leftarrow \bigcup_{i:\text{next}(i,j)} F[i]$ 
6:     repeat
7:       Check AF  $F[j]$  with nuXmv
8:       if Found counterexample  $\sigma$  then
9:          $F[j] \leftarrow F[j] \cup \sigma(|\sigma| - 1)$ 
10:      end if
11:    until No more counterexample
12:  end for
13: until  $F$  does not change

```

---

**Abstraction:** We abstract away CRC computation by choosing the values of the CRC fields nondeterministically (i.e., unconstrained by the payload), which over-approximates the set of possible messages. To model CRC error detection, we over-approximate receivers’ behaviors by allowing them to nondeterministically choose whether to detect a CRC error at protocol fields where CRC checks occur. Although these abstractions limit our ability to prove CRC-algorithm-specific properties (e.g., what payload modifications cause CRC collisions), they are suitable for discovering MAC sub-layer vulnerabilities as they do not rule out any MAC sub-layer behaviors. They may introduce spurious counterexamples, but we experimentally validate all discovered vulnerabilities (Sec. 7.1) to ensure all findings correspond to real protocol behaviors.

## 4.4 Analysis Workflow

We design a formal analysis workflow leveraging CAN XL’s sequential field-oriented structure (Fig. 5).

**Find Reachable States.** We first compute the protocol’s reachable states (Fig. 5-①) using the decomposed models by solving the least fixpoint of Eq. 1. This is done once before verifying any properties. We use an automated script to perform this computation using Algorithm 1. We initialize the arbitration model’s starting states to idle and all other fields’ starting and ending states to the empty set (lines 1-2). For each model, we set its starting states to the union of all its immediate predecessor fields’ ending states (line 5). To model unbounded message exchange, the starting states of the arbitration model include ending states of the ACK-intermission models. We use nuXmv to verify the property **AF**  $F$ , which states the model always finally reaches a state within the set  $F$ . If  $F$  does not cover all ending states, nuXmv generates a counterexample. We extract its ending states, add them to  $F$ , and repeatedly verify **AF**  $F$  until no counterexamples are found (lines 6-11). We repeat this for all models until no model’s ending states change, at which point the decomposed models collectively capture all reachable states of the protocol.





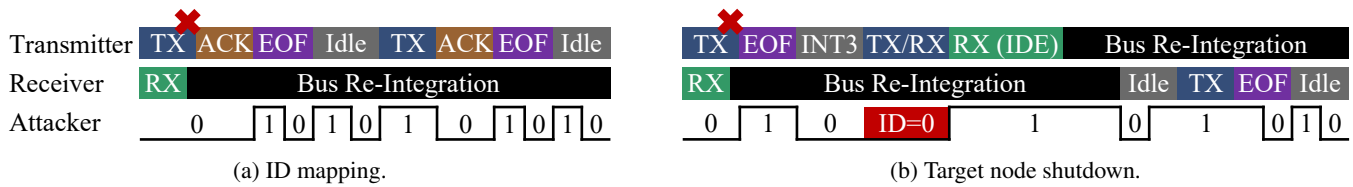


Figure 7: Tx/Rx state mismatch (V2) exploits.

keep receivers in bus re-integration.

**Step ③:** The attacker keeps sending small pulses when the bus is idle to keep receivers in bus re-integration. If he receives an ID, it must be from the same transmitter as the initially chosen ID. The attacker records this and repeats step ② on the message, but without injecting an error in it.

**Step ④:** The attacker continues step ③ for a duration longer than the longest message period recorded in step ① to ensure he can observe all periodic IDs. The attacker then stops sending pulses to bring all nodes back into normal operation.

**Step ⑤:** The attacker picks an unmapped ID, performs steps ② to ④, and repeats until all periodic IDs are mapped.

**Step ⑥:** To map aperiodic IDs, an attacker uses step ② to shut down all nodes except its transmitter when he sees an aperiodic ID. He performs step ③ until he sees a mapped periodic ID and maps the aperiodic ID to the same transmitter.

**Impact.** ID mapping is useful for reverse engineering, as information on what messages each node transmits is proprietary. It can help attackers identify each node's functionality to launch targeted attacks. For CAN CC, the same goal can be achieved by exploiting error states (Va), which are removed from CAN XL's error signalling disabled specifications. However, CAN XL's new features make it vulnerable to the same attack.

**Exploit 2: Target Node Shutdown.** V2 also allows an LL attacker to disable a single node (Fig. 6b). With V2.1, it enables persistent target node shutdown using these steps (Fig. 7b):

**Step ①:** The attacker injects an error in the victim's message and sends a 0 in INT3 to keep receivers in bus re-integration.

**Step ②:** After INT3, the attacker sends a partial message with ID 0 ending with a recessive IDE (Fig. 2) to win arbitration. As a receiver, the victim detects a form error at the recessive IDE and enters bus re-integration.

**Step ③:** The attacker waits for 10 bits until other nodes resume normal operation. The attacker then sends short pulses in every message's EOF and whenever the bus is idle to keep the victim in bus re-integration.

**Impact.** An attacker can use this to launch targeted DoS attacks by shutting down a node controlling critical system functions. Defenses can use it to disable an attacker node and prevent it from impacting the system's operation. The same can be achieved for CAN CC exploiting error states (Va), which are not present in CAN XL when error signalling is disabled.

**Exploit 3: Message Replacement.** Exploiting V2 with V1, LL attackers can launch a new attack not possible for CAN CC. An attacker first intercepts (V1) a message, then exploits

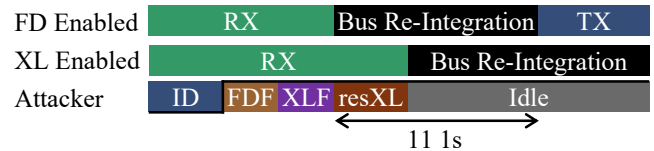


Figure 8: Mixed network state mismatch (V3.1).

V2.3 to send a 0 in INT3 after the intercepted message and push its transmitter into bus re-integration. As soon as receivers resume operation while the transmitter is still in bus re-integration, the attacker injects a fake message with the same ID as the transmitter's message.

**Impact.** This is a stealthy spoofing attack that can evade several types of defenses, such as those monitoring message timings [31,38] or letting each ECU monitor for its ID [14,22]. This is because the attack effectively replaces the original message with a fake one as receivers only see the fake message transmitted at almost the expected time. Moreover, the transmitter in bus re-integration cannot observe the fake message.

**Exploit 4: Data Consistency Violation.** V2 allows LL attackers to compromise network-wide data consistency in many ways, e.g., disabling selected nodes for arbitrary durations to prevent them from receiving certain messages or inject messages that only the remaining ones can receive.

**Impact.** This violates CAN's fundamental network-wide data consistency property. It undermines higher-layer control and protocols that rely on consistency at the lower layer, causing inconsistent system states and control decisions.

## 5.2 Error Signalling Enabled Vulnerabilities

**V3. Mixed Network State Mismatch.** In a mixed network, we discover scenarios violating  $\mathcal{P}_3$ , where some nodes are in bus re-integration while others are not when a message arrives, even if they support the frame. This also violates  $\mathcal{P}_4$  as only nodes not in bus re-integration validate the message.

**(V3.1) Partial Frames:** The standard specifies that different node types detect unsupported frames at distinct protocol fields (Sec. 2). Partial unsupported frames ending with 1s cause them to enter and leave bus re-integration at different times (e.g., FD enabled nodes enter and leave bus re-integration earlier than XL enabled nodes by one bit when receiving a partial frame ending with 1 in resXL—Fig. 8). This timing mismatch may cause some node types to be in bus

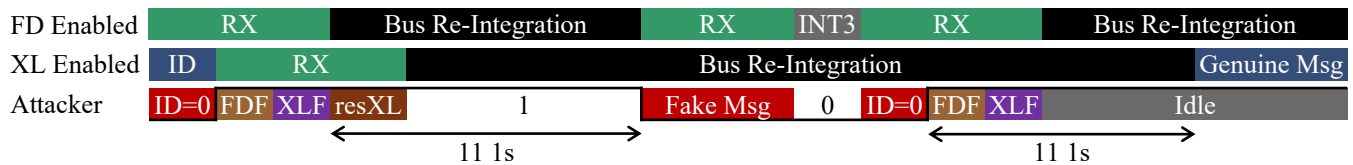


Figure 9: Mixed network state mismatch (V3) exploit: inconsistent frames.

re-integration while others are not when a message is transmitted. An attacker can exploit this to push selected node types into bus re-integration.

**(V3.2) 0 in INT3 or Short Pulses Between Messages:** Similar to V2.1 and V2.3, injecting 0s in INT3 or short pulses between messages forces nodes that do not support the frame to stay in bus re-integration when a new message arrives.

**Exploit 1: Node Type Identification.** An LL attacker can exploit V3.2 to test whether a node has a specific type or another type supporting more frame formats. When the attacker sees a message transmitted by the node, he wins arbitration with a higher priority ID and sends a potentially unsupported frame (e.g., XL frame for FD enabled nodes) and a 0 in INT3. He then checks if the node retransmits the message after INT3. The absence of a retransmission indicates that the node has the expected type, does not support the frame, and has entered bus re-integration. Otherwise, it supports the frame and has another type. Interference from other traffic during the test is handled as described in App. C–V3.

**Impact.** Node type identification reveals proprietary network configuration information remotely. It is the prerequisite for targeted attacks on specific node types as described below.

**Exploit 2: Inconsistent Frames.** LL attackers can exploit V3.1 to launch a new attack that lets nodes of the same type as a message’s transmitter receive the genuine message while nodes of other types receive a fake one. To illustrate its steps, Fig. 9 shows an example where an attacker lets all XL enabled nodes receive a genuine message while other nodes receive a fake one when an XL enabled node sends a message:

**Step ①:** The attacker wins arbitration over the target message and pushes all nodes into bus re-integration (e.g., by sending a partial XL frame with 1 in resXL).

**Step ②:** When nodes of a type different from the transmitter’s (i.e., FD enabled) resume operation and nodes of the same type are in bus re-integration, the attacker sends a fake message with the same ID so different-type nodes receive it.

**Step ③:** The attacker sends a 0 in the fake message’s INT3 and pushes different-type nodes into bus re-integration (e.g., by sending a partial XL frame with ID 0 ending with XLF).

**Step ④:** The transmitter and nodes of the same type (i.e., XL enabled) exit bus re-integration earlier. It retransmits the genuine message and all nodes of its type receive it.

**Impact.** This is a targeted spoofing attack that deceives selected nodes. Victims develop an inconsistent view of the system state relative to other nodes, leading to conflicting

control decisions. Moreover, it is useful for evading payload-inspecting IDSs [21,36] by letting the IDS receive the genuine message while nodes of other types receive a fake one.

**Target Node Shutdown, Advanced Spoofing, and Data Consistency Violation.** LL attackers can exploit V3.2 to shut down target nodes by keeping them in bus re-integration. However, the victim may not be a single node, but all nodes of a specific type. This is still useful in many scenarios. For example, if a network has only one node of a particular type. Similar to message replacement (V2), LL attackers can exploit V3.1 to spoof a message after pushing its transmitter into bus re-integration. Although it cannot let the transmitter believe it has transmitted the message successfully, it is still useful for evading defenses that let transmitters monitor for their own IDs. LL attackers can also exploit V3.1 to violate data consistency among different types of nodes in various ways (e.g., drop messages only for specific types or inject messages only specific types can receive).

**V4. Premature Re-Integration Exit.** In a mixed network, we discover that when a message is interrupted with errors and all nodes in normal operation are error-passive (Sec. 2), nodes that do not support the frame exit bus re-integration before the passive error frame ends, violating  $\mathcal{P}_2$ . This is because passive error frames contain 14 1s, more than what is needed to exit bus re-integration (11 1s). If a node transmits a message right after exiting bus re-integration, it collides with error-passive nodes’ error delimiter and error-passive nodes detect another error, violating  $\mathcal{P}_9$ . Due to the error, error-passive nodes cannot receive the message. They transmit another passive error frame which is invisible to other nodes. Thus, other nodes can still receive the message, violating  $\mathcal{P}_4$ .

**Exploit: Data Consistency Violation.** After pushing XL enabled nodes with protocol exception disabled (Sec. 2) into error-passive, LL attackers can exploit V4 to violate data consistency between them and other nodes in many ways. Fig. 10 shows an example. When a node with protocol exception enabled transmits a message, the attacker can let nodes with protocol exception disabled receive a fake message, but other nodes receive the genuine one:

**Step ①:** The attacker wins arbitration with a higher-priority ID after seeing the target message’s ID.

**Step ②:** The attacker sends a partial XL frame ending with a 1 in resXL to push nodes with protocol exception enabled into bus re-integration. Nodes with protocol exception disabled transmit a passive error frame.



a single message and do not need to inject repeated collisions to keep nodes in the error-passive state. Furthermore, the time to map a CAN CC ECU is proportional to the number of IDs it transmits [28]. However, the time to map a CAN XL ECU depends only on its longest message period. The attack is also less detectable, as an IDS monitoring the higher layer only observes a single error instead of repeated collisions.

**Node Type Identification.** A unique attack for mixed networks, it is possible by exploiting V3.2.

**Advanced Spoofing.** For CAN CC and CAN XL with error signalling enabled, attackers can exploit  $V_a$  to launch advanced spoofing attacks by pushing a node to bus-off/error-passive and spoofing/hijacking its messages. However, pushing nodes into different error states takes time and causes multiple errors, making the attacks detectable. Contrastingly, when error signalling is disabled, attackers can overwrite (V1.2) or replace (V2) messages. They are faster and less detectable as they do not require additional setup steps and cause no or only one error. In mixed networks, attackers can spoof messages after pushing the transmitter into bus re-integration (V3.2) or launch inconsistent frame attacks (V3.1), without requiring additional setup time or generating errors.

**Data Consistency Violation.** By violating  $\mathcal{P}_4$  only,  $V_b$  allows one form of data consistency violation in CAN CC: letting receivers validate a message but the transmitter does not. This can be achieved for CAN XL when error signalling is enabled as it inherits  $V_b$ . In contrast, many more forms of consistency violations are possible among arbitrary nodes when error signalling is disabled by exploiting V2, or among nodes of different types in mixed networks by exploiting V3.1 and V4.

**Full Network DoS.**  $V_c$  allows LL attackers to launch full network DoS attacks on CAN CC and CAN XL with error signalling enabled. With error signalling disabled, the same can be achieved by exploiting V5. In contrast, it only requires HL control to introduce controlled, finite bus delays (App. C–V5).

**Target Node Shutdown.**  $V_a$  allows HL attackers to shut down target nodes by pushing them into the bus-off state for CAN CC and CAN XL when error signalling is enabled. When error signalling is disabled or in mixed networks, LL attackers can shut down nodes by exploiting V2 or V3. This is often faster than CAN CC. The fastest technique for CAN CC needs 255 nominal bit durations [27]. V2 or V3 allows attackers to disable a node in one message duration plus 16 bits (SOF-IDE in a message with ID 0, Fig. 7b), which is often faster. It may be less persistent as pulses must be injected in every message’s EOF and whenever the bus is idle to keep nodes in bus re-integration. However, in Sec. 7.2.1, we achieve a high suppression rate experimentally, making the attack highly effective in practice. Finally, this attack is less detectable as it causes fewer errors: one error when error signalling is disabled or no errors in mixed networks. Contrastingly, its CAN CC counterpart incurs at least 32 errors.

**Interception.** This attack is only possible for CAN XL networks with error signalling disabled by exploiting V1.1.

**Sniffing, Injection, Flooding, and Error Injection.** For completeness, we note that HL attackers can launch these attacks for both CAN CC and CAN XL using the same techniques (by sniffing/injecting messages directly or errors with simultaneous transmission) with the same effectiveness.

**Conclusion.** In both error signalling configurations, CAN XL’s MAC sub-layer vulnerabilities enable all CAN CC attack categories. Moreover, when error signalling is disabled or in mixed networks, its new vulnerabilities enable new attacks and make several attacks faster and stealthier. Although some new or improved attacks require LL control, this is achievable remotely on many standard ECUs and has been demonstrated in practice (Sec. 4.2), and does not make them less feasible. Thus, we argue that *from a strict MAC sub-layer standpoint, changes by CAN XL worsen security.*

## 7 Evaluation

We perform two evaluations. First, we validate the presence of the discovered vulnerabilities in commercial CAN XL protocol controllers. Next, we show the vulnerabilities’ exploitability by using them to construct and evaluate two multi-stage attacks from reconnaissance to execution.

### 7.1 Vulnerability Presence Validation

**Setup.** We use two commercial CAN XL interfaces: PEAK-System PCAN-USB XL and Vector VN1641. We set up a network with each device and connect an STM32 microcontroller as an attacker. Our formal analysis identifies one or more counterexamples per vulnerability, each triggered by different attacker actions. For each counterexample, we set up nodes to transmit/receive messages and the attacker to inject signals as specified, and verify whether the observed node behaviors match the counterexample.

**Standard Vulnerabilities.** Tbl. 4 shows their presence.

Table 4: Standard vulnerabilities in commercial devices.

| Device  | V1 | V2 | V3 | V4 | V5 | V6 | V7 |
|---------|----|----|----|----|----|----|----|
| PCAN XL | ✓  | ✓  | ✓  | ✓  | ✓  | ✓  | ✓  |
| VN1641  | ✓* | ✓* | ✓  | ✓  | ✓  | ✓  | ✓  |

*PCAN XL:* All vulnerabilities can be triggered by the same attacker actions as our formal analysis discovers.

*VN1641:* All vulnerabilities are present but V1 and V2 require different attacker actions to trigger because its implementation deviates from the standard as explained below.

**Deviations From the Standard.** Apart from standard-level vulnerabilities, we also notice behavioral deviations between the devices and our standard model in some error corner cases of the protocol. These deviations are intentional by vendors or due to implementation bugs.







off-the-shelf CAN controller implementations. Since implementations must comply with the standard to ensure interoperability, standard-level vulnerabilities will be present in all compliant implementations unless implementations deviate from the standard. Such deviations, which are expected with a new protocol, are more likely to introduce additional vulnerabilities than to avoid standard ones. Our formal model is proven effective at exposing such deviations, as it accurately captures protocol corner cases and helps discover bugs that vendors may not discover otherwise (Sec. 7.1). Formally analyzing these differences allows us to understand their security impacts, as demonstrated by the VN1641 example.

**Analysis Approach Generalizability.** While certain parts of our analysis approach are designed specifically for CAN XL, similar sequential decomposition techniques could apply to other protocols with sequential fields/phases (e.g., automotive Ethernet). Human expertise is required to identify appropriate boundaries to decompose the protocol into separate FSMs. Then, if Algorithm 1 is modified to accept the protocol's initial states and field/phase order as inputs, it can be reused to analyze other protocols given their decomposed FSMs.

**Other CAN XL Layers.** Since our analysis focuses on the MAC sub-layer of CAN XL compared to CAN CC, we must make assumptions about the operation of other layers and do not attempt to discover attacks exploiting other layers. For the physical coding sub-layer (PCS), which manages bit sampling and synchronization, we assume nodes are always synchronized to the bit stream on the bus (Sec. 4.3). As CAN XL modifies the PCS minimally, prior attacks exploiting CAN CC's PCS [24, 32] carry over to CAN XL. For the physical layer, we assume it operates in SIC mode, as it is common to all versions, mandated for critical frame fields, and is the only mode currently available in off-the-shelf transceivers. This makes it the natural baseline for meaningful comparison and evaluation. However, investigating the new fast mode, currently only available in non-commercial transceivers, offers an interesting future research direction. As it introduces new voltage symbols and unpredictable collision behavior, it will likely introduce vulnerabilities previously unseen in CAN CC. Regarding its impact on the vulnerabilities discovered in this paper, we expect it to only lessen the impact of VI.2, as VI.2 requires deterministically overwriting level\_1 with level\_0. All other vulnerabilities remain intact, as they either relate solely to SIC mode or occur in frame fields that are required to always use SIC mode (e.g., INT).

**Physical Access Attackers.** Several real-world attacks have been launched by attackers who obtain physical access to the CAN bus and attach specialized equipment. These attackers can break any protocol rules and modify signals on the bus arbitrarily. To incorporate them in the analysis, we can modify the model (Fig. 4) to directly connect the attacker's outputs to nodes. We choose not to do so because CAN security literature still considers physical access a very strong privilege, and the attacker's strong capabilities make it unlikely any security

properties can be guaranteed.

## 10 Conclusion

We present the first standard-level security analysis of CAN XL, the latest CAN generation. We focus specifically on its MAC sub-layer, a historically vulnerable part of the standard. We build the first formal model of the MAC sub-layer across all CAN generations (CC/FD/XL) that precisely captures its bit-oriented operation and thoroughly covers cross-generation interactions. We conduct its first formal analysis and address tractability challenges by designing an analysis workflow that decomposes the protocol into analyzable portions based on its field-oriented structure. Our analysis shows that CAN XL inherits all known CAN CC vulnerabilities and introduces seven new standard-level vulnerabilities. These weaknesses enable all classical CAN CC attacks as well as new attacks, several faster, more flexible, and less detectable than their CAN CC counterparts. We argue that despite its extensive modifications and attempts to improve security, from a strict MAC sub-layer standpoint, CAN XL's security has worsened. We validate all vulnerabilities on off-the-shelf CAN XL controllers and demonstrate their exploitability through two multi-stage attacks. Finally, we propose mitigations, including formally verifying standard revisions to prevent several attacks.

## Acknowledgments

We thank the anonymous reviewers for their valuable feedback. This work was supported in part by the National Science Foundation (NSF) under grants CNS-2144645, CNS-2333488, and CNS-2231651, Hyundai America Technical Center (HATCI), and Georgia Department of Transportation. Any opinions, findings, and conclusions or recommendations expressed in this material are those of the authors and do not necessarily reflect the views of our sponsors.

## Ethical Considerations

**Stakeholder Identification.** A direct stakeholder of this research is the standardization body of CAN XL. Since CAN XL's main use case is in the automotive domain, another direct stakeholder is the automotive industry. Moreover, vehicle occupants and road users are indirect stakeholders.

**Impacts.** This research can benefit the CAN XL ecosystem by improving its security before widespread adoption. Our analysis helps the standardization body and automotive industry identify design and implementation weaknesses early and deploy more robust defenses in future products. Contrastingly, the research may reduce the effort required for adversaries to target CAN XL systems once they are deployed. For vehicle occupants and road users, successful exploitation could result in safety risks, operational disruption, or financial losses. For

the automotive industry, such exploits could impose incident-response costs or expose proprietary information.

**Mitigations.** In accordance with responsible disclosure practices, we have reported all discovered vulnerabilities in the CAN XL standard to the Robert Bosch Product Security Incident Response Team (PSIRT) since the CAN protocols are patented products by Bosch. PSIRT has acknowledged the receipt of our report and is investigating the issues internally. We document concrete countermeasures in the paper for the discovered vulnerabilities to further mitigate potential harm. All discovered bugs in commercial CAN XL controllers have been disclosed to and acknowledged by the vendors.

**Decision.** Since CAN XL has not yet been deployed in production systems, our findings will not cause immediate real-world harm. Instead, our goal is to proactively identify vulnerabilities to support the design of defensive measures before CAN XL's widespread adoption. Publishing our results at this stage gives relevant stakeholders time to assess exposure and implement countermeasures prior to large-scale deployment. Overall, we judge that such benefits outweigh potential risk.

## Open Science

We provide the source code of our CAN XL protocol models and analysis scripts. We also provide the implementation of the attacks evaluated in Sec. 7.2. Our artifacts can be accessed at <https://zenodo.org/records/20262062>.

## References

- [1] Future of CAN networking in VW's passenger cars. <https://tinyurl.com/54xhzd7w>, 2021.
- [2] *Road vehicles — Controller area network (CAN). Part 1: Data link layer and physical coding sublayer, ISO 11898-1:2024*. 2024.
- [3] Apple carplay vulnerability exploited to gain root access. <https://tinyurl.com/2r5rmybs>, 2025.
- [4] Application of CAN XL communication technology in automotive millimeter-wave radar (2). <https://tinyurl.com/2rpp3e25>, 2025.
- [5] Application of CAN XL communication technology in vehicle-mounted millimeter-wave radar (1). <https://tinyurl.com/5d8h83ch>, 2025.
- [6] CAN XL in the automotive industry: History, technical capabilities, and market outlook. <https://tinyurl.com/84f66bnc>, 2025.
- [7] *CiA 613-2: CAN XL add-on services — Part 2: CANsec data plane*. 2025.
- [8] Rohit Bhatia, Vireshwar Kumar, Khaled Serag, Z Berkay Celik, Mathias Payer, and Dongyan Xu. Evading voltage-based intrusion detection on automotive CAN. In *Network and Distributed System Security Symposium (NDSS)*, 2021.
- [9] Tim Brom. Cant. <https://github.com/bitbane/CANT>, 2018.
- [10] Zhiqiang Cai, Aohui Wang, Wenkai Zhang, Michael Gruffke, and Hendrick Schweppe. 0-days & mitigations: Roadways to exploit and secure connected bmw cars. *Black Hat USA*, 2019.
- [11] Roberto Cavada, Alessandro Cimatti, Michele Dorigatti, Alberto Griggio, Alessandro Mariotti, Andrea Micheli, Sergio Mover, Marco Roveri, and Stefano Tonetta. The nuXmv symbolic model checker. In *Computer Aided Verification*, 2014.
- [12] Stephen Checkoway, Damon McCoy, Brian Kantor, Danny Anderson, Hovav Shacham, Stefan Savage, Karl Koscher, Alexei Czeskis, Franziska Roesner, and Tadayoshi Kohno. Comprehensive experimental analyses of automotive attack surfaces. In *USENIX Security Symposium*, 2011.
- [13] Kyong-Tak Cho and Kang G. Shin. Error handling of in-vehicle networks makes them vulnerable. In *ACM SIGSAC Conference on Computer and Communications Security (CCS)*, 2016.
- [14] Tsvika Dagan and Avishai Wool. Parrot, a software-only anti-spoofing defense system for the CAN bus. *ESCAR EUROPE*, 2016.
- [15] Alvise de Faveri Tron, Stefano Longari, Michele Carminati, Mario Polino, and Stefano Zanero. Canflit: Exploiting peripheral conflicts for data-link layer attacks on automotive networks. In *ACM SIGSAC Conference on Computer and Communications Security (CCS)*, 2022.
- [16] Magnus-Maria Hell. The physical layer in the CAN XL world. In *The international CAN Conference (iCC)*, 2020.
- [17] Pengfei Jing, Zhiqiang Cai, Yingjie Cao, Le Yu, Yuefeng Du, Wenkai Zhang, Chenxiong Qian, Xiapu Luo, Sen Nie, and Shi Wu. Revisiting automotive attack surfaces: A practitioners' perspective. In *IEEE Symposium on Security and Privacy (SP)*, 2024.
- [18] Karl Koscher, Alexei Czeskis, Franziska Roesner, Shwetak Patel, Tadayoshi Kohno, Stephen Checkoway, Damon McCoy, Brian Kantor, Danny Anderson, Hovav Shacham, et al. Experimental security analysis of a modern automobile. In *IEEE Symposium on Security and Privacy (SP)*, 2010.

- [19] Sekar Kulandaivel, Shalabh Jain, Jorge Guajardo, and Vyas Sekar. CANNON: reliable and stealthy remote shutdown attacks via unaltered automotive microcontrollers. In *IEEE Symposium on Security and Privacy (SP)*, 2021.
- [20] Stefano Longari, Matteo Penco, Michele Carminati, and Stefano Zanero. Copycan: An error-handling protocol based intrusion detection system for controller area network. In *ACM Workshop on Cyber-Physical Systems Security & Privacy*, 2019.
- [21] Stefano Longari, Carlo Alberto Pozzoli, Alessandro Nichelini, Michele Carminati, and Stefano Zanero. Candito: improving payload-based detection of attacks on controller area networks. In *International Symposium on Cyber Security, Cryptology, and Machine Learning*, 2023.
- [22] Tsutomu Matsumoto, Masato Hata, Masato Tanabe, Katsunari Yoshioka, and Kazuomi Oishi. A method of preventing unauthorized data transmission in controller area network. In *IEEE 75th Vehicular Technology Conference (VTC Spring)*, 2012.
- [23] Charlie Miller and Chris Valasek. Remote exploitation of an unaltered passenger vehicle. *Black Hat USA*, 2015.
- [24] Pal-Stefan Murvay and Bogdan Groza. Dos attacks on controller area networks by fault injections from the software layer. In *International Conference on Availability, Reliability and Security (ARES)*, 2017.
- [25] Sen Nie, Ling Liu, Yuefeng Du, and Wenkai Zhang. Over-the-air: How we remotely compromised the gateway, bcm, and autopilot ecus of tesla cars. *Black Hat USA*, 2018.
- [26] Khaled Serag. Proactive Vulnerability Identification and Defense Construction – The Case For CAN. 2023.
- [27] Khaled Serag, Rohit Bhatia, Akram Faqih, Muslim Ozgur Ozmen, Vireshwar Kumar, Z. Berkay Celik, and Dongyan Xu. ZBCAN: A zero-byte CAN defense system. In *USENIX Security Symposium*, 2023.
- [28] Khaled Serag, Rohit Bhatia, Vireshwar Kumar, Z. Berkay Celik, and Dongyan Xu. Exposing new vulnerabilities of error handling mechanism in CAN. In *USENIX Security Symposium*, 2021.
- [29] Khaled Serag, Vireshwar Kumar, Z Berkay Celik, Rohit Bhatia, Mathias Payer, and Dongyan Xu. Attacks on can error handling mechanism. In *Workshop on Automotive and Autonomous Vehicle Security (AutoSec)*, 2022.
- [30] Khaled Serag, Zhaozhou Tang, Sungwoo Kim, Vireshwar Kumar, Saman Zonouz, Raheem Beyah, Dongyan Xu, Z Berkay Celik, et al. SoK: Kicking CAN down the road. Systematizing CAN security knowledge. *arXiv preprint arXiv:2510.02960*, 2025.
- [31] Hyun Min Song, Ha Rang Kim, and Huy Kang Kim. Intrusion detection system based on the analysis of time intervals of CAN messages for in-vehicle network. In *International Conference on Information Networking (ICOIN)*, 2016.
- [32] Zhaozhou Tang, Khaled Serag, Saman Zonouz, Z. Berkay Celik, Dongyan Xu, and Raheem Beyah. ERACAN: Defending Against an Emerging CAN Threat Model. In *ACM SIGSAC Conference on Computer and Communications Security (CCS)*, 2024.
- [33] Ken Tindell. CAN bus security: Attacks on CAN bus and their mitigations. Technical report, Canis Automotive Labs, February 2020.
- [34] Ken Tindell. Canhack. <https://github.com/kentindell/canhack>, 2020.
- [35] Ken Tindell. Three new can protocol hacks. <https://kentindell.github.io/2020/01/20/new-can-hacks/>, 2020.
- [36] Armin Wasicek, Mert D Pesé, André Weimerskirch, Yelizaveta Burakova, and Karan Singh. Context-aware intrusion detection in automotive control systems. *ES-CAR USA*, 2017.
- [37] Fuyu Yang. A bus off case of can error passive transmitter. *EDN Technical paper*, 2009.
- [38] Clinton Young, Habeeb Olufowobi, Gedare Bloom, and Joseph Zambreno. Automotive intrusion detection based on constant can message frequencies across vehicle driving modes. In *ACM Workshop on Automotive Cybersecurity*, 2019.

## A Formal Property Specifications

Tbl. 11 shows the properties in Sec. 4.1 expressed in CTL using the syntax of nuXmv. We explain them in detail below:  $\mathcal{P}_1$ : This states that at all times, if the node  $n1$  is in state `bus_reintegration`, then from this point,  $n1$  can always finally leave `bus_reintegration`. The same property, and other properties below, are defined symmetrically for  $n2$ , the second node we model (Fig. 4). We omit them for brevity.  $\mathcal{P}_2$ : At all times, if  $n1$  is in `bus_reintegration` but  $n2$  is not in `int3` (the last field before a new frame), idle, or also in `bus_reintegration`,  $n1$  always stays in `bus_reintegration` in the next cycle. This ensures if  $n1$  enters `bus_reintegration` during a frame, it stays in this state until the frame completes.





Figure 14: Message interception by HL attackers.

corner case with the same impacts.

**V7. Error Regeneration.** When error signalling is enabled, we discover a violation of  $\mathcal{P}_9$  where a node encounters further errors when sending an error frame. It applies to XL frames that have 9 bits in FCRC after the last fixed stuff bit and an XL data bit rate equal to twice the nominal bit rate. As shown in Fig. 13b, an attacker injects an error in the first of these 9 bits. Depending on its error state, the transmitter sends an active or a passive error frame. The error flag starts at the next bit and ends at the bit corresponding to FCP0, followed by the recessive error delimiter. This violates the fixed FCP format and receivers detect an error. This is treated as a CRC error and not signalled immediately after FCP. Instead, receivers signal it with error frames after AL1. They collide with the transmitter's error delimiter, causing it to detect another error. If the transmitter is error-passive, back-to-back transmissions by receivers after the error frame further increment its error counter, possibly pushing it into bus-off.

**Impact.** Attackers can exploit V7 to rapidly set a victim's error counter and achieve many purposes such as quick targeted DoS by pushing the victim to bus-off. It is exploitable only by LL attackers, not HL attackers. This is because to cause an error only in FCRC, the attacker needs to transmit a message with the same data but a different FCRC than the victim's, which is impossible with a standard CAN controller. Similar goals can be achieved by exploiting passive error regeneration ( $V_a$ ), but  $V_a$  only applies to error-passive victims while V7 applies regardless of victims' error states. Moreover, V7 undermines the robustness of CAN XL communication, introducing a new corner case where a single genuine error leads to rapid error counter increments and possible transmitter bus-off.

## C Additional Exploit Descriptions

**V1. HL Message Interception.** Exploiting  $V1.1$ , HL attackers can intercept messages using simultaneous transmission. As shown in Fig. 14, the attacker simultaneously transmits a message that is longer than a transmitter's message by 2 nominal bit durations. With disabled error signalling, both the attacker and the transmitter can transmit their messages until completion. Collisions between two different messages let receivers detect errors. When the transmitter is expecting the ACK, the attacker is transmitting the dominant AL1 bit, which the transmitter treats as a valid ACK.

**V3. Node Type Identification.** During node type identification, other nodes may also transmit and win arbitration over the target node's retransmission. Thus, for each consecutive



Figure 15: Full network DoS by HL attackers.

message, the attacker records its ID and sends a 0 in INT3 until the bus becomes idle. If the target node's retransmission is observed, the node has the expected type. Otherwise, it has another type supporting more frame formats. If XL frames are observed when testing if a node is FD tolerant, the attacker cannot conclude that the node is FD tolerant if he does not see its retransmission. This may indicate that the node is FD tolerant and enters bus re-integration due to the attacker's FD frame, or is FD enabled and enters bus re-integration due to another node's XL frame. The test needs to be repeated until it is not interfered with by XL frames. To reduce such interference, the attacker should pick a high-priority ID that is not typically transmitted at a similar time as other nodes' messages to test. Exploiting V3.2, an attacker can also check if an XL enabled node enables or disables protocol exception (Sec. 2). Similarly, the attacker wins arbitration over the node's transmission and sends a partial XL frame with 1 in resXL, followed by an active error frame, and a 0 in INT3.

**V5. HL Full Network DoS.** Exploiting V5, HL attackers can delay bus traffic for finite durations by simultaneously transmitting a longer message with any node. As shown in Fig. 15, the collision causes the transmitter to detect errors in the EOF and receivers earlier, while the attacker continues transmission. Nodes stay in bus re-integration during the attacker's message (as verified by  $\mathcal{P}_2$ ). The delay can be arbitrary up to the maximum length of a CAN XL message.