

UNIVERSIDAD MIGUEL HERNÁNDEZ DE ELCHE

ESCUELA POLITÉCNICA SUPERIOR DE ELCHE

GRADO EN INGENIERÍA DE TECNOLOGÍAS DE
TELECOMUNICACIÓN



"MODELADO ANALÍTICO DE
COMUNICACIONES 5G NR V2X PARA
VEHÍCULO AUTÓNOMO
CONECTADO"

TRABAJO FIN DE GRADO

Junio - 2026

AUTOR: Sergio Alacid Riquelme

DIRECTOR/ES: Miguel Sepulcre Ribes
Luca Lusvarghi

RESUMEN

El desarrollo del vehículo conectado y autónomo requiere que los vehículos no dependan únicamente de sus propios sensores, sino que también sean capaces de intercambiar información con otros vehículos, con la infraestructura, con los peatones y con la red. En este contexto, las comunicaciones V2X permiten ampliar la percepción del entorno y habilitar aplicaciones orientadas a mejorar la seguridad vial, como la coordinación de maniobras, la percepción cooperativa o la gestión inteligente del tráfico. Para que estas aplicaciones sean viables, las comunicaciones vehiculares deben ofrecer baja latencia, alta fiabilidad y capacidad suficiente en escenarios densos y dinámicos.

Dentro de esta evolución, 5G NR V2X Modo 2 permite que los vehículos seleccionen de forma autónoma los recursos radio empleados en sus transmisiones directas. Sin embargo, esta gestión distribuida introduce un problema fundamental: el rendimiento del sistema depende de cómo se seleccionan, reservan, mantienen, descartan o reseleccionan dichos recursos. Una gestión ineficiente puede provocar colisiones, pérdidas de paquetes, incremento de la latencia o degradación de la fiabilidad. Además, el problema se vuelve más complejo cuando se consideran modelos de tráfico aperiódico y tamaños de paquete variables, ya que pueden aparecer reservas no utilizadas, paquetes que no caben en los recursos reservados o transmisiones que no pueden esperar a la siguiente reserva sin superar el límite máximo de latencia permitido.

El objetivo de este Trabajo Fin de Grado es modelar de forma analítica la gestión de recursos en comunicaciones 5G NR V2X Modo 2. Para ello, se propone una formulación probabilística basada en una cadena de Markov que representa el proceso de reserva, uso, descarte y reelección de recursos. A partir de este modelo se obtienen magnitudes como la probabilidad de reelección, la probabilidad de reservar un determinado número de subcanales y la probabilidad de que una reserva sea descartada. Estas probabilidades se emplean posteriormente para caracterizar métricas de rendimiento adicionales, concretamente la distribución de la latencia y el número de paquetes perdidos consecutivamente entre dos recepciones correctas.

La validación del modelo se realiza mediante una doble comparación. En primer lugar, se emplea un simulador propio desarrollado en MATLAB, diseñado para reproducir de forma controlada la lógica temporal considerada en el modelo analítico. En segundo lugar, los resultados se contrastan con los obtenidos mediante un simulador de comunicaciones vehiculares desarrollado en la UMH, basado en OMNeT++, INET y

Veins. De este modo, el trabajo proporciona una herramienta analítica de bajo coste computacional que permite estudiar de forma rápida e interpretable la gestión de recursos en 5G NR V2X Modo 2 y analizar su impacto sobre métricas relevantes para las comunicaciones vehiculares.



ABSTRACT

The development of connected and autonomous vehicles requires vehicles not to rely solely on the information obtained from their own sensors, but also to exchange information with other vehicles, infrastructure, pedestrians and the communication network. In this context, V2X communications extend the perception of the environment and enable applications aimed at improving road safety, such as maneuver coordination, cooperative perception and intelligent traffic management. To make these applications feasible, vehicular communications must provide low latency, high reliability and sufficient capacity in dense and dynamic scenarios.

In this context, this work focuses on 5G NR V2X Mode 2, where vehicles autonomously select the radio resources used for their direct transmissions. However, this distributed operation introduces a fundamental problem: the system performance depends on how these resources are selected, reserved, maintained, discarded or reselected. Inefficient resource management may lead to collisions, packet losses, increased latency or degraded communication reliability. Moreover, the problem becomes more complex when aperiodic traffic models and variable packet sizes are considered, since unused reservations, packets that do not fit within the reserved resources, or transmissions that cannot wait for the next reservation without exceeding the maximum latency limit may arise.

The objective of this Bachelor's Thesis is to analytically model resource management in 5G NR V2X Mode 2 communications. To this end, a probabilistic formulation based on a Markov chain is proposed to represent the process of resource reservation, usage, discarding and reselection. From this model, relevant metrics are obtained, such as the resource reselection probability, the probability of reserving a given number of subchannels and the probability that a reservation is discarded. These probabilities are then used to characterize additional performance metrics, specifically the latency distribution and the number of consecutive packet losses between two successful receptions.

The model is validated through a twofold comparison. First, a MATLAB simulator is developed to reproduce, in a controlled manner, the temporal logic considered in the analytical model. Second, the results are compared with those obtained from a vehicular communications simulator developed at UMH and based on OMNeT++, INET and Veins.

In this way, this work provides a low-computational-cost analytical tool that enables the fast and interpretable study of resource management in 5G NR V2X Mode 2 and its impact on relevant metrics for vehicular communications.



ÍNDICE

1. INTRODUCCIÓN.....	13
2. ESTADO DEL ARTE	17
2.1. 5G NR V2X.....	17
2.1.1. Introducción.....	17
2.1.2. Capa Física	18
2.1.3. Capa de Enlace	23
2.2. Modelos analíticos existentes	29
3. MODELADO ANALÍTICO	35
3.1. Modelo de gestión de recursos	35
3.1.1. Antecedentes y planteamiento del modelo	35
3.1.2. Definición de la cadena de Markov y probabilidades asociadas	38
3.1.3. Cálculo de probabilidades de descarte y reelección por latencia.....	45
3.2. Modelo de distribución de la latencia.....	53
3.2.1. Casos con distribución uniforme de latencia.....	55
3.2.2. Distribución de latencia en modelos aperiódicos con $RRI=2PDB$	59
3.3. Modelo de pérdidas entre recepciones correctas	63
3.3.1. Planteamiento del modelo	63
3.3.2. Cadena de Markov y distribución de pérdidas consecutivas.....	65
4. RESULTADOS Y VALIDACIÓN	71
4.1. Herramientas y modelos de tráfico considerados	71
4.1.1. Herramientas de validación	71
4.1.2. Modelos de tráfico considerados	73
4.2. Validación de los modelos analíticos	75
4.2.1. Validación del modelo de gestión de recursos	75
4.2.2. Validación del modelo de distribución de la latencia.....	81
4.2.3. Validación del modelo de pérdidas entre recepciones correctas	85

5. CONCLUSIONES Y LÍNEAS FUTURAS	94
6. REFERENCIAS	97



ÍNDICE DE FIGURAS

Figura 1. Estructura tiempo-frecuencia de los recursos <i>sidelink</i> en 5G NR V2X.	20
Figura 2. Ventanas de sensado y selección en NR V2X Modo 2 para el caso $T2 = PDB$	26
Figura 3. Esquema de transmisiones.	36
Figura 4. Esquema de reselecciones por latencia.	37
Figura 5. Esquema de reservas descartadas.	37
Figura 6. Cadena de Markov para el Modelo Analítico de gestión de recursos.	39
Figura 7. Cadena de Markov para el Modelo Analítico de pérdidas entre recepciones correctas.	68
Figura 8. Probabilidades de reselección y descarte para todos los modelos de tráfico con $\mu = 1$ y MCS=13.	76
Figura 9. Distribución del número de subcanales reservados para todos los modelos de tráfico con $\mu = 1$ y MCS=13.	77
Figura 10. Error absoluto Modelo Analítico frente a Simulador MATLAB.	79
Figura 11. Error absoluto Modelo Analítico frente a Simulador Referencia.	79
Figura 12. Distribución de la latencia para todos los modelos de tráfico con $\mu = 1$ y MCS = 13.	82
Figura 13. Error absoluto medio (MAE) y error absoluto máximo del modelo analítico de latencia respecto a ambos simuladores para las configuraciones evaluadas.	84
Figura 14. Distribución de pérdidas consecutivas para los modelos de tráfico periódicos.	86
Figura 15. Distribución de pérdidas consecutivas para los modelos de tráfico AP1 y AP2.	87
Figura 16. Distribución de pérdidas consecutivas para los modelos de tráfico AP3 y AP4.	88
Figura 17. Probabilidad de perder N_{loss} paquetes consecutivos en función de la distancia para el modelo P1.	89
Figura 18. Probabilidad de perder N_{loss} paquetes consecutivos en función de la distancia para el modelo P2.	89
Figura 19. Probabilidad de perder N_{loss} paquetes consecutivos en función de la distancia para el modelo P3.	89

Figura 20. Probabilidad de perder <i>N</i> loss paquetes consecutivos en función de la distancia para el modelo AP1.	90
Figura 21. Probabilidad de perder <i>N</i> loss paquetes consecutivos en función de la distancia para el modelo AP2.	90
Figura 22. Probabilidad de perder <i>N</i> loss paquetes consecutivos en función de la distancia para el modelo AP3.	90
Figura 23. Probabilidad de perder <i>N</i> loss paquetes consecutivos en función de la distancia para el modelo AP4.	91
Figura 24. Error absoluto medio del Modelo Analítico de Pérdidas consecutivas respecto al simulador de referencia para los modelos y configuraciones evaluadas.....	92



ÍNDICE DE TABLAS

Tabla 1. Modelos de tráfico para validación del modelo analítico.....	73
Tabla 2. Configuración inicial de validación.....	76



GLOSARIO DE TÉRMINOS

3GPP: 3rd Generation Partnership Project

5G: Fifth Generation

BWP: Bandwidth Part

CAS: Cooperative Awareness Service

CPS: Collective Perception Service

CBR: Channel Busy Ratio

C-V2X: Cellular Vehicle-to-Everything

DMRS: Demodulation Reference Signals

gNB: next generation Node B

HARQ: Hybrid Automatic Repeat Request

IEEE: Institute of Electrical and Electronics Engineers

INET: Framework de simulación de redes para OMNeT++

ITS-G5: Intelligent Transport Systems operating in the 5 GHz band

LTE: Long Term Evolution

MAC: Medium Access Control

MAE: Mean Absolute Error

MATLAB: MATrix LABoratory

MCS: Modulation and Coding Scheme

NR: New Radio

OFDM: Orthogonal Frequency Division Multiplexing

OMNeT++: Objective Modular Network Testbed in C++

PDB: Packet Delay Budget

PDF: Probability Density Function

PDR: Packet Delivery Ratio

PIR: Packet Inter-Reception

PRB: Physical Resource Block

PRR: Packet Reception Ratio

PSCCH: Physical Sidelink Control Channel

PSFCH: Physical Sidelink Feedback Channel

PSSCH: Physical Sidelink Shared Channel

QAM: Quadrature Amplitude Modulation

RB: Resource Block

RC: Reselection Counter
RRI: Resource Reservation Interval
RSRP: Reference Signal Received Power
RSSI: Received Signal Strength Indicator
SCI: Sidelink Control Information
SCS: Subcarrier Spacing
SPS: Semi-Persistent Scheduling
TB: Transport Block
V2V: Vehicle-to-Vehicle
V2X: Vehicle-to-Everything
Veins: Vehicles in Network Simulation



1. INTRODUCCIÓN

El transporte por carretera constituye uno de los pilares fundamentales de la movilidad actual, pero también plantea desafíos importantes relacionados con la congestión, la eficiencia y, especialmente, la seguridad vial. Una parte significativa de los accidentes de tráfico está asociada al factor humano, incluyendo distracciones, errores de percepción o decisiones incorrectas en situaciones complejas [1]. Esta realidad ha impulsado el desarrollo de sistemas avanzados de asistencia a la conducción y de soluciones orientadas al vehículo conectado y autónomo, donde la percepción del entorno, la toma de decisiones y la coordinación con otros usuarios de la vía adquieren un papel fundamental. El desarrollo del vehículo autónomo se ha apoyado tradicionalmente en sensores instalados en el propio vehículo, como cámaras, radares o sistemas LiDAR, que permiten al vehículo construir una representación de su entorno inmediato. Sin embargo, estos sistemas presentan limitaciones asociadas a condiciones meteorológicas adversas, obstáculos físicos, oclusiones, alcance limitado o errores de interpretación del entorno [2]. Además, un vehículo basado únicamente en percepción local no siempre puede anticipar situaciones que ocurren fuera de su campo de visión.

Para superar estas restricciones, resulta necesario que los vehículos intercambien información de forma explícita con otros vehículos, con la infraestructura vial, con los peatones y con la red de comunicaciones. De esta necesidad surge el paradigma de las comunicaciones V2X, del inglés “*Vehicle-to-Everything*” [3]. Las comunicaciones V2X permiten habilitar aplicaciones como los avisos cooperativos de emergencia, la asistencia en intersecciones, la coordinación de maniobras, la percepción cooperativa, el *platooning* o la distribución de información sobre el estado del tráfico [4]. En estos escenarios, los vehículos pueden compartir información como su posición, velocidad, dirección, estado de frenado o intención de maniobra.

Para que estas aplicaciones sean viables, las comunicaciones vehiculares deben cumplir requisitos exigentes de latencia, fiabilidad y capacidad, especialmente en escenarios densos y dinámicos. Un mensaje recibido demasiado tarde puede perder su utilidad, y una pérdida prolongada de información puede impedir que el sistema disponga de una visión actualizada del entorno.

Por este motivo, la evolución de las comunicaciones vehiculares ha estado ligada al desarrollo de distintas tecnologías de acceso radio. Inicialmente, una parte importante de los estudios y despliegues se centró en tecnologías como ITS-G5, basada en IEEE

802.11p, y posteriormente en LTE-V2X, introducida dentro del ecosistema celular. Estas tecnologías supusieron avances importantes para la comunicación directa entre vehículos, permitiendo el intercambio de información sin necesidad de que los datos transmitidos pasaran por la infraestructura de red, lo que se conoce como comunicación *sidelink*. Sin embargo, tanto ITS-G5 como LTE-V2X presentaban limitaciones en términos de flexibilidad, capacidad, latencia y fiabilidad [5].

En este contexto, 5G NR V2X representa una evolución de las comunicaciones vehiculares celulares, incorporando nuevas capacidades orientadas a soportar casos de uso más exigentes [6]. Dentro de esta tecnología, el presente trabajo se centra especialmente en el Modo 2, donde la asignación de recursos no está controlada directamente por la red, sino que cada vehículo selecciona de forma autónoma los recursos que emplea para sus transmisiones directas. Esta autonomía introduce un reto fundamental: la gestión distribuida de los recursos radio. En escenarios vehiculares densos, múltiples vehículos comparten un mismo conjunto de recursos tiempo-frecuencia para transmitir sus mensajes. Cada vehículo debe decidir qué recursos utilizar intentando evitar interferencias y colisiones con las transmisiones de otros usuarios. Esta decisión se realiza además en un entorno altamente dinámico, donde los vehículos entran y salen de la zona de comunicación y las condiciones del canal cambian con el tiempo.

En este modo de operación, el rendimiento del sistema depende de mecanismos como la selección inicial de recursos, el mantenimiento de reservas, la reelección de nuevos recursos y la capacidad de ajustar dichas reservas al tamaño y al instante de generación de los paquetes. Una asignación ineficiente puede provocar colisiones, pérdidas de paquetes, incremento de la latencia o degradación de la fiabilidad de la comunicación. Estos efectos son especialmente críticos en comunicaciones vehiculares, donde la información recibida debe mantenerse actualizada para que el vehículo pueda interpretar correctamente su entorno [4].

El análisis de estos mecanismos puede abordarse principalmente mediante simulación o mediante modelos analíticos. Los simuladores permiten representar con un alto nivel de detalle la movilidad de los vehículos, el canal radio, los protocolos de comunicación y los procedimientos de acceso al medio. Sin embargo, este nivel de detalle suele implicar un coste computacional elevado, especialmente cuando se desea estudiar un gran número de configuraciones, realizar barridos paramétricos o evaluar escenarios con alta densidad de usuarios. Por ello, aunque la simulación resulta fundamental para validar y estudiar casos

concretos, no siempre es la herramienta más eficiente para analizar de forma rápida la influencia de los distintos parámetros del sistema.

Frente a este enfoque, los modelos analíticos ofrecen una alternativa complementaria, ya que permiten representar matemáticamente el comportamiento del sistema y obtener métricas de rendimiento con un coste computacional reducido [7].

En este contexto se sitúa el presente Trabajo Fin de Grado, cuyo objetivo principal es proporcionar una herramienta analítica de bajo coste computacional para estudiar la gestión de recursos en 5G NR V2X Modo 2 y evaluar su efecto sobre métricas relevantes de rendimiento. El trabajo se centra en el comportamiento de las reservas de recursos, las reselecciones y las reservas descartadas, considerando distintos modelos de tráfico y configuraciones temporales. Para ello, se propone una formulación probabilística basada en una cadena de Markov, capaz de representar la evolución del proceso de reserva y transmisión de paquetes.

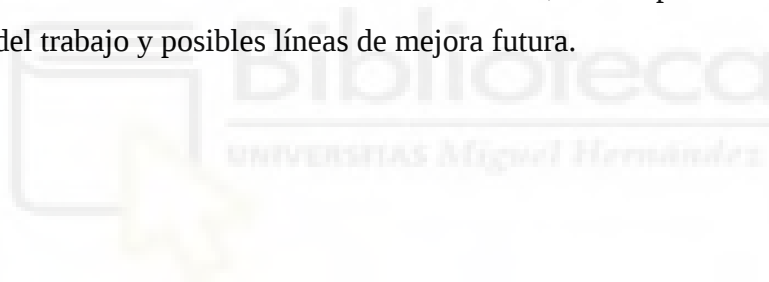
El modelo desarrollado permite obtener probabilidades asociadas a eventos relevantes del mecanismo de gestión de recursos, como la probabilidad de reselección, la probabilidad de reservar un determinado número de subcanales o la probabilidad de que una reserva no llegue a utilizarse. Además, el trabajo considera tanto tráfico periódico como tráfico aperiódico, ya que este último introduce situaciones adicionales de interés, como la posibilidad de que una reserva previamente anunciada se descarte o que el paquete generado no cumpla las condiciones de latencia necesarias para transmitirse en los recursos reservados.

Tras el modelado de la gestión de recursos, se aborda la caracterización de métricas temporales asociadas a la comunicación vehicular. En particular, se estudia la distribución de la latencia y el número de paquetes perdidos consecutivamente entre dos recepciones correctas. Estas métricas permiten ir más allá de una evaluación basada únicamente en la probabilidad media de recepción correcta, ya que en aplicaciones V2X no solo importa cuántos paquetes se reciben, sino también cómo se distribuyen temporalmente las pérdidas. Desde el punto de vista de la seguridad, no es equivalente perder paquetes aislados que sufrir ráfagas de pérdidas consecutivas durante un intervalo prolongado.

Con el objetivo de validar el modelo analítico propuesto, se desarrolla también un simulador de transmisiones en MATLAB. Este simulador reproduce, de forma controlada y con bajo coste computacional, el comportamiento del sistema desde el punto de vista de la generación de paquetes, la reserva de recursos, las reselecciones y las reservas descartadas. Su finalidad no es sustituir a simuladores de red completos, sino

proporcionar una herramienta intermedia que permita comprobar si las expresiones analíticas obtenidas describen correctamente el comportamiento esperado. Además, los resultados se contrastan con un simulador de referencia de mayor nivel de detalle, concretamente un simulador de comunicaciones vehiculares desarrollado en la UMH y basado en OMNeT++, INET y Veins, que implementa distintas tecnologías V2X. De este modo, se evalúa la validez del modelo analítico en un entorno de simulación más completo y representativo.

La memoria se organiza de la siguiente forma. Tras esta introducción, se presentan los fundamentos de las comunicaciones 5G NR V2X y se revisan los principales modelos analíticos existentes en la literatura. A continuación, se desarrolla el modelo analítico propuesto para la gestión de recursos y se formulan los modelos asociados a la latencia y a las pérdidas entre recepciones correctas. Posteriormente, se describen las herramientas y configuraciones empleadas para la validación, y se comparan los resultados analíticos con los obtenidos mediante simulación. Finalmente, se exponen las principales conclusiones del trabajo y posibles líneas de mejora futura.



2. ESTADO DEL ARTE

2.1. 5G NR V2X

En este capítulo se presentan los aspectos clave del 5G NR V2X en mayor detalle con el fin de entender el contexto de las próximas secciones y tener una visión global de esta tecnología.

2.1.1. Introducción

Como se ha introducido en el capítulo anterior, las comunicaciones V2X permiten el intercambio de información entre el vehículo y los distintos elementos de su entorno, ampliando su capacidad de percepción y coordinación más allá de la información obtenida mediante sensores instalados en el propio vehículo [3]. En este apartado se concreta esta idea en el contexto de 5G NR V2X, prestando atención a los elementos técnicos necesarios para el modelo desarrollado posteriormente: el enlace *sidelink*, los modos de asignación de recursos, la estructura física de los recursos y la relación entre los patrones de generación de paquetes y la gestión de recursos radio.

Dentro del ecosistema celular, las comunicaciones vehiculares han evolucionado desde LTE-V2X hasta 5G NR V2X. Esta evolución responde a la necesidad de soportar servicios vehiculares más exigentes, como la percepción cooperativa, la coordinación de maniobras, el *platooning*, la conducción remota o la protección de usuarios vulnerables de la vía [4], [6]. Frente a LTE-V2X, 5G NR V2X introduce una interfaz radio más flexible, definida por 3GPP en la Release 16, y orientada a escenarios con mayores requisitos de latencia, fiabilidad, capacidad y adaptación a distintos tipos de tráfico [5], [6].

En las comunicaciones vehiculares celulares, el enlace *sidelink* permite que los vehículos y otros dispositivos cercanos intercambien información directamente mediante la interfaz PC5, sin que los datos tengan que pasar necesariamente por la estación base. Esta capacidad es especialmente relevante en escenarios donde la cobertura de red no está garantizada, donde se requiere una latencia reducida o donde la comunicación debe mantenerse incluso en ausencia de infraestructura celular. No obstante, la red puede seguir desempeñando un papel importante cuando está disponible, por ejemplo, proporcionando configuración, asistencia o asignación de recursos en determinados modos de operación [5], [13].

En las comunicaciones *sidelink* de 5G NR V2X se definen dos modos principales de asignación de recursos. De forma general, el Modo 1 se basa en una asignación asistida por la red, mientras que el Modo 2 permite que los vehículos seleccionen autónomamente los recursos dentro de un conjunto preconfigurado. Este segundo modo constituye el caso de interés principal en este trabajo, ya que concentra el problema de gestión distribuida de recursos que se desarrollará posteriormente [8].

El tráfico que deben transmitir los vehículos puede presentar además distintos patrones de generación. En algunos casos, los paquetes se generan de forma periódica, con intervalos relativamente estables, mientras que en otros aparecen de manera aperiódica o asociados a eventos concretos. A esto se le añade que el tamaño de los mensajes puede variar en función de la información transmitida. Esta variabilidad afecta directamente a la gestión de recursos, ya que los recursos reservados para una transmisión pueden no ser adecuados para paquetes posteriores si cambia el tamaño del mensaje o si el instante de generación no coincide con la reserva previamente establecida [7].

Por este motivo, el estudio de 5G NR V2X Modo 2 requiere analizar de forma conjunta varios aspectos: la estructura física de los recursos, los canales de control y datos, el procedimiento de selección autónoma, los mecanismos de reserva y reelección, y los modelos de tráfico considerados, que determinan el instante de generación y el tamaño de los paquetes transmitidos. Todos estos elementos influyen en métricas como la probabilidad de recepción correcta, la latencia, la ocupación de recursos, la probabilidad de colisión o la continuidad temporal de las recepciones.

Los siguientes apartados presentan los elementos técnicos necesarios para contextualizar el modelo analítico desarrollado posteriormente. En primer lugar, se describe la capa física, incluyendo la organización tiempo-frecuencia, las numerologías, los subcanales y los canales físicos del enlace *sidelink*. A continuación, se analiza la capa de enlace, prestando especial atención al funcionamiento de la subcapa MAC, los modos de asignación de recursos, el mecanismo de selección autónoma y los eventos que pueden provocar cambios en las reservas.

2.1.2. Capa Física

La capa física de 5G NR V2X constituye la base sobre la que se realizan las transmisiones directas entre vehículos mediante el enlace *sidelink*. Su función principal es transformar la información procedente de capas superiores en señales radio que puedan transmitirse a

través del medio inalámbrico, así como recuperar en recepción la información contenida en dichas señales. No obstante, en el contexto de 5G NR V2X, la capa física no solo interviene en la transmisión y recepción de datos, sino que también proporciona información esencial para la gestión de recursos, la estimación de la calidad del enlace y la detección de ocupación del canal.

Esta relación entre capa física y gestión de recursos resulta especialmente importante en comunicaciones *sidelink* Modo 2, donde son los vehículos los que de forma autónoma gestionan los recursos. Aunque la decisión final de selección y reserva de recursos se realiza principalmente en la subcapa MAC, esta decisión se apoya en medidas obtenidas a nivel físico, como la potencia recibida, la detección de señales de referencia, la información de control transmitida por otros usuarios o la ocupación estimada de los recursos disponibles. Por tanto, para comprender el funcionamiento de la selección autónoma de recursos es necesario describir previamente cómo se organizan los recursos físicos en 5G NR V2X [9], [11], [12].

La transmisión en 5G NR V2X se basa en OFDM, del inglés “*Orthogonal Frequency Division Multiplexing*”. Esta técnica divide el ancho de banda disponible en múltiples subportadoras ortogonales entre sí, permitiendo transmitir información en paralelo sobre diferentes componentes frecuenciales. El uso de OFDM facilita una asignación flexible de recursos en el dominio tiempo-frecuencia y permite adaptar la transmisión a distintos anchos de banda, numerologías y condiciones de canal. Además, en la transmisión OFDM se emplea un prefijo cíclico para reducir los efectos de la interferencia entre símbolos provocada por la propagación multicamino [9].

En el dominio temporal, la estructura de 5G NR se organiza en tramas de 10 ms, divididas en 10 subtramas de 1 ms. A su vez, cada subtrama se divide en un número de *slots* que depende de la numerología utilizada. Esta numerología se identifica mediante el parámetro μ , que determina el espaciado entre subportadoras, o *Subcarrier Spacing* (SCS), según la expresión:

$$f_{SCS} = 2^{\mu} \cdot 15 \text{ kHz} \quad (1)$$

De este modo, al aumentar la numerología aumenta el espaciado entre subportadoras y disminuye la duración temporal del *slot*. Esta flexibilidad permite adaptar la granularidad temporal de los recursos a distintos requisitos de servicio. Por ejemplo, numerologías con mayor SCS permiten *slots* más cortos, lo que puede resultar útil para aplicaciones con requisitos estrictos de latencia, mientras que numerologías con menor SCS pueden ser

adecuadas en escenarios donde se prioriza una mayor robustez frente a determinados efectos del canal. Esta flexibilidad constituye una diferencia relevante respecto a LTE-V2X, donde la estructura temporal estaba más limitada por el uso de subtramas de 1 ms. En 5G NR V2X, la posibilidad de configurar diferentes numerologías permite ajustar con mayor precisión la duración de los recursos, la latencia de transmisión y la capacidad disponible. Esto resulta especialmente importante en comunicaciones vehiculares, donde pueden coexistir servicios con requisitos muy distintos, desde mensajes periódicos de estado hasta comunicaciones asociadas a percepción cooperativa o maniobras coordinadas [5].

En el dominio frecuencial, los recursos se organizan a partir de bloques de recursos físicos, o PRBs (*Physical Resource Blocks*). Cada PRB está formado por 12 subportadoras consecutivas en frecuencia. Estos PRBs se agrupan a su vez en subcanales, que constituyen una unidad fundamental para la transmisión de datos en *sidelink*. En 5G NR V2X, cada subcanal está compuesto por un número configurable de PRBs, normalmente denotado como M_{sub} , que depende de la configuración del *resource pool* [9], [11]. Un *resource pool* define el conjunto de recursos tiempo-frecuencia que pueden utilizarse para transmisiones directas entre dispositivos. Esta estructura se puede observar en la Figura 1:

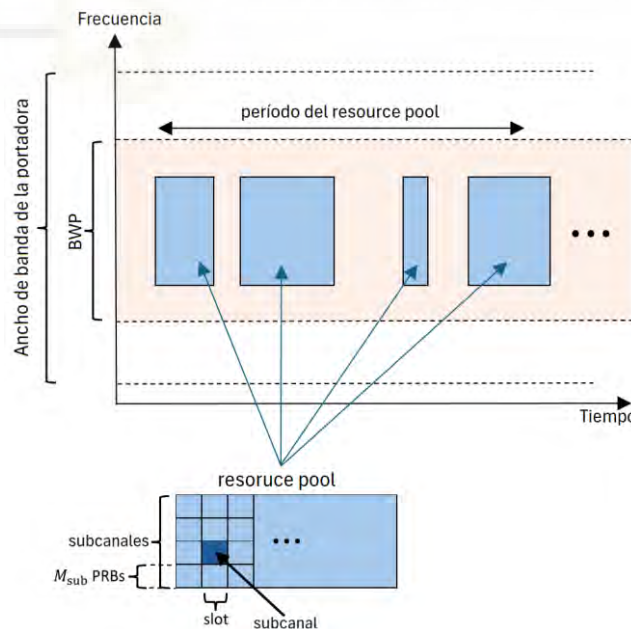


Figura 1. Estructura tiempo-frecuencia de los recursos *sidelink* en 5G NR V2X.

Dentro del *resource pool*, los vehículos seleccionan los recursos que emplearán para transmitir sus paquetes. A nivel frecuencial, este conjunto de recursos se encuentra asociado a una *Bandwidth Part* o BWP, que representa una porción contigua del ancho

de banda configurada para un determinado tipo de operación (ver Figura 1). En el caso de las comunicaciones *sidelink*, la BWP asignada al *resource pool* se divide en un número determinado de subcanales contiguos, sobre los que se realizará posteriormente el procedimiento de selección de recursos [8], [11].

La utilización de subcanales es especialmente importante desde el punto de vista del modelado de la gestión de recursos. Un paquete transmitido en 5G NR V2X puede ocupar uno o varios subcanales dentro de un mismo *slot*, dependiendo principalmente del tamaño del bloque de transporte, del esquema de modulación y codificación utilizado y de la configuración del sistema. Por tanto, el número de subcanales necesarios para transmitir un paquete no es necesariamente constante, sino que puede variar en función del tamaño del mensaje generado por las capas superiores, del MCS empleado y de la configuración de los recursos disponibles. Esta relación entre el tamaño del paquete, la configuración de transmisión y el número de subcanales ocupados es uno de los aspectos que se utilizará posteriormente en el modelo analítico desarrollado en este trabajo.

Por otro lado, la transmisión de información en 5G NR V2X se realiza mediante bloques de transporte, o TBs (*Transport Blocks*). Cada bloque de transporte contiene los datos que deben enviarse en una transmisión y se asocia a información de control que permite al receptor interpretar correctamente dicha transmisión. Esta información de control se denomina SCI, del inglés *Sidelink Control Information*. La SCI indica, entre otros aspectos, qué recursos se están utilizando, qué parámetros de transmisión se aplican y cómo debe procesarse el paquete recibido. Además, parte de esta información puede ser empleada por otros vehículos durante el sensado del canal, aspecto que se desarrollará con mayor detalle en la capa de enlace [9], [11], [13].

El enlace *sidelink* define varios canales físicos, cada uno con una función concreta. El primero de ellos es el PSCCH, o *Physical Sidelink Control Channel*. Este canal transporta información de control asociada a la transmisión *sidelink*, incluyendo información necesaria para localizar e interpretar el paquete transmitido. Desde el punto de vista del receptor, el PSCCH permite conocer parámetros esenciales de la transmisión. Además, la información transportada en este canal puede ser utilizada posteriormente en los procedimientos de sensado y selección de recursos descritos en la capa de enlace.

El segundo canal fundamental es el PSSCH, o *Physical Sidelink Shared Channel*. Este canal transporta los datos de usuario y parte de la información de control asociada a la transmisión. A través del PSSCH se transmiten los bloques de transporte generados por las capas superiores. La transmisión en este canal puede utilizar distintos esquemas de

modulación y codificación, definidos mediante el MCS (*Modulation and Coding Scheme*). El MCS determina conjuntamente la modulación empleada y la tasa de codificación, por lo que afecta directamente a la eficiencia espectral y a la robustez de la comunicación. Modulaciones de mayor orden permiten transmitir más bits por símbolo, pero requieren mejores condiciones de canal para mantener una probabilidad de error reducida. En cambio, modulaciones más robustas reducen la tasa de transmisión, pero permiten mejorar la fiabilidad en condiciones radio más degradadas [9], [11].

Además del PSCCH y el PSSCH, 5G NR V2X introduce el PSFCH, o *Physical Sidelink Feedback Channel*. Este canal se utiliza para transmitir realimentación HARQ en determinadas configuraciones de comunicación *sidelink*. HARQ, del inglés “*Hybrid Automatic Repeat Request*”, combina mecanismos de detección de errores y retransmisión con técnicas de codificación para mejorar la fiabilidad de la comunicación. Mediante el PSFCH, un receptor puede informar al transmisor sobre si un paquete ha sido recibido correctamente o no, permitiendo retransmisiones cuando la configuración del sistema lo contempla [9], [10], [11].

Las señales de referencia, como las DMRS (*Demodulation Reference Signals*), permiten estimar el canal radio y obtener medidas de potencia útiles para la recepción y el sensado. A partir de las señales recibidas pueden obtenerse medidas de potencia como el RSSI (*Received Signal Strength Indicator*) y el RSRP (*Reference Signal Received Power*). El RSSI proporciona una medida de la potencia recibida sobre un conjunto de recursos, mientras que el RSRP mide la potencia recibida asociada a señales de referencia. Estas magnitudes son relevantes porque proporcionan información física sobre la ocupación y la calidad de los recursos radio, que posteriormente puede ser utilizada por los procedimientos de selección de recursos en Modo 2 [9], [11], [12].

Desde el punto de vista del rendimiento, la capa física condiciona directamente la probabilidad de recepción correcta de un paquete. Incluso aunque un recurso haya sido seleccionado correctamente desde el punto de vista del acceso al medio, la transmisión puede fallar por distintos motivos físicos o de acceso al medio, como potencia recibida insuficiente, interferencia, errores de propagación, *half-duplex* (el vehículo no puede transmitir y recibir a la vez) o colisión con otra transmisión.

En 5G NR V2X, la posibilidad de que un bloque de transporte ocupe uno o varios subcanales hace que las colisiones puedan tener distinta naturaleza. Si dos paquetes se solapan completamente en los mismos recursos tiempo-frecuencia, se produce una colisión total. Sin embargo, si el solapamiento afecta únicamente a una parte de los PRBs

o subcanales ocupados por el paquete, se produce una colisión parcial. Esta distinción es relevante porque una colisión parcial no siempre tiene el mismo efecto que una colisión total: su impacto dependerá del número de recursos afectados, del MCS utilizado, de la codificación y de la robustez del enlace. Por tanto, la organización física de los recursos influye directamente en la forma en que deben interpretarse las pérdidas por colisión [7]. Aunque físicamente puedan distinguirse colisiones totales y parciales, en el modelo de pérdidas entre recepciones correctas desarrollado posteriormente estos efectos se incorporan de forma conjunta dentro de la probabilidad global de pérdida por colisión.

2.1.3. Capa de Enlace

La capa de enlace de datos en 5G NR V2X desempeña un papel fundamental en el funcionamiento de las comunicaciones *sidelink*, especialmente en todo lo relacionado con el acceso al medio, la selección de recursos, la reserva de transmisiones futuras, la interacción con los procesos HARQ y el control de congestión. Aunque la capa física define la estructura tiempo-frecuencia y los canales físicos sobre los que se transmite la información, es la subcapa MAC la que determina cómo accede cada vehículo al canal, qué recursos utiliza, durante cuánto tiempo mantiene una reserva y en qué condiciones debe modificarla. Por este motivo, el estudio de la capa de enlace está estrechamente ligado a la gestión distribuida de recursos, ya que su comportamiento influye directamente en métricas como la probabilidad de recepción correcta, la latencia, la ocupación del canal, la eficiencia espectral y la continuidad temporal de las transmisiones [13].

Desde el punto de vista de la asignación de recursos, los modos 1 y 2 de 5G NR V2X definen dos principios de operación en comunicaciones *sidelink*, pudiendo interpretarse como la evolución de los modos 3 y 4 de LTE-V2X, respectivamente. En el Modo 1, la asignación de recursos está controlada por la red, de forma que una estación base, generalmente un gNB, proporciona al vehículo los recursos que debe utilizar para transmitir por *sidelink*. Este modo requiere cobertura de red y permite una planificación centralizada, lo que puede reducir la probabilidad de conflictos entre usuarios siempre que la infraestructura tenga información suficiente del estado del sistema [5], [8]. La red puede asignar recursos mediante concesiones dinámicas, en las que el vehículo solicita recursos cuando necesita transmitir, o mediante concesiones configuradas, donde se preasignan recursos con una determinada periodicidad. La primera opción proporciona

mayor adaptación al tráfico instantáneo, mientras que la segunda reduce la señalización y la latencia de acceso, aunque puede desaprovechar recursos si el tráfico previsto no llega a generarse [13].

En el Modo 2, por el contrario, los vehículos seleccionan de forma autónoma los recursos de transmisión dentro de un *resource pool* configurado. En este caso no existe una entidad central que asigne los recursos a cada usuario en cada instante, por lo que cada vehículo debe decidir qué recursos utilizar basándose en la información obtenida mediante el sensado del canal y la información de control transmitida por otros usuarios. Esta autonomía permite el funcionamiento fuera de cobertura o en escenarios donde no se desea depender de una asignación directa por parte de la infraestructura [5], [13].

Dentro del Modo 2 se definen dos esquemas principales de planificación de recursos: la planificación dinámica y la planificación semi-persistente basada en sensado, conocida habitualmente como “*Sensing-Based Semi-Persistent Scheduling*” o SPS. En la planificación dinámica, el vehículo selecciona nuevos recursos para cada bloque de transporte, reservando recursos futuros únicamente para posibles retransmisiones del mismo paquete. Este esquema proporciona flexibilidad ante tráfico irregular o impredecible, ya que no presupone que los paquetes futuros vayan a generarse con una periodicidad determinada. En cambio, en el esquema SPS, el vehículo selecciona un conjunto de recursos y los reutiliza durante varias transmisiones consecutivas, lo que resulta especialmente adecuado para servicios V2X con tráfico periódico. Sin embargo, esta misma persistencia puede provocar problemas si la reserva seleccionada resulta conflictiva, si los paquetes no se generan en el instante esperado o si el tamaño de los paquetes varía con el tiempo [5], [7], [13], [14].

El funcionamiento del SPS se apoya principalmente en dos parámetros: el *Resource Reservation Interval* (RRI) y el *Reselection Counter* (RC). El RRI define el intervalo temporal entre dos transmisiones consecutivas que reutilizan el mismo patrón de recursos. Por tanto, debe estar relacionado con las características temporales del tráfico generado por la aplicación, ya que una mala elección puede provocar reservas demasiado frecuentes, oportunidades reservadas sin paquete disponible para transmitir o transmisiones que no cumplen los requisitos de latencia.

El RC determina durante cuántas transmisiones consecutivas se mantiene una reserva antes de considerar una posible reelección de recursos. Cuando un vehículo selecciona nuevos recursos mediante SPS, inicializa el RC con un valor aleatorio dentro de un intervalo definido por el estándar y dependiente del RRI. Tras cada transmisión realizada

con la reserva actual, el contador se decrementa en una unidad. Cuando alcanza el valor cero, el vehículo decide si mantiene los recursos previamente seleccionados o si realiza una nueva selección, en función de una probabilidad de permanencia, P [7], [13].

Este mecanismo introduce aleatoriedad en la duración de las reservas y evita que un vehículo mantenga indefinidamente los mismos recursos. Si la probabilidad de permanencia es elevada, el entorno de reservas es más estable y los vehículos vecinos pueden predecir mejor el uso futuro del canal. Sin embargo, si dos vehículos han seleccionado recursos incompatibles, una probabilidad de permanencia alta puede prolongar la colisión durante varias transmisiones consecutivas. Por el contrario, forzar reselecciones más frecuentes puede reducir la duración de colisiones persistentes, pero también aumenta la variabilidad del sistema y la probabilidad de que varios vehículos vuelvan a seleccionar recursos coincidentes. Para el modelado analítico de este trabajo, se asume $P = 0$, por lo que, cuando el RC alcanza cero, se produce siempre una reselección de recursos.

El procedimiento de selección de recursos en Modo 2 se inicia cuando el vehículo necesita transmitir un paquete y no dispone de una reserva válida, o cuando se produce algún evento que obliga a abandonar los recursos previamente reservados. En ese instante, el vehículo define una ventana de selección comprendida entre los instantes $n + T_1$ y $n + T_2$, donde n representa el instante en el que se activa el proceso de selección. El parámetro T_1 está relacionado con el tiempo mínimo necesario para procesar la selección, mientras que T_2 debe situarse dentro del presupuesto de retardo permitido por el paquete, conocido como “*Packet Delay Budget*” o PDB. Por tanto, la ventana de selección delimita los recursos que podrían emplearse para transmitir el paquete sin incumplir su restricción máxima de latencia [5], [7], [11]. En la Figura 2 se muestra un ejemplo para $T_2 = \text{PDB}$.

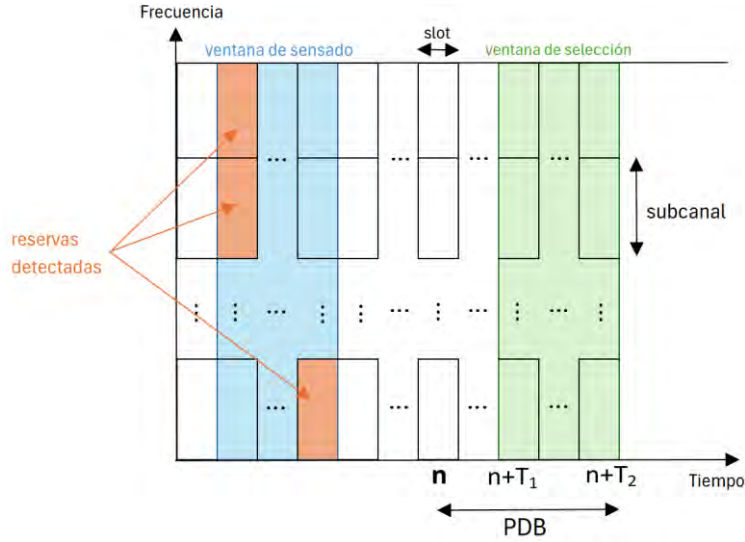


Figura 2. Ventanas de sensado y selección en NR V2X Modo 2 para el caso $T_2 = PDB$

Dentro de esta ventana, el vehículo construye un conjunto de recursos candidatos. Cada recurso candidato está definido por una posición temporal y por un conjunto de subcanales contiguos en frecuencia. Como se ha explicado en la capa física, el número de subcanales necesarios depende del tamaño del bloque de transporte, del MCS utilizado y de la configuración del *resource pool*. Esta relación es especialmente relevante para este trabajo, ya que una reserva realizada para un paquete de cierto tamaño puede no ser suficiente para transmitir un paquete posterior de mayor tamaño. En ese caso, el vehículo deberá seleccionar nuevos recursos con capacidad suficiente.

La selección de recursos se basa en un proceso de sensado del canal. Durante una ventana temporal anterior al instante de selección (ver Figura 2), el vehículo observa los recursos *sidelink* disponibles, decodifica la información de control transmitida por otros usuarios y utiliza medidas físicas como el RSRP o el RSSI para valorar la potencia recibida de las transmisiones detectadas. A partir de esta información, la subcapa MAC excluye del conjunto de candidatos aquellos recursos que presentan mayor riesgo de estar ocupados o de provocar colisiones.

El procedimiento de selección puede interpretarse en dos fases. En una primera fase se eliminan los recursos candidatos que aparecen reservados por otros vehículos mediante la SCI y cuya potencia recibida supera un determinado umbral. También pueden excluirse recursos asociados a intervalos en los que el vehículo no pudo sensar el canal por encontrarse transmitiendo, debido a la restricción *half-duplex*. En esos intervalos, el vehículo puede perder información relevante sobre reservas futuras, por lo que dichos

recursos pueden considerarse menos fiables dentro del proceso de selección. Si tras esta primera fase el número de recursos candidatos disponibles es demasiado bajo, el procedimiento puede repetirse relajando progresivamente el umbral de exclusión. De este modo, se garantiza que el vehículo disponga de un conjunto mínimo de recursos candidatos entre los que elegir. En una segunda fase, el vehículo selecciona aleatoriamente uno o varios recursos dentro del conjunto no excluido. Esta aleatoriedad es necesaria porque varios vehículos cercanos pueden observar condiciones de canal similares y, por tanto, podrían llegar a conjuntos de candidatos parecidos. Si todos seleccionaran siempre el recurso aparentemente óptimo, la probabilidad de colisión aumentaría [7], [11], [12].

En este procedimiento resulta especialmente importante la organización de la SCI en dos etapas. La primera etapa contiene información necesaria para el sensado y la selección de recursos, como la asignación en tiempo y frecuencia, la prioridad del paquete y determinados parámetros de reserva. La segunda etapa contiene información más específica de la transmisión, como campos relacionados con HARQ, identificadores de origen y destino o información asociada al tipo de comunicación [5], [11], [13]. De este modo, la SCI permite que la selección autónoma de recursos no dependa únicamente de medidas de potencia, sino también de información explícita sobre el uso presente y futuro de los recursos.

Además del procedimiento básico de selección, 5G NR V2X incorpora mecanismos destinados a mejorar la robustez de la asignación de recursos. Uno de ellos es la re-evaluación, que permite comprobar si un recurso previamente seleccionado sigue siendo válido antes del instante de transmisión. Si durante ese intervalo el vehículo detecta que el recurso ha pasado a estar ocupado o debe excluirse según las reglas de selección, puede ejecutar de nuevo el procedimiento de selección y escoger un recurso diferente. También existe el mecanismo de *preemption*, o liberación de recursos por prioridad, mediante el cual transmisiones de mayor prioridad pueden obligar a modificar recursos previamente seleccionados por transmisiones de menor prioridad [5], [11]. Estos mecanismos proporcionan mayor capacidad de adaptación ante cambios en la ocupación del canal, aunque no se modelan de forma explícita en este trabajo.

Dentro del funcionamiento SPS aparecen distintos eventos que pueden obligar a abandonar una reserva y seleccionar nuevos recursos. En este trabajo se consideran especialmente relevantes tres causas de reelección: la reelección por agotamiento del contador, la reelección por tamaño y la reelección por latencia. Estas tres causas

conectan directamente el comportamiento de la subcapa MAC con el modelo analítico desarrollado posteriormente, ya que determinan cuándo el vehículo mantiene su reserva y cuándo debe iniciar un nuevo proceso de selección [7], [14].

La reelección por contador es la más directamente asociada al funcionamiento básico del SPS. Como se ha explicado, cuando el *Reselection Counter* alcanza el valor cero, el vehículo puede abandonar la reserva actual y seleccionar nuevos recursos. En el modelo analítico desarrollado en este trabajo, este evento se representa como una de las causas de cambio de estado dentro de la cadena de Markov.

La reelección por tamaño aparece cuando el nuevo paquete generado por las capas superiores no cabe en los recursos previamente reservados. Esto puede ocurrir cuando el tamaño de los paquetes es variable. Por ejemplo, si un vehículo reserva un determinado número de subcanales para transmitir un paquete pequeño, pero posteriormente se genera un paquete mayor que requiere más subcanales, la reserva anterior deja de ser suficiente. En ese caso, el vehículo debe seleccionar nuevos recursos que permitan transmitir el nuevo bloque de transporte. Este fenómeno es especialmente importante en modelos de tráfico con tamaños de mensaje variables, ya que introduce reelecciones adicionales que no aparecerían en un escenario de tamaño constante.

La reelección por latencia aparece cuando la próxima reserva disponible no permite cumplir el PDB del paquete. La latencia de transmisión puede definirse como la diferencia entre el instante en que el paquete se genera en capas superiores y el instante en que se transmite por el medio radio. Si un paquete se genera en un instante determinado y el siguiente recurso reservado cae demasiado tarde, de forma que la transmisión superaría la latencia máxima permitida, el vehículo no puede esperar a dicha reserva. En consecuencia, debe seleccionar nuevos recursos dentro de la ventana temporal válida. Este fenómeno aparece en tráfico aperiódico, donde los instantes de generación de paquetes no coinciden necesariamente con la periodicidad de las reservas SPS.

En escenarios periódicos ideales, donde el intervalo de generación de paquetes coincide con el RRI y el tamaño de paquete permanece constante, el comportamiento del SPS es relativamente estable. La reserva seleccionada puede reutilizarse durante varias transmisiones consecutivas y las principales reelecciones se deben al agotamiento del contador. Sin embargo, en escenarios más realistas, los paquetes pueden generarse de forma aperiódica, presentar tamaños variables o tener requisitos de latencia estrictos. En estas condiciones, la reserva previamente seleccionada puede dejar de ser válida, provocando reelecciones adicionales por tamaño o por latencia.

Además de las reselecciones, en el funcionamiento SPS pueden aparecer oportunidades de transmisión previamente reservadas que finalmente no lleguen a transportar ningún paquete. En la especificación 3GPP, este fenómeno se relaciona con las *unused transmission opportunities* asociadas al *selected sidelink grant*. En particular, 3GPP contabiliza como oportunidades de transmisión no utilizadas aquellas en las que ninguno de los recursos del *selected sidelink grant* dentro de un RRI llega a utilizarse [13]. Por un lado, un vehículo puede abandonar los recursos previamente seleccionados antes del instante de transmisión si se activa una nueva reselección. Por otro lado, puede ocurrir que el instante reservado llegue antes de que se haya generado un nuevo paquete en las capas superiores. En este último caso, la reserva existe desde el punto de vista del mecanismo SPS, pero no hay ningún bloque de transporte disponible para transmitir [7], [14].

No obstante, en este trabajo se analiza específicamente la probabilidad de que una reserva no se utilice por no existir ningún paquete disponible para transmitir en el instante reservado. Para evitar ambigüedad con el concepto general de oportunidad de transmisión no utilizada, se emplea el término reserva descartada para referirse exclusivamente a este caso.

Además de los mecanismos anteriores, la capa de enlace también contempla procedimientos de control de congestión, cuyo objetivo es limitar la ocupación del canal cuando aumenta el número de usuarios o transmisiones. De forma general, estos mecanismos utilizan medidas de ocupación del canal, como el *Channel Busy Ratio* (CBR), para adaptar el uso de los recursos y evitar una sobrecarga excesiva del medio [13].

En conjunto, la capa de enlace permite situar el problema abordado en este trabajo dentro de la gestión distribuida de recursos en 5G NR V2X. De los modos y esquemas descritos, el análisis se centrará en el Modo 2 con SPS, al ser el caso más relevante desde el punto de vista de la gestión autónoma de reservas de recursos. Este planteamiento proporciona la base conceptual para el modelo analítico desarrollado en el capítulo 3.

2.2. Modelos analíticos existentes

Como se ha introducido previamente, la evaluación de las comunicaciones vehiculares puede abordarse mediante simulación o mediante modelado analítico. En este apartado

se revisan los modelos analíticos más relevantes relacionados con las comunicaciones *sidelink* autónomas y, en particular, con la gestión distribuida de recursos en 5G NR V2X Modo 2. El objetivo es identificar qué aspectos han sido ya tratados en la literatura y qué limitaciones justifican el modelo desarrollado en este trabajo.

Los primeros antecedentes relevantes no se desarrollaron directamente para 5G NR V2X, sino para LTE-V2X Modo 4, que constituye el precursor del Modo 2 de NR V2X en comunicaciones *sidelink* autónomas. En ambos casos, los vehículos seleccionan recursos de forma distribuida y utilizan información obtenida mediante sensado del canal para reducir la probabilidad de seleccionar recursos ocupados por otros usuarios. Por ello, varios modelos formulados inicialmente para LTE-V2X han servido como base conceptual para trabajos posteriores sobre NR V2X.

Uno de los trabajos más influyentes en el modelado analítico de LTE-V2X Modo 4 es el presentado en [15]. En dicho trabajo se propone un modelo analítico para caracterizar la *Packet Delivery Ratio* o PDR, equivalente en este contexto a la *Packet Reception Ratio* o PRR, en función de la distancia entre transmisor y receptor. Una de sus principales aportaciones es la descomposición de la probabilidad de error en varias causas diferenciadas: errores por *half-duplex*, errores debidos a potencia recibida por debajo del umbral de sensibilidad, errores de propagación y errores por colisión. Esta separación resulta especialmente relevante porque permite identificar de forma individual qué fenómenos limitan la recepción correcta de los paquetes. Aunque el modelo fue diseñado para LTE-V2X Modo 4 y considera tráfico periódico con tamaño de paquete constante, su estructura ha servido como punto de partida para modelos posteriores de comunicaciones *sidelink* autónomas, incluidos algunos modelos de 5G NR V2X Modo 2. Además de los modelos centrados en la PDR o PRR en función de la distancia, existen trabajos que modelan el comportamiento de la capa MAC. En [16] se propone un modelo basado en una cadena de Markov discreta multidimensional para comparar el rendimiento de LTE-V2X Modo 4 con IEEE 802.11p. Este modelo considera mensajes periódicos y mensajes generados por evento de tamaño constante, y permite obtener métricas como el retardo medio de acceso al canal, la utilización del canal y la probabilidad de colisión. Su interés principal radica en la caracterización del comportamiento MAC, pero su ámbito es más limitado desde el punto de vista radio, ya que no incorpora de forma explícita la distancia entre vehículos como variable de análisis ni modela con detalle fenómenos como la propagación, la sensibilidad de recepción o los terminales ocultos. Por tanto,

aunque resulta útil para estudiar el acceso al medio, no proporciona una caracterización completa de la PDR frente a la distancia transmisor-receptor.

Otros trabajos dentro del ámbito LTE-V2X se centran en escenarios o casos de uso más específicos. En [17] se presenta un marco analítico basado en geometría estocástica para estudiar comunicaciones C-V2X sobre canales compartidos entre comunicaciones V2V y enlace ascendente celular. En este planteamiento, los vehículos pueden entregar información mediante comunicación directa o a través de la infraestructura, lo que introduce una perspectiva diferente respecto a los modelos puramente *sidelink*. Por su parte, [18] propone un modelo analítico basado en cadenas de Markov para evaluar el mecanismo de reserva de recursos de LTE-V2X Modo 4 aplicado a *platooning* vehicular. En este caso, el interés principal se sitúa en un escenario cooperativo específico, donde la coordinación y regularidad de las transmisiones dentro del pelotón son aspectos fundamentales. Estos trabajos muestran la diversidad de enfoques analíticos existentes en C-V2X, pero no abordan todavía las particularidades introducidas posteriormente por NR V2X Modo 2.

Con la aparición de 5G NR V2X Modo 2 comenzaron a desarrollarse modelos analíticos específicos para esta tecnología. Uno de los primeros es el presentado en [19], donde se propone un modelo de PRR y *throughput* para NR C-V2X Modo 2. Este modelo, parcialmente derivado de trabajos previos sobre LTE-V2X Modo 4, se centra en el esquema SPS con tráfico periódico y tamaño de paquete constante. Su interés principal radica en trasladar al contexto NR una formulación analítica de la PRR y el *throughput* en comunicaciones *sidelink* autónomas. Sin embargo, no incorpora de forma completa algunas características propias de NR V2X, como la numerología o las retransmisiones, y su precisión puede verse limitada en escenarios con elevada carga de canal. Además, no aborda la generación aperiódica de paquetes, la variabilidad del tamaño de los mensajes ni las reelecciones asociadas a tamaño insuficiente o a restricciones de latencia. En [20] se presentan modelos analíticos para estudiar la asignación distribuida de recursos en NR V2X Modo 2. Este trabajo analiza la influencia de la numerología y de la configuración de recursos sobre el rendimiento, y constituye una de las primeras aproximaciones analíticas orientadas específicamente a NR V2X. Su principal aportación es introducir parámetros propios de NR, como la separación entre subportadoras, la duración de *slot* y el número de subcanales disponibles, permitiendo estudiar cómo afectan a la probabilidad de recepción. No obstante, su alcance sigue estando centrado en configuraciones relativamente simplificadas, con tráfico periódico y tamaño de paquete

constante. Además, no incorpora de forma completa todos los mecanismos de reelección introducidos por NR V2X Modo 2, por lo que no resulta suficiente para caracterizar escenarios donde los paquetes generados por los vehículos presentan tamaños variables o instantes de generación no periódicos.

Otro trabajo cercano al modelado analítico de NR V2X Modo 2 es [21], donde se evalúa el rendimiento del Modo 2 mediante modelado analítico y simulaciones. Este modelo estudia la PRR, equivalente a la PDR, para los esquemas SPS y *Dynamic scheduling*. Su interés principal radica en que aborda específicamente NR V2X Modo 2 y permite comparar dos esquemas de planificación definidos para *sidelink*. Sin embargo, su alcance se limita a tráfico periódico con tamaño de paquete constante. Además, en el caso de SPS, el modelo considera únicamente las reelecciones asociadas al agotamiento del *Reselection Counter*, sin incorporar reelecciones por tamaño insuficiente de la reserva ni por incumplimiento de restricciones de latencia. Por tanto, aunque constituye uno de los antecedentes más próximos dentro del modelado analítico de NR V2X Modo 2, no cubre los escenarios de tráfico aperiódico y tamaño variable considerados en este trabajo. También se han propuesto modelos analíticos orientados específicamente al tráfico esporádico. En [22] se desarrolla un modelo analítico de 5G V2X Modo 2 para estimar la tasa de pérdida de paquetes y la capacidad de red en escenarios con tráfico esporádico o generado por evento. Este trabajo es relevante porque introduce el análisis de tráfico no periódico en NR V2X Modo 2, lo que lo convierte en un antecedente parcial importante para modelos que consideren generación aperiódica de paquetes. Sin embargo, el modelo está orientado a situaciones de baja carga y adopta hipótesis simplificadoras, considerando únicamente algunos mecanismos de pérdida, como los errores por *half-duplex* y por colisión. Además, no modela de forma completa la dinámica asociada al mantenimiento, abandono o reutilización de reservas SPS. Por tanto, aunque supone un avance respecto a modelos puramente periódicos, no cubre el caso general de tráfico aperiódico con tamaño variable ni sus efectos sobre las reelecciones de recursos.

Desde el punto de vista de las métricas temporales, un trabajo especialmente relacionado con el *Packet Inter-Reception* (PIR) es [23]. En este artículo se analiza la planificación dinámica en NR V2X Modo 2 y se obtiene la probabilidad de entrega correcta de paquetes, así como la estadística del intervalo temporal entre dos recepciones correctas consecutivas de un mismo flujo. Este aspecto es relevante porque va más allá de la PRR media y aborda la continuidad temporal de la recepción, que resulta fundamental en servicios vehiculares donde no solo importa cuántos paquetes se reciben, sino también

cuánto tiempo transcurre entre recepciones exitosas. No obstante, el modelo se centra en el esquema de planificación dinámica, en el que se seleccionan nuevos recursos para cada paquete y se reservan recursos únicamente para posibles retransmisiones. Por tanto, aunque constituye un antecedente directo para el análisis del Packet Inter-Reception, no aborda el comportamiento de SPS con reservas mantenidas durante varios paquetes ni las reselecciones provocadas por la evolución del tamaño de los paquetes o por restricciones de latencia.

Por último, [24] propone un modelo analítico de NR V2X Modo 2 con mecanismo de *re-evaluation*, incorporando un generador de mensajes basado en cadenas de Markov discretas para representar patrones de tráfico recomendados por 3GPP. Este trabajo resulta relevante porque aborda de forma analítica un mecanismo específico de NR V2X y estudia métricas MAC como la probabilidad de colisión y la latencia. Además, reconoce la importancia de los mensajes aperiódicos y del tráfico variable en la aparición de colisiones. Sin embargo, se trata de un trabajo centrado principalmente en el comportamiento MAC asociado a la *re-evaluation*. No proporciona una caracterización completa de la gestión de recursos SPS con tamaños variables de paquete, reservas no utilizadas y reselecciones por tamaño y latencia.

La comparación de estos trabajos permite identificar varias limitaciones comunes en la literatura existente. En primer lugar, muchos modelos analíticos relevantes se desarrollaron originalmente para LTE-V2X Modo 4 y, aunque son útiles como base conceptual, no incorporan de manera completa las nuevas funcionalidades de NR V2X Modo 2. En segundo lugar, buena parte de los modelos específicos para NR V2X consideran tráfico periódico y tamaño de paquete constante, lo que simplifica el análisis, pero no representa adecuadamente servicios vehiculares avanzados, donde pueden aparecer paquetes aperiódicos y tamaños variables. En tercer lugar, algunos modelos se concentran en métricas MAC, como la probabilidad de colisión, la utilización del canal o el retardo de acceso, pero no proporcionan una caracterización completa de la PRR o PDR frente a la distancia. En cuarto lugar, otros trabajos sí consideran tráfico aperiódico o generado por evento, pero lo hacen bajo hipótesis simplificadas o sin modelar de forma detallada la evolución de las reservas SPS.

Conviene matizar, no obstante, que esta conclusión no implica la ausencia total de trabajos sobre tráfico aperiódico, tamaño variable o métricas temporales en comunicaciones V2X. Como se ha indicado, existen modelos analíticos que cubren partes concretas del problema. Sin embargo, no se han identificado modelos previos que

integren en una misma formulación la generación aperiódica de paquetes, la variabilidad del tamaño, las reservas no utilizadas, las reselecciones por agotamiento del contador, por tamaño insuficiente y por restricciones de latencia, junto con el modelado posterior de métricas temporales como la latencia y las pérdidas entre recepciones correctas.

En este contexto se sitúa la contribución del presente trabajo. Partiendo de los modelos analíticos previos de comunicaciones *sidelink* autónomas y de las limitaciones identificadas en NR V2X Modo 2, se desarrolla un modelo basado en una cadena de Markov capaz de representar la evolución de la reserva y uso de recursos cuando los paquetes presentan tamaños variables y generación aperiódica. Este enfoque permite obtener probabilidades asociadas al proceso de gestión de recursos, como la probabilidad de reselección, la probabilidad de reserva descartada por ausencia de paquete generado y la probabilidad de que sea necesario reservar un determinado número de subcanales. A diferencia de modelos que consideran únicamente el agotamiento del *Reselection Counter*, este planteamiento permite incorporar otros disparadores de reselección relevantes en NR V2X, como la insuficiencia de recursos reservados para el nuevo tamaño de paquete o el incumplimiento de las restricciones de latencia. Con estas magnitudes adicionales, el modelo permite caracterizar la latencia y el número de paquetes perdidos consecutivamente entre dos recepciones correctas para modelos con tráfico aperiódico y tamaño de paquete variable.

3. MODELADO ANALÍTICO

El objetivo de este capítulo es presentar los modelos analíticos desarrollados para caracterizar la gestión de recursos en 5G NR V2X Modo 2 y su impacto sobre distintas métricas de rendimiento. En primer lugar, se desarrolla un modelo de gestión de recursos basado en una cadena de Markov, orientado a representar la evolución de las reservas SPS, las reselecciones de recursos, las reservas descartadas y el número de subcanales reservados. Posteriormente, las probabilidades obtenidas mediante este modelo se utilizan como base para formular dos modelos adicionales: un modelo de distribución de la latencia y un modelo de pérdidas consecutivas entre recepciones correctas. De este modo, el capítulo proporciona una formulación analítica que permite estudiar el comportamiento del mecanismo de gestión de recursos y sus efectos sobre métricas relevantes sin depender exclusivamente de simulaciones de mayor coste computacional.

3.1. Modelo de gestión de recursos

3.1.1. Antecedentes y planteamiento del modelo

El modelo desarrollado en este trabajo parte de una formulación analítica previa desarrollada en el grupo de investigación para modelar la reserva y reelección de recursos en 5G NR V2X con tráfico periódico y tamaños de paquete variables. En dicho modelo, los paquetes se generan con un intervalo constante entre transmisiones, pero su tamaño puede variar de forma probabilística e independiente entre paquetes, siguiendo los modelos de tráfico periódicos propuestos por 3GPP. Esta formulación permite obtener magnitudes como la probabilidad de que un vehículo reserve un determinado número de subcanales y la probabilidad de que tenga que realizar una reelección de recursos. Ambas magnitudes son relevantes, ya que el número de subcanales reservados determina la ocupación de recursos del sistema, mientras que la probabilidad de reelección influye directamente en la posibilidad de que dos nodos seleccionen recursos coincidentes.

En este modelo de partida, cada paquete generado puede tener un tamaño diferente, y dicho tamaño se traduce posteriormente en un número de subcanales necesarios para su transmisión. Aunque los tamaños de los paquetes se especifican inicialmente en bytes, el modelo requiere como entrada la probabilidad de que un nuevo paquete necesite s subcanales para su transmisión, con $1 \leq s \leq S$, donde S representa el número total de subcanales disponibles en el ancho de banda considerado. Esta probabilidad depende de

la distribución de tamaños de paquete y de los parámetros de transmisión empleados, como el MCS, la numerología y la configuración de subcanales, como se ha explicado en la capa física.

El comportamiento del mecanismo SPS se representa mediante una cadena de Markov, a partir de cuya matriz de transición se obtiene la distribución estacionaria del proceso. Con dicha distribución pueden calcularse la probabilidad de selección o reelección de recursos y la probabilidad de reservar un determinado número de subcanales. Esta cadena constituye una base adecuada para representar el comportamiento del sistema cuando el tráfico es periódico y la variabilidad procede únicamente del tamaño de los paquetes. Sin embargo, este planteamiento no contempla de forma completa los modelos de tráfico aperiódicos, que son necesarios para el análisis realizado en este trabajo.

En los modelos periódicos, el intervalo entre generaciones consecutivas de paquetes es constante. Como consecuencia, una vez realizada una reserva de recursos, la relación temporal entre la generación del paquete y su transmisión (lo que llamamos latencia de transmisión) se mantiene constante mientras no se produzca una reelección, como se observa en la Figura 3.

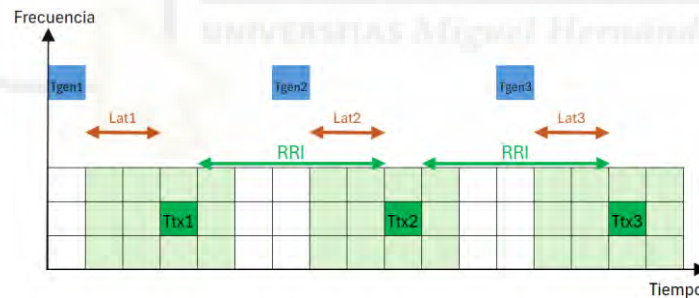


Figura 3. Esquema de transmisiones.

En cambio, en los modelos aperiódicos, el intervalo entre generaciones de paquetes es aleatorio, lo que introduce nuevos eventos que no aparecen en el modelo de partida.

Para los modelos de tráfico aperiódico, 3GPP no fija un único valor de RRI , sino que este puede configurarse en función del escenario y del tipo de tráfico considerado. En este trabajo se parte de los modelos de tráfico propuestos por 3GPP, pero se amplía el análisis mediante un conjunto de modelos de menor a mayor complejidad. En concreto, se consideran modelos periódicos y aperiódicos con $RRI = PDB$ y modelos aperiódicos con $RRI = 2PDB$, completando finalmente un conjunto de siete modelos de estudio. La descripción detallada de estos modelos se desarrolla en el capítulo de validación, mientras

que en este capítulo se presenta la formulación analítica general que permite tratarlos dentro de una misma cadena de Markov.

En los modelos aperiódicos, el instante en que el siguiente paquete está disponible puede adelantarse o retrasarse respecto al recurso previamente reservado, como ya se ha explicado en el apartado de capa de enlace. Como consecuencia, la latencia deja de ser constante y pasa a depender de la realización concreta del proceso aleatorio de generación de paquetes, como se observa en la Figura 4.

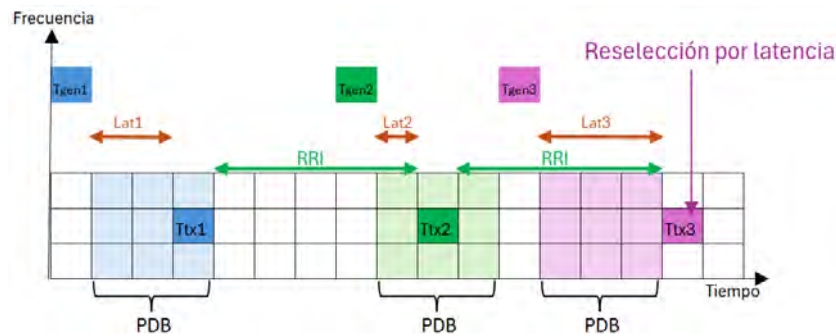


Figura 4. Esquema de reselecciones por latencia.

Además, existe otro escenario propio del tráfico aperiódico que tampoco aparece en el modelo original: las reservas descartadas. Como se ha definido anteriormente, este evento se produce cuando una reserva alcanza su instante de transmisión sin que exista todavía un paquete disponible para transmitir. Este comportamiento se representa esquemáticamente en la Figura 5, donde se observa que, tras transmitir el primer paquete en T_{tx} , la reserva realizada en $T_{tx} + RRI$ no puede utilizarse porque el siguiente paquete aún no ha sido generado, por lo que la transmisión se retrasa hasta una reserva posterior.

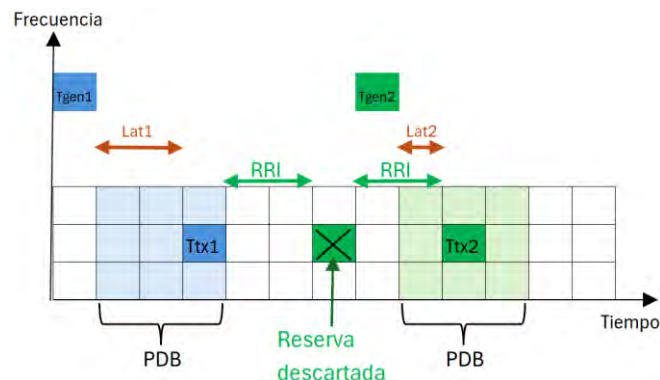


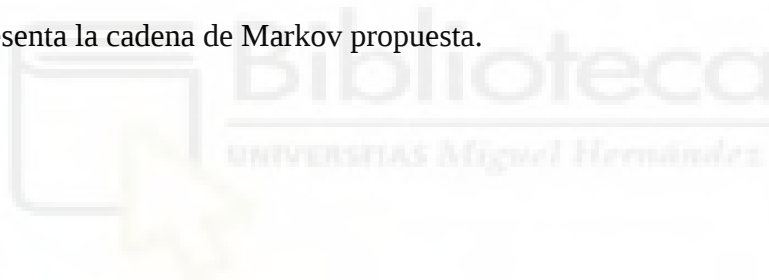
Figura 5. Esquema de reservas descartadas.

Por tanto, para modelar correctamente el comportamiento del SPS en tráfico aperiódico, no basta con representar únicamente las reselecciones por contador y por tamaño. Es

necesario incorporar también las reselecciones por latencia, las reselecciones simultáneas por tamaño y latencia, y las reservas descartadas por ausencia de paquete generado. Las reservas no utilizadas por reselección quedan representadas a través de los propios estados de reselección, ya que cada reselección implica abandonar una reserva previamente existente para seleccionar nuevos recursos.

3.1.2. Definición de la cadena de Markov y probabilidades asociadas

Con estas bases, se ha diseñado una cadena de Markov para modelar el proceso SPS y extraer información sobre la reserva, el uso, el descarte y la reselección de recursos. El modelo es válido para escenarios sin memoria, donde tanto el tamaño de cada paquete como el intervalo de generación respecto al paquete anterior son independientes del tamaño y del intervalo del paquete generado previamente. Esta hipótesis permite representar la evolución del mecanismo mediante una cadena de Markov, ya que las probabilidades de transición dependen únicamente del estado actual del sistema. La Figura 6 representa la cadena de Markov propuesta.



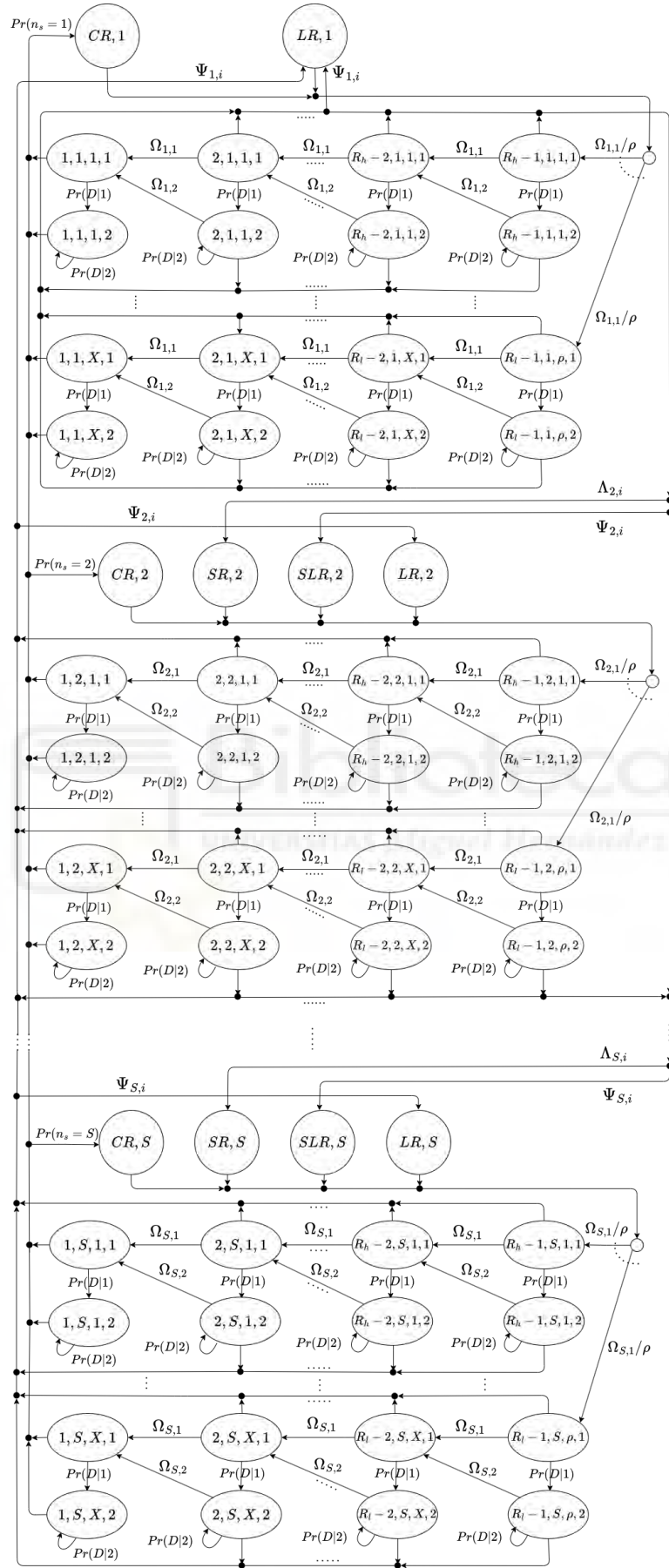


Figura 6. Cadena de Markov para el Modelo Analítico de gestión de recursos.

La cadena de Markov propuesta tiene tres tipos diferentes de estados, que representan el estado interno del mecanismo SPS: estados de reserva utilizada, estados de reserva descartada y estados de reelección. Los estados de reserva, tanto utilizada como descartada, se alcanzan cuando el SPS ya ha realizado una reserva de recursos. Se llega a un estado de reserva utilizada cuando se transmite un paquete empleando el recurso previamente reservado. En cambio, una transición hacia un estado de reserva descartada ocurre cuando el recurso previamente reservado no puede emplearse debido a la llegada tardía del siguiente paquete desde las capas superiores.

Todos los estados de reserva en la cadena de Markov se caracterizan mediante la tupla de cuatro elementos $\langle c, s, k, i \rangle$. En esta tupla, c denota el valor del Reselection Counter, mientras que s representa el número de subcanales actualmente reservados en el esquema SPS. El componente k es un número entero asociado al valor inicial del Reselection Counter, que refleja el valor aleatorio inicial que determina el estado seleccionado. Este parámetro oscila entre 1 y:

$$\rho = R_h - R_l + 1 \quad (2)$$

El valor inicial del Reselection Counter se selecciona de forma aleatoria uniforme dentro de un intervalo definido por el estándar y dependiente del RRI considerado. En NR V2X Modo 2, si $RRI \geq 100$ ms, dicho contador se selecciona dentro del intervalo $[5, 15]$. En cambio, si $RRI < 100$ ms, se selecciona dentro del intervalo $[5C, 15C]$, donde $C = \frac{100}{\max(20, RRI)}$ [5]. En el modelo, se denotan como R_l y R_h los límites inferior y superior enteros del conjunto de valores posibles del contador. El índice k permite identificar cada posible valor inicial del contador dentro de este intervalo, de forma que $k = 1$ corresponde al primer valor posible y $k = \rho$ al último. Cuando se alcanza un estado de reelección, el vehículo selecciona una nueva reserva y se reinicializa el contador tomando de nuevo uno de estos valores posibles.

Finalmente, i identifica si una reserva determinada se utiliza o se descarta. En concreto, $i = 1$ corresponde a un estado de reserva utilizada, mientras que $i = 2$ corresponde a un estado de reserva descartada por ausencia de paquete generado.

Los estados de reelección se alcanzan cuando se produce una reelección de recursos. Estos estados se representan mediante la tupla de dos elementos: $\langle Res, s \rangle$, donde Res indica el tipo de reelección y s denota el nuevo número de subcanales requeridos para el próximo paquete. Los tipos de reelección soportados por la cadena de Markov propuesta

son: reelección por contador (CR), reelección por tamaño (SR), reelección por latencia (LR) y reelección simultánea por tamaño y latencia (SLR). Por ejemplo, el estado $\langle CR, 2 \rangle$ se alcanza cuando ocurre una reelección por contador y el siguiente paquete requiere dos subcanales. Desde un estado de reelección, el vehículo selecciona una nueva reserva capaz de alojar un paquete que requiere s subcanales. Además, se reinicializa el Reselection Counter seleccionando de forma uniforme uno de los valores posibles dentro del intervalo $[R_l, R_h]$. Por tanto, la transición desde un estado de reelección hacia los estados de reserva se reparte entre las filas asociadas a los distintos valores iniciales del contador, representados mediante el índice k . A partir de ese instante, el proceso continúa de nuevo como una reserva activa dentro del mecanismo SPS.

La cadena de Markov está organizada en S bloques de estados, donde S es el número total de subcanales en el ancho de banda considerado. El primer bloque corresponde al caso en que la reserva incluye solo un subcanal, el segundo bloque corresponde al caso con dos subcanales y, en general, el bloque s -ésimo corresponde a s subcanales reservados. Cada bloque contiene $R_h - R_l + 1$ pares de filas de estados de reserva, más una fila adicional con los correspondientes estados de reelección. Cada par de filas de estados de reserva se identifica mediante el componente k de la tupla $\langle c, s, k, i \rangle$. Para cada uno de estos pares, los estados de la fila superior, $\langle c, s, k, 1 \rangle$, son estados de reserva utilizada, mientras que los estados de la fila inferior, $\langle c, s, k, 2 \rangle$, representan estados de reserva descartada. Por simplicidad, y dadas las limitaciones de representación de la figura, la Figura 6 muestra únicamente algunos bloques de estados y el primer y último par de filas de cada bloque, aunque la estructura se repite para todos los valores posibles de s y k .

La cadena de Markov propuesta diferencia tres tipos principales de transiciones entre estados. El primer tipo de transición se produce cuando se utiliza el recurso previamente reservado y, por tanto, el estado de destino es un estado de reserva utilizada, situado en la fila superior de cualquiera de los pares de filas. Este tipo de transición ocurre cuando el siguiente paquete se genera a tiempo para ser transmitido en la reserva actual, su latencia es inferior al PDB , el paquete cabe en los recursos reservados y el Reselection Counter es mayor que uno, es decir, $c > 1$. En este primer tipo de transición son posibles dos subcasos. En el primero, la transición correspondiente en la cadena de Markov va del estado $\langle c, s, k, 1 \rangle$ al estado $\langle c - 1, s, k, 1 \rangle$. Ambos estados se encuentran en la fila superior, ya que en los dos casos se trata de reservas utilizadas. En el segundo subcaso, la transición se produce desde el estado $\langle c, s, k, 2 \rangle$ al estado $\langle c - 1, s, k, 1 \rangle$, es decir, desde

la fila inferior hacia la fila superior. En ambos casos, el estado de destino es una reserva utilizada y el Reselection Counter se decrementa en una unidad.

La probabilidad de que ocurra esta transición es igual a la probabilidad de que no haya reelección por tamaño ni por latencia, y de que además se utilice la reserva. Esta probabilidad se expresa como:

$$\Omega_{s,i} = (1 - Pr(SR | s))(1 - Pr(LR | i))(1 - Pr(D | i)) \quad (3)$$

En esta ecuación, $Pr(SR | s)$ es la probabilidad de reelección por tamaño dado que la reserva actual tiene s subcanales. Puede calcularse como la probabilidad de que el próximo paquete requiera más de s subcanales:

$$Pr(SR|s) = \sum_{j=s+1}^s p_r(j) \quad (4)$$

La probabilidad de reelección por latencia y la probabilidad de que el recurso reservado actual sea descartado se denotan como $Pr(LR | i)$ y $Pr(D | i)$, respectivamente. Ambas dependen del proceso temporal de generación de paquetes y se calculan posteriormente. El parámetro i permite diferenciar si el estado de origen es una reserva utilizada, $i = 1$, o una reserva descartada, $i = 2$. El segundo tipo de transición se produce cuando el recurso reservado actualmente se descarta y no hay reelección de recursos. Esto ocurre cuando el siguiente paquete es generado por las capas superiores después de la reserva actual, pero aún podría transmitirse en la siguiente oportunidad, tras un periodo adicional de RRI . En este caso, el estado de destino es un estado de reserva descartada, situado en la fila inferior de cualquiera de los pares de filas.

Dentro de este segundo tipo de transición también pueden distinguirse dos subcasos. En el primero, la transición ocurre desde el estado $\langle c, s, k, 1 \rangle$ al estado $\langle c, s, k, 2 \rangle$, es decir, desde la fila superior hacia la fila inferior, con probabilidad $Pr(D | 1)$. En el segundo, el estado no cambia, por lo que el estado de origen y el estado de destino son ambos $\langle c, s, k, 2 \rangle$, con una probabilidad de transición igual a $Pr(D | 2)$. Este último caso representa la posibilidad de tener reservas descartadas consecutivas.

El tercer tipo de transición en la cadena de Markov se produce cuando se dispara una reelección de recursos. Esto ocurre cuando el tamaño del siguiente paquete no cabe en

la reserva actual, cuando la latencia supera el *PDB* y/o cuando el *Reselection Counter* alcanza su valor final. Si el contador es igual a uno, es decir, si el estado de origen es de la forma $\langle 1, s, k, i \rangle$, se transita al estado $\langle CR, j \rangle$ con probabilidad $p_r(j)$, ya que el siguiente recurso no se mantiene y, por tanto, no existe posibilidad de conservar la reserva anterior. En cualquier otro estado de reserva, ya sea utilizada o descartada, existe cierta probabilidad de reelección por tamaño y/o por latencia. La única excepción se produce en los estados $\langle c, S, k, i \rangle$, donde la probabilidad de reelección por tamaño es cero porque la reserva ya cuenta con todos los subcanales disponibles, S , y por tanto ningún paquete puede requerir más subcanales que los actualmente reservados. Las probabilidades de transición a los estados $\langle SR, s \rangle$, que implican una reelección por tamaño pero no por latencia, se calculan como:

$$\Lambda_{s,i} = Pr(n_s = s)(1 - Pr(LR | i))(1 - Pr(D | i)) \quad (5)$$

En esta ecuación, $Pr(n_s = s) = p_r(s)$, donde s es el número de subcanales requeridos por el nuevo paquete. En este caso, dicho número de subcanales debe ser mayor que el número de subcanales reservados actualmente, ya que de lo contrario no existiría reelección por tamaño. El parámetro i diferencia el estado de origen, siendo $i = 1$ para estados de reserva utilizada e $i = 2$ para estados de reserva descartada.

Las probabilidades de transición a estados de reelección que implican una reelección por latencia se calculan como:

$$\Psi_{s,i} = Pr(n_s = s) \cdot Pr(LR|i) \cdot (1 - Pr(D|i)) \quad (6)$$

En este caso, cuando s es menor o igual que el número de subcanales reservados actualmente, no hay reelección por tamaño y el estado de destino es $\langle LR, s \rangle$. Sin embargo, cuando s es mayor que el número de subcanales reservados actualmente, también se produce una reelección por tamaño y el estado de destino es $\langle SLR, s \rangle$. De este modo, la cadena distingue entre las reelecciones debidas únicamente a latencia y aquellas en las que se combinan simultáneamente la latencia y el aumento del tamaño del paquete. La cadena de Markov propuesta se utiliza para calcular la probabilidad de que se active una reelección de recursos, denotada como p_{rsl} . Esta probabilidad es necesaria para determinar la frecuencia con la que el vehículo abandona los recursos previamente

reservados y selecciona una nueva reserva. Dicha probabilidad es igual a la probabilidad de que el proceso de Markov se encuentre en alguno de los estados de reelección y, por tanto, puede calcularse como:

$$p_{rsl} = \sum_{s=1}^S Pr(\langle CR, s \rangle) + \sum_{s=1}^S Pr(\langle LR, s \rangle) + \sum_{s=2}^S Pr(\langle SR, s \rangle) + \sum_{s=2}^S Pr(\langle SLR, s \rangle) \quad (7)$$

En esta ecuación, el primer sumatorio representa la probabilidad de que se active una reelección de recursos debido al agotamiento del *Reselection Counter*. El segundo sumatorio representa la probabilidad de que una reserva exceda la latencia máxima permitida. El tercer sumatorio representa la probabilidad de que se active una reelección por tamaño, y el cuarto sumatorio representa la probabilidad de que se activen simultáneamente las reelecciones por tamaño y latencia.

También se utiliza la cadena de Markov propuesta para calcular la probabilidad de que se reserve un número determinado de subcanales, $p_R(s)$. Esta probabilidad es necesaria para caracterizar la ocupación de recursos del sistema. Para $s = 1$, esta probabilidad se calcula como:

$$p_R(1) = Pr(\langle CR, 1 \rangle) + Pr(\langle LR, 1 \rangle) + Pr(\langle \cdot, 1, \cdot, \cdot \rangle) \quad (8)$$

donde $Pr(\langle \cdot, 1, \cdot, \cdot \rangle)$ representa la suma de las probabilidades estacionarias de todos los estados $\langle c, 1, k, i \rangle$ para todos los valores válidos de c , k e i . Del mismo modo, cuando el número de subcanales es $s > 1$, la probabilidad se calcula como:

$$p_R(s) = Pr(\langle CR, s \rangle) + Pr(\langle LR, s \rangle) + Pr(\langle SR, s \rangle) + Pr(\langle SLR, s \rangle) + Pr(\langle \cdot, s, \cdot, \cdot \rangle), s > 1 \quad (9)$$

Por último, también se utiliza la cadena de Markov propuesta para calcular la probabilidad de que una reserva sea descartada por ausencia de paquete generado. Esta probabilidad se denota como p_d y se obtiene como la suma de las probabilidades estacionarias de todos los estados $\langle c, s, k, i \rangle$ para todos los valores válidos de c , s y k , cuando $i = 2$:

$$p_d = Pr(\langle \cdot, \cdot, \cdot, 2 \rangle) \quad (10)$$

p_d no representa la probabilidad total de reserva no utilizada, sino la probabilidad de reserva descartada por ausencia de paquete generado. Las reservas no utilizadas por reelección quedan representadas mediante los estados de reelección de la cadena, ya que en esos casos el vehículo abandona una reserva previa y selecciona nuevos recursos. Si se desea agrupar ambos efectos en una única magnitud auxiliar, puede definirse la probabilidad de no utilizar una reserva, p_{ur} , como la probabilidad de estar en un estado de reelección o en un estado de reserva descartada ($i=2$):

$$p_{ur} = p_{rst} + p_d \quad (11)$$

Para resolver las expresiones anteriores, es necesario calcular la distribución estacionaria, π , es decir, el vector con todas las probabilidades de que el proceso de Markov se encuentre en cada estado. Para ello, se hace uso de la cadena de Markov propuesta y de su matriz de transición, que describe todas las probabilidades de transición entre dos estados cualesquiera. Si se denota dicha matriz como M , la distribución estacionaria puede calcularse resolviendo el siguiente sistema lineal:

$$\pi = M^T \pi \quad (12)$$

junto con la condición de normalización:

$$\sum_{q \in \mathcal{X}} \pi_q = 1, \quad \pi_q \geq 0 \quad (13)$$

donde $\mathcal{X}$ representa el conjunto de estados de la cadena de Markov.

3.1.3. Cálculo de probabilidades de descarte y reelección por latencia

Una vez definida la estructura de la cadena de Markov y el papel de las probabilidades $\Pr(D \mid i)$ y $\Pr(LR \mid i)$, se detalla a continuación cómo se calculan dichas probabilidades para los modelos de tráfico aperiódico. Estas probabilidades permiten representar dos eventos propios de este tipo de tráfico: el descarte de una reserva por ausencia de paquete generado y la reelección por latencia cuando la reserva disponible no permite cumplir el PDB . En los modelos periódicos considerados, donde la generación de paquetes coincide

con la periodicidad de las reservas, estas probabilidades se anulan, ya que no aparecen reservas descartadas por ausencia de paquete ni reselecciones por latencia debidas a generación aperiódica.

En un tratamiento exacto, estas probabilidades dependerían de toda la historia temporal del proceso. Es decir, sería necesario conocer cuántas oportunidades de reserva han transcurrido desde la última transmisión efectiva y cuántas reservas han sido descartadas consecutivamente. Sin embargo, incluir esta memoria temporal en la cadena no es trivial debido al elevado número de estados que serían necesarios. Para mantener el modelo en una forma tratable, se adopta una aproximación basada en conservar únicamente la información temporal más relevante: si el estado anterior corresponde a una reserva utilizada o a una reserva descartada. De este modo, todas las transiciones desde estados de reserva utilizada emplean las probabilidades $\Pr(D | 1)$ y $\Pr(LR | 1)$, mientras que las transiciones desde estados de reserva descartada emplean $\Pr(D | 2)$ y $\Pr(LR | 2)$.

Para el cálculo de estas probabilidades, se denota Δt_L a la latencia del último paquete transmitido, definida como la diferencia entre su instante de transmisión y su instante de generación:

$$\Delta t_L(j) = t_{tx}(j) - t_{gen}(j) \quad (14)$$

Tras una reselección de recursos, el instante de transmisión se selecciona dentro de la ventana temporal permitida. Por ello, la latencia del paquete transmitido tras una reselección se modela como una variable aleatoria uniforme entre la latencia mínima considerada por el sistema y el PDB :

$$\Delta t_L \sim U(t_{\min}, PDB) \quad (15)$$

Por otro lado, llamamos Δt_G al intervalo de generación entre dos paquetes consecutivos:

$$\Delta t_G = t_{gen}(j+1) - t_{gen}(j) \quad (16)$$

En los modelos aperiódicos considerados, este intervalo se modela mediante una distribución exponencial desplazada:

$$\Delta t_G = K + Z \quad (17)$$

donde K es la componente constante del modelo de tráfico y Z es una variable aleatoria exponencial de media K . Por tanto, la función de distribución acumulada de Δt_G , denotada como $F_G(t)$, viene dada por:

$$F_G(t) = \begin{cases} 0, & t < K \\ 1 - e^{-\frac{t-K}{K}}, & t \geq K \end{cases} \quad (18)$$

Esta función permite calcular directamente la probabilidad de que el siguiente paquete se haya generado antes de un determinado instante temporal.

Con esto definido, podemos calcular las probabilidades. En primer lugar, se analiza el caso en el que el estado de origen corresponde a una reserva utilizada, es decir, $i = 1$. Después de una transmisión efectiva, la siguiente oportunidad de reserva se sitúa un intervalo RRI después del instante de transmisión anterior. Si se toma como referencia el instante de generación del último paquete transmitido, dicha reserva aparece en: $\Delta t_L + RRI$. La reserva será descartada si el siguiente paquete se genera después de esa oportunidad reservada. En ese caso, cuando llega el instante de la reserva, todavía no existe ningún paquete disponible para transmitir. Por tanto, el evento de descarte tras una reserva utilizada se define como:

$$D_1: \Delta t_G > \Delta t_L + RRI \quad (19)$$

Para obtener la probabilidad de este evento, debe calcularse la probabilidad de que el intervalo de generación Δt_G sea mayor que el instante de la reserva, que depende a su vez de la latencia anterior Δt_L . Matemáticamente, cuando se quiere calcular la probabilidad de que una variable aleatoria sea mayor que otra, se integra la densidad conjunta en la región donde se cumple dicha condición. Es decir, si x es un posible valor de Δt_L , y g un posible valor de Δt_G , el descarte se produce cuando:

$$g > x + RRI \quad (20)$$

Por tanto, la probabilidad de descarte puede escribirse inicialmente como:

$$\Pr(D_1) = \int_{t_{\min}}^{PDB} \int_{x+RRI}^{\infty} f_{\Delta t_L}(x) \cdot f_{\Delta t_G}(g) dg dx \quad (21)$$

donde $f_{\Delta t_L}(x)$ es la densidad de la latencia anterior y $f_{\Delta t_G}(g)$ es la densidad del intervalo de generación entre paquetes. Como Δt_L se modela como una variable uniforme en el intervalo $[t_{\min}, PDB]$, su función de densidad es:

$$f_{\Delta t_L}(x) = \frac{1}{PDB - t_{\min}} \quad (22)$$

Sustituyendo:

$$\Pr(D_1) = \frac{1}{PDB - t_{\min}} \int_{t_{\min}}^{PDB} \int_{x+RRI}^{\infty} f_{\Delta t_G}(g) dg dx \quad (23)$$

La integral interior representa la probabilidad de que Δt_G sea mayor que $(x + RRI)$. Usando la función de distribución acumulada de Δt_G (18), dicha probabilidad puede escribirse como:

$$\int_{x+RRI}^{\infty} f_{\Delta t_G}(g) dg = 1 - F_G(x + RRI) \quad (24)$$

Sustituyendo este resultado en la expresión anterior, se obtiene:

$$\Pr(D | 1) = \frac{1}{PDB - t_{\min}} \int_{t_{\min}}^{PDB} [1 - F_G(x + RRI)] dx \quad (25)$$

donde $1 - F_G(x + RRI)$ representa la probabilidad de que el siguiente paquete todavía no se haya generado en el instante de la reserva, para un valor concreto de la latencia anterior x . La integral promedia esta probabilidad sobre todos los posibles valores de Δt_L .

Además del descarte, puede producirse una reelección por latencia. Este evento ocurre cuando el siguiente paquete sí se genera antes de la reserva, pero esperar hasta dicha reserva provocaría una latencia superior al PDB . Si el paquete se transmitiera en la reserva situada en $\Delta t_L + RRI$, su latencia sería:

$$\Delta t_L + RRI - \Delta t_G \quad (26)$$

Por tanto, se produce una reelección por latencia cuando:

$$\Delta t_L + RRI - \Delta t_G > PDB \quad (27)$$

equivalentemente:

$$\Delta t_G < \Delta t_L + RRI - PDB \quad (28)$$

De forma análoga a lo hecho con el evento de descarte, el evento de reelección por latencia asociado a la primera reserva tras una transmisión efectiva se define entonces como:

$$LR_1: \Delta t_G < \Delta t_L + RRI - PDB \quad (29)$$

Aplicando el mismo razonamiento anterior, la probabilidad de este evento se obtiene integrando en la región donde $g < x + RRI - PDB$. Es decir:

$$\Pr(LR_1) = \int_{t_{min}}^{PDB} \int_{-\infty}^{x+RRI-PDB} f_{\Delta t_L}(x) \cdot f_{\Delta t_G}(g) dg dx \quad (30)$$

Como Δt_L es uniforme, queda:

$$\Pr(LR_1) = \frac{1}{PDB - t_{min}} \int_{t_{min}}^{PDB} \int_{-\infty}^{x+RRI-PDB} f_{\Delta t_G}(g) dg dx \quad (31)$$

La integral interior es ahora la función de distribución acumulada de Δt_G evaluada en $x + RRI - PDB$:

$$\int_{-\infty}^{x+RRI-PDB} f_{\Delta t_G}(g) dg = F_G(x + RRI - PDB) \quad (32)$$

Por tanto, la probabilidad efectiva de reelección por latencia en esta primera reserva es:

$$\Pr(LR_1) = \frac{1}{PDB - t_{\min}} \int_{t_{\min}}^{PDB} F_G(x + RRI - PDB) dx \quad (33)$$

En esta expresión, $F_G(x + RRI - PDB)$ representa la probabilidad de que el paquete se genere suficientemente pronto como para que esperar hasta la reserva supere el PDB .

Los eventos de descarte y reelección por latencia son excluyentes. Si una reserva se descarta, significa que el paquete todavía no ha sido generado en el instante reservado, por lo que no puede evaluarse una reelección por latencia sobre esa misma oportunidad. Por ello, la matriz de transición de la cadena de Markov separa ambos eventos: primero considera si la reserva se descarta y, solo en caso de que no se descarte, evalúa si se produce una reelección por latencia. Esta separación aparece directamente en las transiciones de la cadena, donde las probabilidades de reelección por latencia se multiplican por el término $1 - \Pr(D | i)$. Por tanto, la probabilidad $\Pr(LR | i)$ utilizada en la matriz no representa la probabilidad efectiva total de reelección por latencia, sino la probabilidad de reelección por latencia dentro de los casos en los que la reserva no se ha descartado. Así, para los estados procedentes de una reserva utilizada, la probabilidad de reelección por latencia empleada en la matriz se define como:

$$\Pr(LR | 1) = \frac{\Pr(LR_1)}{1 - \Pr(D | 1)} \quad (34)$$

A continuación, se analiza el caso en el que el estado de origen corresponde a una reserva descartada, es decir, $i = 2$. En este caso, el proceso no se reinicia completamente, ya que se sabe que la reserva anterior no pudo utilizarse porque el paquete todavía no había sido generado. Por tanto, las probabilidades deben calcularse condicionadas a que el primer descarte ya ha ocurrido.

La siguiente oportunidad de reserva tras un primer descarte, medida desde la última transmisión efectiva, se sitúa en: $\Delta t_L + 2RRI$. La reserva volverá a descartarse si el paquete sigue sin haberse generado antes de esta segunda oportunidad. El evento de segundo descarte se define como:

$$D_2: \Delta t_G > \Delta t_L + 2RRI \quad (35)$$

Siguiendo el mismo principio utilizado para D_1 , la probabilidad de que el paquete no haya llegado antes de la segunda oportunidad reservada es:

$$\frac{1}{PDB - t_{\min}} \int_{t_{\min}}^{PDB} [1 - F_G(x + 2RRI)] dx \quad (36)$$

Sin embargo, para llegar a un estado de reserva descartada ya se sabe que la primera reserva fue descartada. Por tanto, la probabilidad de nuevo descarte debe calcularse condicionada a dicho primer descarte. Así, se obtiene:

$$\Pr(D | 2) = \frac{\frac{1}{PDB - t_{\min}} \int_{t_{\min}}^{PDB} [1 - F_G(x + 2RRI)] dx}{\Pr(D | 1)} \quad (37)$$

El numerador representa la probabilidad de que el paquete no haya llegado antes de la segunda oportunidad reservada. El denominador introduce la condición de que la primera reserva ya había sido descartada. De este modo, $\Pr(D | 2)$ representa la probabilidad de descartar de nuevo una reserva sabiendo que la reserva anterior ya fue descartada. También puede ocurrir que, tras descartar la primera reserva, el paquete llegue antes de la segunda oportunidad reservada, pero que transmitirlo en dicha segunda oportunidad incumpla el PDB . Para que esto ocurra deben cumplirse dos condiciones: en primer lugar, el paquete debe generarse después de la primera reserva, ya que esta fue descartada:

$$\Delta t_G > \Delta t_L + RRI \quad (38)$$

En segundo lugar, debe generarse suficientemente pronto como para que esperar hasta la segunda reserva produzca una latencia superior al PDB :

$$\Delta t_L + 2RRI - \Delta t_G > PDB \quad (39)$$

lo que equivale a:

$$\Delta t_G < \Delta t_L + 2RRI - PDB \quad (40)$$

Por tanto, el evento de reelección por latencia tras una reserva descartada queda definido por el intervalo:

$$\Delta t_L + RRI < \Delta t_G < \Delta t_L + 2RRI - PDB \quad (41)$$

La probabilidad de este evento se obtiene integrando la densidad de Δt_G entre esos dos límites. Para un valor concreto $\Delta t_L = x$, dicha probabilidad es:

$$F_G(x + 2RRI - PDB) - F_G(x + RRI) \quad (42)$$

Por tanto, promediando sobre todos los valores posibles de Δt_L y condicionando a que la primera reserva ya fue descartada, se obtiene:

$$\Pr(LR_2 | D_1) = \frac{\frac{1}{PDB - t_{\min}} \int_{t_{\min}}^{PDB} [F_G(x + 2RRI - PDB) - F_G(x + RRI)] dx}{\Pr(D | 1)} \quad (43)$$

El término:

$$F_G(x + 2RRI - PDB) - F_G(x + RRI) \quad (44)$$

representa la probabilidad de que el paquete se genere después de la primera reserva, pero suficientemente pronto como para que esperar hasta la segunda reserva supere el PDB .

De forma análoga al caso anterior, la matriz de transición separa primero el evento de descarte y después el evento de reelección por latencia. Por ello, la probabilidad de reelección por latencia desde un estado de reserva descartada se define dentro de los casos en los que la nueva reserva evaluada no vuelve a descartarse:

$$\Pr(LR | 2) = \frac{\Pr(LR_2 | D_1)}{1 - \Pr(D | 2)} \quad (45)$$

Con esta formulación, las probabilidades de transición quedan normalizadas y son coherentes con la estructura de la cadena de Markov. Desde un estado de reserva utilizada se emplean $\Pr(D | 1)$ y $\Pr(LR | 1)$, mientras que desde un estado de reserva descartada se emplean $\Pr(D | 2)$ y $\Pr(LR | 2)$. La cadena no almacena el número exacto de descartes consecutivos desde la última transmisión, sino únicamente si la reserva anterior fue

utilizada o descartada. Esta aproximación permite capturar el efecto principal del tráfico aperiódico, incluyendo descartes consecutivos y reselecciones por latencia, sin aumentar excesivamente la dimensión de la matriz de transición.

Con las probabilidades de reselección por contador, reselección por tamaño, reselección por latencia y descarte ya definidas, la cadena de Markov queda completamente especificada. A partir de su matriz de transición se obtiene la distribución estacionaria y, con ella, las magnitudes principales del modelo: la probabilidad global de reselección de recursos, la probabilidad de reservar un determinado número de subcanales y la probabilidad de que una reserva sea descartada por ausencia de paquete generado. Estas probabilidades constituyen la base del análisis posterior, ya que permiten caracterizar de forma analítica el comportamiento del mecanismo SPS en presencia de tráfico periódico y aperiódico, y sirven como entrada para el estudio de métricas como la latencia y las pérdidas consecutivas entre recepciones correctas.

3.2. Modelo de distribución de la latencia

Una vez obtenidas mediante la cadena de Markov las probabilidades asociadas al proceso de reserva y reselección de recursos, es posible caracterizar analíticamente la distribución de la latencia experimentada por los paquetes transmitidos. Dichas probabilidades no representan por sí solas la distribución de la latencia, sino que deben combinarse con el mecanismo temporal de generación de paquetes y con la forma en que se utilizan las reservas de recursos.

El propósito de este apartado es obtener una expresión analítica que permita representar la distribución de la latencia, la cual se obtiene a partir de los mecanismos principales del modelo: la generación de paquetes, la periodicidad de las reservas, la reselección de recursos y la aparición de reservas descartadas por ausencia de paquete generado.

La latencia de un paquete se define como la diferencia entre el instante en el que dicho paquete se transmite y el instante en el que se genera:

$$L = T_{tx} - T_{gen} \quad (46)$$

donde T_{gen} representa el instante de generación del paquete, T_{tx} el instante de transmisión y L la latencia resultante. A lo largo del desarrollo se utilizará l para denotar un valor concreto de latencia sobre el que se evalúa la función de densidad de probabilidad. Es

decir, L es la variable aleatoria, mientras que l es la variable continua del eje horizontal de la distribución.

La latencia estará definida entre una latencia mínima, $t_{\min}$, y una latencia máxima dada por el Packet Delay Budget, PDB . Por tanto, de forma general:

$$t_{\min} \leq L \leq PDB \quad (47)$$

El comportamiento de la latencia depende de la forma en que se transmite cada paquete. En el modelo considerado existen dos situaciones fundamentales: la primera ocurre cuando se produce una reelección de recursos, la segunda ocurre cuando no hay reelección y el paquete utiliza una reserva periódica previamente establecida. La distribución final de la latencia se obtiene, por tanto, como una combinación de ambas contribuciones.

Cuando se produce una reelección, el sistema selecciona un nuevo recurso dentro de la ventana de latencia permitida. Esto equivale a fijar el instante de transmisión como:

$$T_{tx} = T_{gen} + U \quad (48)$$

donde U es una variable aleatoria uniforme en el intervalo $[t_{\min}, PDB]$. Por tanto, en caso de reelección:

$$L = T_{tx} - T_{gen} = U \quad (49)$$

De esta forma, la latencia condicionada a que haya reelección sigue una distribución uniforme, $U(t_{\min}, PDB)$. La función de densidad de probabilidad asociada a estos paquetes es:

$$f_{rst}(l) = \frac{1}{PDB - t_{\min}}, t_{\min} \leq l \leq PDB \quad (50)$$

Este resultado es independiente del motivo que haya provocado la reelección. En el sistema considerado puede haber reelección por agotamiento del contador, por tamaño de paquete, por latencia o por una combinación simultánea de tamaño y latencia. Sin embargo, una vez que se produce la reelección, el comportamiento temporal es el mismo: se abandona la reserva anterior y se selecciona una nueva reserva dentro de la ventana permitida. Por ello, todos los tipos de reelección aportan la misma forma de distribución

de latencia: una distribución uniforme en $[t_{\min}, PDB]$. Esta probabilidad de reelección, p_{rsl} , se obtiene directamente de forma analítica del modelo explicado en el apartado anterior.

3.2.1. Casos con distribución uniforme de latencia

En los modelos periódicos, el intervalo de generación de paquetes coincide con la periodicidad de las reservas, de forma que el desfase temporal entre la generación y la transmisión de los paquetes se mantiene constante mientras no se produzca una reelección de recursos. Cuando se transmite el primer paquete tras una reelección, se escoge una latencia aleatoria dentro de la ventana permitida. A partir de ese instante, si los paquetes se generan periódicamente y las reservas también se mantienen periódicas con el mismo intervalo, todos los paquetes siguientes conservan la misma latencia. Por tanto, cada ciclo entre reelecciones tiene una latencia constante, pero dicha latencia se escoge aleatoriamente al comienzo del ciclo. Como el valor inicial se selecciona uniformemente en la ventana $[t_{\min}, PDB]$, y las reelecciones vuelven a seleccionar un nuevo valor uniforme en esa misma ventana, la distribución global de la latencia sobre muchos paquetes es uniforme, siguiendo la ecuación (50).

Este resultado no depende de cuántas reelecciones se produzcan, siempre que el valor seleccionado tras cada reelección sea uniforme dentro de la ventana permitida. Las reelecciones modifican cuándo cambia la latencia, pero no modifican la distribución global de los valores que puede tomar.

En los modelos aperiódicos, el intervalo de generación de paquetes ya no es constante. En particular, tomando como base los modelos de tráfico definidos por 3GPP, se considera un modelo de generación con tiempo de generación entre paquetes, T , de la forma:

$$T = K + E \quad (51)$$

donde E representa una variable aleatoria exponencial de media K , $E \sim \text{Exp}(K)$. Esto significa que entre dos paquetes consecutivos transcurre siempre, como mínimo, un tiempo K , al que se suma una componente aleatoria exponencial. Este trabajo se centra en caracterizar dos configuraciones características, tanto cuando RRI es igual al mínimo del intervalo de generación, $RRI = K$, como cuando es igual a la media de este intervalo,

$RRI = 2 \cdot K$. En ambos casos se considera $K = PDB$. Además, por simplicidad, se asume $t_{\min} = 0$.

Empezando por los modelos aperiódicos donde RRI es igual al mínimo del intervalo de generación, supongamos que el paquete i se genera en $T_{gen,i}$ y se transmite con latencia L_i . Entonces:

$$T_{tx,i} = T_{gen,i} + L_i \quad (52)$$

Tras esa transmisión, la siguiente reserva queda programada RRI milisegundos después:

$$T_{res,i+1} = T_{tx,i} + RRI \quad (53)$$

Para el caso considerado $RRI = K$, quedaría:

$$T_{res,i+1} = T_{gen,i} + L_i + K \quad (54)$$

Por otro lado, el siguiente paquete se genera en:

$$T_{gen,i+1} = T_{gen,i} + K + E_i \quad (55)$$

donde $E_i \sim \text{Exp}(K)$. Si el paquete $i + 1$ pudiera transmitirse directamente en la reserva $T_{res,i+1}$, su latencia candidata sería:

$$L_{i+1} = T_{res,i+1} - T_{gen,i+1} \quad (56)$$

Sustituyendo las expresiones anteriores:

$$L_{i+1} = (T_{gen,i} + L_i + K) - (T_{gen,i} + K + E_i) = L_i - E_i \quad (57)$$

Esta expresión indica que, antes de considerar descartes, la nueva latencia sería la latencia anterior menos la componente exponencial adicional de la generación aperiódica. Si $L_i - E_i$ es positivo, la reserva sigue estando después de la generación del paquete y puede utilizarse. Sin embargo, si esta latencia es negativa, entonces la reserva ocurre antes de que el paquete se haya generado. En ese caso, dicha reserva no puede utilizarse y se

descarta. El sistema salta entonces a la siguiente reserva, situada K milisegundos después. Si esta nueva reserva también quedase antes de la generación, se descartaría de nuevo, y así sucesivamente.

Este proceso equivale a plegar la latencia dentro del intervalo $[0, K]$, lo que puede expresarse mediante una operación módulo K :

$$L_{i+1} = (L_i - E_i) \bmod K \quad (58)$$

La operación módulo representa el hecho de que, cada vez que una reserva ocurre antes de que el paquete se haya generado, se suma un nuevo intervalo K hasta que la transmisión vuelve a quedar dentro de la ventana válida de latencia.

Este comportamiento conserva la uniformidad de la distribución. Es decir, si L_i es uniforme en $[0, K]$, y E_i es independiente de L_i , entonces:

$$(L_i - E_i) \bmod K \sim U(0, K) \quad (59)$$

Esta propiedad se basa en la invariancia de la distribución uniforme módulo 1, la cual establece que, si una variable X está uniformemente distribuida módulo 1 y Y es una variable aleatoria independiente de X , entonces $X + Y$ también está uniformemente distribuida módulo 1 [25]. Esta propiedad puede entenderse como una extensión de la invariancia frente a traslaciones fijas de la distribución uniforme módulo 1 recogida en [26].

Aunque la propiedad anterior se formula para variables uniformemente distribuidas módulo 1, puede aplicarse al caso módulo K mediante una normalización. Para ello, se define:

$$X = \frac{L_i}{K} \quad (60)$$

De este modo, si $L_i \sim U(0, K)$, entonces $X \sim U(0,1)$.

Asimismo, definiendo:

$$Y = -\frac{E_i}{K} \quad (61)$$

se tiene que Y es independiente de X , ya que E_i es independiente de L_i . Por tanto, aplicando la propiedad de invariancia de la distribución uniforme módulo 1 frente a la suma de una variable aleatoria independiente, se obtiene:

$$(X + Y) \bmod 1 \sim U(0,1) \quad (62)$$

es decir:

$$\left(\frac{L_i}{K} - \frac{E_i}{K}\right) \bmod 1 \sim U(0,1) \quad (63)$$

Finalmente, al desnormalizar multiplicando por K , se obtiene:

$$(L_i - E_i) \bmod K \sim U(0, K) \quad (64)$$

que coincide con la expresión (59). Por tanto, en el caso $RRI = K = PDB$ y $t_{\min} = 0$, la rama sin reelección conserva una distribución uniforme:

$$f_{NR}(l) = \frac{1}{K}, \quad 0 \leq l \leq K \quad (65)$$

Además, la rama con reelección también es uniforme, ya que cuando se produce una reelección se selecciona un nuevo recurso dentro de la ventana de latencia permitida, siguiendo la misma expresión que (65), de modo que $f_{rs}(l) = f_{NR}(l)$. Al mezclar dos distribuciones uniformes definidas sobre el mismo intervalo, la distribución total de la latencia también es uniforme:

$$f_L(l) = \frac{1}{K}, \quad 0 \leq l \leq K \quad (66)$$

Por tanto, aunque el tráfico sea aperiódico, si $RRI = K = PDB$ y $t_{\min} = 0$, el efecto combinado de las reservas periódicas y de las reservas descartadas conserva la uniformidad de la latencia.

Este resultado no debe extenderse directamente a valores $t_{\min} > 0$. En ese caso, la ventana válida de latencia ya no coincide con el ciclo completo de duración K , por lo que el argumento basado en la invariancia de la distribución uniforme módulo K deja de ser

directamente aplicable. Por ello, la uniformidad obtenida en este caso está asociada específicamente a la hipótesis $t_{\min} = 0$, definiéndose así la latencia entre 0 y PDB.

3.2.2. Distribución de latencia en modelos aperiódicos con $RRI=2PDB$

El caso más relevante aparece en los modelos aperiódicos con $RRI = 2K$. En estos escenarios, la latencia deja de ser uniforme y aparece una componente creciente en la función de densidad. Esta forma se debe a la relación entre la generación aperiódica $K + E_i$ y las reservas periódicas separadas por $2K$. De nuevo, se considera el caso $t_{\min} = 0$ y $PDB = K$, de modo que la latencia válida pertenece al intervalo $0 \leq L \leq K$.

En primer lugar, se analiza la rama sin reelección, siguiendo un razonamiento similar al caso anterior. En ausencia de reelección, el paquete no selecciona una nueva reserva aleatoria, sino que se transmite en una reserva periódica ya existente.

Sea L_i la latencia del paquete transmitido previamente. Tras una transmisión, la siguiente reserva se programa un intervalo $RRI=2K$ después:

$$T_{res,i+1} = T_{tx,i} + 2K \quad (67)$$

Como:

$$T_{tx,i} = T_{gen,i} + L_i \quad (68)$$

se obtiene:

$$T_{res,i+1} = T_{gen,i} + L_i + 2K \quad (69)$$

Por otro lado, el siguiente paquete se genera en:

$$T_{gen,i+1} = T_{gen,i} + K + E_i \quad (70)$$

donde $E_i \sim \text{Exp}(K)$ representa la componente exponencial adicional del intervalo de generación aperiódico. Si este paquete se transmitiera en la reserva mantenida, su latencia sería:

$$L_{i+1} = T_{res,i+1} - T_{gen,i+1} \quad (71)$$

Sustituyendo las expresiones anteriores:

$$L_{i+1} = (T_{\text{gen},i} + L_i + 2K) - (T_{\text{gen},i} + K + E_i) = L_i + K - E_i \quad (72)$$

Para que esta transmisión pueda realizarse sin reelección y utilizando la reserva periódica mantenida, la latencia resultante debe quedar dentro de la ventana permitida. Como se está considerando $t_{\min} = 0$ y $PDB = K$, dicha condición viene dada por:

$$0 \leq L_{i+1} \leq K \quad (73)$$

Sustituyendo la expresión anterior:

$$0 \leq L_i + K - E_i \leq K \quad (74)$$

lo que equivale a:

$$L_i \leq E_i \leq L_i + K \quad (75)$$

Se puede razonar que, si $E_i < L_i$, el paquete se genera demasiado pronto respecto a la reserva mantenida, por lo que esperar hasta dicha reserva produciría una latencia superior al PDB . En ese caso, sería necesaria una reelección por latencia. En cambio, si $E_i > L_i + K$, la reserva considerada ocurre antes de que el paquete se haya generado, por lo que dicha reserva se descarta por ausencia de paquete. Finalmente, cuando se cumple (75), el paquete puede transmitirse en la reserva periódica mantenida sin necesidad de reeleccionar.

Para obtener la distribución de la latencia en esta rama, se parte de la relación (72). Así, para obtener una latencia concreta l , debe cumplirse:

$$E_i = L_i + K - l \quad (76)$$

La variable E_i sigue una distribución exponencial de media K , por lo que su función de densidad es:

$$f_E(e) = \frac{1}{K} e^{-\frac{e}{K}}, \quad e \geq 0 \quad (77)$$

Sustituyendo $e = L_i + K - l$, la contribución asociada a una latencia l es proporcional a:

$$f_E(L_i + K - l) = \frac{1}{K} e^{-(L_i + K - l)/K} \quad (78)$$

Separando los términos de (78):

$$f_E(L_i + K - l) = \frac{1}{K} e^{-L_i/K} e^{-1+l/K} \quad (79)$$

El término $e^{-L_i/K}$ depende de la latencia anterior, pero no de la nueva latencia l . Por tanto, al considerar los posibles valores de L_i , dicho término afecta a la constante de normalización, pero no a la forma de la dependencia con l . En cambio, la dependencia con la latencia viene dada por el término $e^{-1+l/K}$. Por ello, la función de densidad de la rama sin reselección es proporcional a:

$$f_{NR}(l) \propto e^{-1+l/K}, \quad 0 \leq l \leq K \quad (80)$$

Para obtener la PDF normalizada, se escribe:

$$f_{NR}(l) = C e^{-1+l/K}, \quad 0 \leq l \leq K \quad (81)$$

donde C es una constante de normalización. Imponiendo que la densidad integre uno en el intervalo $[0, K]$:

$$\int_0^K C e^{-1+l/K} dl = 1 \quad (82)$$

Como:

$$\int_0^K e^{-1+l/K} dl = K(1 - e^{-1}) \quad (83)$$

se obtiene:

$$C = \frac{1}{K(1 - e^{-1})} \quad (84)$$

Por tanto, la PDF de latencia condicionada a transmisión sin reselección queda:

$$f_{NR}(l) = \frac{e^{-1+\frac{l}{K}}}{K(1-e^{-1})}, \quad 0 \leq l \leq K \quad (85)$$

Esta expresión es creciente con la latencia. Cuando $l = 0$, el término exponencial vale e^{-1} , mientras que cuando $l = K$, vale 1. Por tanto, la densidad aumenta conforme aumenta l , lo que indica que, en la rama sin reelección, las latencias próximas al *PDB* son más probables que las latencias bajas.

La razón física de este comportamiento es que, con $RRI = 2K$, las reservas periódicas están más separadas que la ventana de latencia permitida. En consecuencia, para que un paquete pueda aprovechar una reserva mantenida sin reeleccionar, su instante de generación debe situarse en una región temporal concreta respecto a dicha reserva. La distribución exponencial de la componente aleatoria E_i hace que esta condición favorezca valores altos de latencia, dando lugar a una PDF creciente.

Una vez obtenidas las dos ramas, la distribución total de latencia se calcula como una mezcla ponderada. Con probabilidad p_{rsl} , el paquete se transmite tras una reelección de recursos. En ese caso, su latencia sigue una distribución uniforme, ya que el nuevo recurso se selecciona dentro de la ventana de latencia permitida:

$$f_{rsl}(l) = \frac{1}{K}, \quad 0 \leq l \leq K \quad (86)$$

Con probabilidad $1 - p_{rsl}$, el paquete se transmite usando una reserva periódica mantenida, sin reelección. En ese caso, su latencia sigue la distribución obtenida anteriormente (85). Por tanto, la PDF total de latencia para modelos aperiódicos con $RRI = 2K$, $t_{\min} = 0$ y $PDB = K$, queda:

$$f_L(l) = p_{rsl} \frac{1}{K} + (1 - p_{rsl}) \frac{e^{-1+\frac{l}{K}}}{K(1-e^{-1})}, \quad 0 \leq l \leq K \quad (87)$$

Esta expresión recoge los dos mecanismos principales que determinan la latencia. La primera componente, uniforme, aparece cuando hay reelección de recursos, ya que el nuevo recurso se selecciona dentro de la ventana permitida. La segunda componente, creciente, aparece cuando no hay reelección y el paquete utiliza una reserva periódica previamente establecida. En este caso, la generación aperiódica de paquetes y el hecho de que $RRI = 2K$ provocan que las latencias altas sean más probables que las bajas.

Con este apartado, la cadena de Markov desarrollada previamente adquiere un papel directo en la caracterización de la latencia. En particular, una vez conocida la probabilidad total de reelección, p_{rsl} , obtenida mediante la expresión (7), la distribución de latencia puede obtenerse analíticamente mediante las fórmulas anteriores, para todos los modelos de tráfico analizados en este trabajo.

3.3. Modelo de pérdidas entre recepciones correctas

Otra aplicación de la cadena de Markov presentada en el apartado 3.1 consiste en utilizar la información obtenida para analizar métricas adicionales relacionadas con la continuidad temporal de la recepción en comunicaciones 5G NR V2X. En particular, en este apartado se propone un modelo analítico para caracterizar la distribución del número de paquetes perdidos consecutivamente entre dos recepciones correctas. Esta métrica resulta especialmente relevante en comunicaciones vehiculares porque no solo importa conocer la probabilidad media de recepción correcta, sino también cómo se distribuyen temporalmente las pérdidas. Dos escenarios pueden presentar una probabilidad media de recepción similar y, sin embargo, tener comportamientos muy distintos desde el punto de vista de la continuidad de recepción. En un caso, las pérdidas pueden aparecer de forma aislada, alternándose con recepciones correctas. En otro, pueden agruparse en ráfagas de varios paquetes consecutivos. Desde el punto de vista de V2X, esta diferencia es importante, ya que la pérdida de varios mensajes consecutivos puede afectar a la actualización del estado de los vehículos vecinos, a la percepción cooperativa del entorno o al tiempo de reacción ante eventos críticos.

3.3.1. Planteamiento del modelo

Para describir este comportamiento se define la variable aleatoria discreta N_{loss} , que representa el número de paquetes perdidos entre dos recepciones correctas consecutivas. Si tras recibir correctamente un paquete el siguiente paquete también se recibe correctamente, no se produce ninguna pérdida intermedia y se tiene: $N_{loss} = 0$. Si entre dos recepciones correctas se pierde un paquete, se tiene: $N_{loss} = 1$. De forma general, $N_{loss} = n$ indica que se han perdido n paquetes consecutivos antes de volver a recibir correctamente. El objetivo del modelo desarrollado en este apartado es obtener analíticamente la distribución de probabilidad de esta variable.

Aunque esta métrica no coincide estrictamente con el Packet Inter-Reception, sí está directamente relacionada con él. El PIR mide el intervalo temporal entre recepciones correctas consecutivas, mientras que N_{loss} mide la longitud de la ráfaga de pérdidas que aparece entre ellas. Por tanto, una mayor probabilidad de valores elevados de N_{loss} implica una mayor probabilidad de intervalos prolongados sin recepciones correctas. En este trabajo no se calcula directamente la distribución temporal del PIR, sino la distribución del número de paquetes perdidos consecutivamente entre recepciones correctas. La extensión del modelo hacia una caracterización temporal completa del PIR se plantea como una posible línea de trabajo futuro.

El modelo analítico de pérdidas consecutivas se construye a partir de dos tipos de información. En primer lugar, utiliza la probabilidad de reelección de recursos obtenida mediante la cadena de Markov descrita en el apartado 3.1, p_{rsi} . En segundo lugar, utiliza las probabilidades asociadas a los distintos tipos de pérdida de paquetes. Estas probabilidades se obtienen a partir del modelo analítico de PDR desarrollado en [7], el cual emplearemos como herramienta auxiliar en este trabajo. En concreto, se consideran las siguientes probabilidades de pérdida: p_{HD} , p_{SEN} , p_{PRO} y p_{COL} , que representan, respectivamente, las probabilidades de pérdida por half-duplex, sensado, propagación y colisión. Estas probabilidades se consideran mutuamente excluyentes dentro del modelo, por lo que la probabilidad de recepción correcta, denotada en adelante como p_{RX} , queda determinada como:

$$p_{\text{RX}} = 1 - p_{\text{HD}} - p_{\text{SEN}} - p_{\text{PRO}} - p_{\text{COL}} \quad (88)$$

Por tanto, la totalidad de resultados posibles de una transmisión queda representada por:

$$p_{\text{RX}} + p_{\text{HD}} + p_{\text{SEN}} + p_{\text{PRO}} + p_{\text{COL}} = 1 \quad (89)$$

Estas probabilidades representan, respectivamente, las pérdidas por *half-duplex*, sensado, propagación y colisión, de acuerdo con la descomposición del modelo auxiliar empleado en [7]. En el modelo propuesto, las pérdidas por *half-duplex*, sensado y propagación se agrupan en una única probabilidad de pérdida no asociada a colisión, mientras que las pérdidas por colisión se mantienen separadas debido a su posible persistencia temporal cuando los vehículos implicados conservan los mismos recursos reservados.

La probabilidad de reelección p_{rsi} resume la tendencia del vehículo a cambiar los recursos que tiene reservados. Su papel en el modelo de pérdidas consecutivas es

fundamental, ya que la reelección afecta a la persistencia temporal de ciertos errores. Si un vehículo mantiene los mismos recursos durante varias transmisiones consecutivas, algunos tipos de pérdida pueden repetirse. En cambio, si se produce una reelección, el patrón de recursos cambia y puede romperse la situación que estaba provocando pérdidas consecutivas.

El modelo analítico empleado para obtener p_{HD} , p_{SEN} , p_{PRO} y p_{COL} no constituye la contribución principal de este trabajo, pero sí está directamente conectado con él. En concreto, dicho modelo utiliza como una de sus entradas la probabilidad de reelección de recursos obtenida mediante la cadena de Markov desarrollada previamente. Por tanto, aunque el desarrollo completo de las probabilidades de pérdida procede de un modelo auxiliar, la cadena de Markov propuesta en este TFG permite parametrizar parte de dicho modelo.

3.3.2. Cadena de Markov y distribución de pérdidas consecutivas

Para modelar los paquetes perdidos consecutivos se propone una cadena de Markov discreta cuyos estados representan el resultado de la transmisión de un paquete. Esta cadena no pretende reproducir de nuevo todo el procedimiento interno de selección, reserva y reelección de recursos, ya que ese comportamiento ha sido modelado previamente. Su objetivo es más compacto: describir cómo se suceden las recepciones correctas y las pérdidas consecutivas.

La cadena propuesta está formada por tres estados: R , P y C . El estado R representa una recepción correcta. El estado P representa una pérdida por una causa distinta de colisión, agrupando las pérdidas por *half-duplex*, sensado y propagación. Finalmente, el estado C representa una pérdida por colisión.

Se define:

$$\begin{aligned} p_P &= p_{HD} + p_{SEN} + p_{PRO} \\ p_C &= p_{COL} \\ p_{RX} &= 1 - p_P - p_C \end{aligned} \tag{90}$$

donde p_P representa la probabilidad de sufrir una pérdida no asociada a colisión, p_C representa la probabilidad de pérdida por colisión y p_{RX} representa la probabilidad de recepción correcta.

El estado R actúa como estado de referencia. Cada vez que la cadena alcanza este estado, se considera que ha finalizado una secuencia de pérdidas consecutivas. Si después de estar en R la cadena vuelve directamente a R , no se ha perdido ningún paquete entre dos recepciones correctas y, por tanto, $N_{\text{loss}} = 0$. Si desde R la cadena pasa a P o a C , comienza una ráfaga de paquetes perdidos. Dicha ráfaga termina cuando la cadena vuelve de nuevo al estado R .

La cadena diferencia entre las pérdidas no asociadas a colisión y las pérdidas por colisión. Las pérdidas agrupadas en P se consideran no persistentes de forma explícita dentro de la cadena. Esto significa que, tras una pérdida de tipo P , el siguiente paquete vuelve a evaluarse de acuerdo con las mismas probabilidades que si el paquete anterior hubiese sido recibido correctamente. En cambio, el estado C sí incorpora memoria temporal, ya que una colisión puede repetirse si los vehículos implicados mantienen sus recursos reservados.

Para modelar la persistencia de las colisiones se define una probabilidad efectiva de renovación del conflicto de recursos. Esta probabilidad se denota como:

$$\eta = 1 - (1 - p_{\text{rsl}})^2 \quad (91)$$

Si p_{rsl} representa la probabilidad de que un vehículo reseleccione su recurso, entonces $1 - p_{\text{rsl}}$ representa la probabilidad de que dicho vehículo mantenga su reserva. En una colisión entre dos vehículos, $(1 - p_{\text{rsl}})^2$ representa la probabilidad de que dos vehículos implicados en una colisión mantengan sus recursos. Así, $1 - (1 - p_{\text{rsl}})^2$ representa la probabilidad de que alguno de los vehículos reseleccione. Para un modelo preciso sería necesario conocer el número de vehículos que colisionan en el mismo recurso, pero dado que el modelo no identifica explícitamente el número de vehículos implicados en cada colisión, se adopta una aproximación par a par, considerando el conflicto dominante entre dos vehículos. Además, se asume que las decisiones de reselección de los dos vehículos implicados son independientes. A partir de esta idea, se introduce una probabilidad auxiliar p_C^* . Esta probabilidad no representa directamente la probabilidad de colisión obtenida del modelo de PDR. La probabilidad de colisión es $p_C = p_{\text{COL}}$. En cambio, p_C^* representa una probabilidad efectiva de aparición o regeneración de una colisión dentro de la cadena, calculada de forma que la probabilidad estacionaria del estado C coincida con la probabilidad marginal de colisión proporcionada por el modelo auxiliar. Esta distinción es necesaria porque el estado C tiene persistencia temporal. Si se utilizase directamente p_C como probabilidad de generar una nueva colisión, la cadena estaría

contando simultáneamente las colisiones nuevas y las colisiones que persisten desde paquetes anteriores. Como consecuencia, la probabilidad estacionaria del estado C no coincidiría necesariamente con la probabilidad de colisión proporcionada por el modelo auxiliar. Por ello, p_C^* se calcula imponiendo que la probabilidad estacionaria del estado C sea precisamente p_C . Por tanto, p_C^* no debe interpretarse como una nueva probabilidad física medida directamente, sino como un parámetro efectivo de la cadena. Su valor se obtiene mediante una condición de estacionariedad, de forma que la persistencia introducida en el estado C no altere la probabilidad marginal de colisión proporcionada por el modelo auxiliar. Con estas definiciones, las transiciones desde R y desde P se plantean de la misma forma, ya que en ambos casos se considera que no existe una colisión persistente activa. Desde estos estados, el siguiente paquete puede perderse por una causa distinta de colisión con probabilidad p_P . Si no se produce una pérdida de tipo P , puede producirse una colisión efectiva con probabilidad p_C^* . Por tanto, la transición al estado C se produce con probabilidad: $(1 - p_P) \cdot p_C^*$. Finalmente, si no se produce una pérdida de tipo P ni se genera una nueva colisión, el paquete se recibe correctamente, con probabilidad $(1 - p_P) \cdot (1 - p_C^*)$.

Desde el estado C , la situación es distinta porque se considera que existe una colisión activa. En primer lugar, puede producirse una pérdida no asociada a colisión, transitando al estado P con probabilidad p_P . Si no se produce una pérdida de tipo P , la evolución depende de la reelección. Para que la colisión desaparezca, debe producirse una reelección por parte de al menos uno de los vehículos implicados y, además, el nuevo patrón de recursos no debe volver a generar una colisión. Esto ocurrirá con probabilidad $(1 - p_P) \cdot \eta \cdot (1 - p_C^*)$, volviendo con esta probabilidad al estado R . Por último, la cadena permanece en C si no se produce una pérdida de tipo P y la colisión sigue presente. Esto puede ocurrir de dos formas: que no haya reelección y, por tanto, se mantenga la colisión previa, o que haya reelección pero el nuevo patrón de recursos vuelva a producir una colisión. Esto ocurrirá con una probabilidad $(1 - p_P)[(1 - \eta) + \eta \cdot p_C^*]$. La cadena de Markov que modela este comportamiento con las probabilidades indicadas se puede observar en la Figura 7.

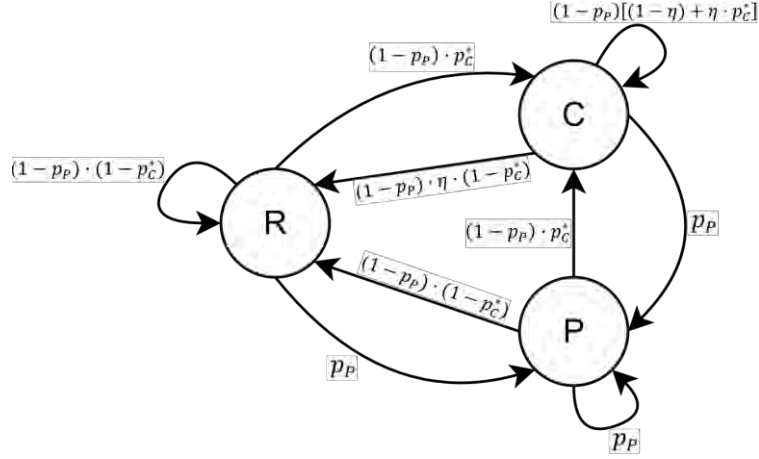


Figura 7. Cadena de Markov para el Modelo Analítico de pérdidas entre recepciones correctas.

Con esto, la matriz de transición de la cadena, siguiendo el orden de estados $[R, P, C]$, queda:

$$M = \begin{bmatrix} (1-p_P)(1-p_C^*) & p_P & (1-p_P)p_C^* \\ (1-p_P)(1-p_C^*) & p_P & (1-p_P)p_C^* \\ (1-p_P)\eta(1-p_C^*) & p_P & (1-p_P)[(1-\eta)+\eta p_C^*] \end{bmatrix} \quad (92)$$

La cadena se construye imponiendo que el vector estacionario coincida con las probabilidades marginales de recepción, pérdida no asociada a colisión y pérdida por colisión:

$$\boldsymbol{\pi} = [p_{RX} \quad p_P \quad p_C] \quad (93)$$

Con esta condición, podemos hallar el valor de p_C^* . En primer lugar, se observa que la probabilidad estacionaria del estado P queda directamente determinada por la matriz, ya que todas las filas tienen probabilidad p_P de transición hacia dicho estado:

$$\pi_P = \pi_R p_P + \pi_P p_P + \pi_C p_P \quad (94)$$

Como:

$$\pi_R + \pi_P + \pi_C = 1 \quad (95)$$

Se obtiene así que $\pi_P = p_P$. Por tanto, la probabilidad estacionaria del estado P coincide directamente con la probabilidad de pérdida no asociada a colisión. Para el estado C , se impone que su probabilidad estacionaria sea p_C . La ecuación de estacionariedad para C es:

$$\pi_C = (\pi_R + \pi_P)(1 - p_P)p_C^* + \pi_C(1 - p_P)[(1 - \eta) + \eta p_C^*] \quad (96)$$

Como se quiere que $\pi_C = p_C$ y que $\pi_R + \pi_P = 1 - p_C$, se obtiene:

$$p_C = (1 - p_C)(1 - p_P)p_C^* + p_C(1 - p_P)[(1 - \eta) + \eta p_C^*] \quad (97)$$

Despejando p_C^* :

$$p_C^* = \frac{p_C[1 - (1 - p_P)(1 - \eta)]}{(1 - p_P)[(1 - p_C) + p_C\eta]} \quad (98)$$

Esta expresión permite calcular la probabilidad efectiva p_C^* a partir de las probabilidades marginales p_P , p_C y de la probabilidad efectiva η asociada a la reelección de los vehículos implicados en una colisión. De este modo, la cadena conserva como probabilidad estacionaria de colisión el valor $p_C = p_{COL}$ proporcionado por el modelo auxiliar. Finalmente, la probabilidad estacionaria del estado R se obtiene como:

$$\pi_R = 1 - \pi_P - \pi_C \quad (99)$$

y, por tanto:

$$\pi_R = 1 - p_P - p_C = p_{RX} \quad (100)$$

Con ello se garantiza que la cadena reproduce a largo plazo las probabilidades de recepción correcta, pérdida no asociada a colisión y pérdida por colisión definidas en la expresión (93).

Una vez construida la matriz de transición, la distribución de las pérdidas consecutivas se obtiene como una distribución de primer retorno al estado R . Si la transición ocurre del estado R al estado R , se tendrán cero pérdidas, con probabilidad igual a la de esta transición. En cambio, si la cadena atraviesa n estados de pérdida antes de volver a R , entonces se han perdido exactamente n paquetes consecutivos. Se puede obtener una expresión analítica que proporcione la distribución de los paquetes perdidos consecutivos, interpretando el retorno al estado R como un problema de tiempo hasta absorción. Para ello, durante el cálculo de la distribución de pérdidas consecutivas se consideran los estados de pérdida P y C como estados transitorios, mientras que el retorno a R actúa como absorción. Esta formulación es equivalente a una distribución fase-tipo discreta, definida como la distribución del número de pasos hasta la absorción en una cadena de

Markov discreta. En dicha formulación, la función de masa puede escribirse mediante un vector inicial sobre los estados transitorios, una matriz de transición entre estados transitorios y un vector de absorción [27].

Para expresar esta distribución de forma compacta, se descompone la matriz de transición en tres elementos. Se define el vector a , que contiene las probabilidades de transición desde R hacia los estados de pérdida:

$$a = [p_P \quad (1 - p_P)p_C^*] \quad (101)$$

Se define también la submatriz Q , que contiene las transiciones entre estados de pérdida:

$$Q = \begin{bmatrix} p_P & (1 - p_P)p_C^* \\ p_P & (1 - p_P)[(1 - \eta) + \eta p_C^*] \end{bmatrix} \quad (102)$$

Finalmente, se define el vector b , que contiene las probabilidades de transición desde los estados de pérdida hacia R :

$$b = \begin{bmatrix} (1 - p_P)(1 - p_C^*) \\ (1 - p_P)\eta(1 - p_C^*) \end{bmatrix} \quad (103)$$

Con esta notación, la distribución de las pérdidas consecutivas, $f_{N_{loss}}(n)$, viene dada por:

$$f_{N_{loss}}(n) = \begin{cases} (1 - p_P)(1 - p_C^*), & n = 0 \\ aQ^{n-1}b, & n \geq 1 \end{cases} \quad (104)$$

Esta expresión proporciona directamente la probabilidad de perder exactamente n paquetes consecutivos entre dos recepciones correctas. El vector a determina cómo comienza una ráfaga de pérdidas, la matriz Q determina cómo se prolonga dicha ráfaga y el vector b determina la probabilidad de recuperar la recepción correcta.

De esta forma, una vez conocidas las probabilidades p_P , p_C , p_{rsi} , y calculadas η y p_C^* , podemos obtener la distribución completa de los paquetes perdidos entre dos recepciones correctas de forma analítica mediante la expresión (104).

4. RESULTADOS Y VALIDACIÓN

En este capítulo se presentan los resultados obtenidos con los modelos analíticos desarrollados en el capítulo anterior y se valida su comportamiento mediante comparación con simulación. El objetivo principal de esta validación es comprobar que las expresiones y cadenas de Markov propuestas reproducen de forma adecuada los mecanismos de gestión de recursos, latencia y pérdidas consecutivas entre recepciones correctas considerados en el sistema.

En primer lugar, se describen las herramientas empleadas para la validación, los modelos de tráfico considerados y las configuraciones radio evaluadas. A continuación, se valida el modelo de gestión de recursos, ya que sus resultados constituyen la base para los modelos posteriores. Finalmente, se validan los modelos analíticos asociados a la distribución de la latencia y al número de paquetes perdidos consecutivamente entre recepciones correctas.

4.1. Herramientas y modelos de tráfico considerados

4.1.1. Herramientas de validación

Para la validación de los modelos analíticos desarrollados en este trabajo se han empleado dos herramientas complementarias. En primer lugar, se ha utilizado un simulador de comunicaciones vehiculares desarrollado en la UMH, que se toma como simulador de referencia. Esta herramienta permite reproducir escenarios V2X con un mayor nivel de detalle, incorporando distintos aspectos de la pila de comunicaciones, de los servicios vehiculares y del comportamiento del sistema. En segundo lugar, se ha desarrollado una herramienta propia en MATLAB, más ligera y específica, cuyo objetivo es reproducir de forma controlada la lógica temporal considerada en el modelo analítico.

Simulador basado en OMNeT++

El simulador de referencia integra distintas tecnologías de comunicación vehicular, incluyendo ITS-G5, LTE-V2X y 5G NR V2X, lo que permite analizar y comparar su rendimiento bajo diferentes condiciones. Además, incorpora implementaciones de servicios cooperativos estandarizados, como el Cooperative Awareness Service (CAS), encargado de la difusión periódica de información de estado de los vehículos, y el Collective Perception Service (CPS), orientado al intercambio de información percibida mediante sensores distribuidos. También permite definir servicios V2X genéricos

configurables, capaces de generar tráfico con tamaños de paquete e intervalos de generación tanto fijos como variables.

Dicho simulador está basado en OMNeT++, un entorno de simulación de eventos discretos ampliamente utilizado en investigación para el modelado de redes de comunicación. Sobre esta plataforma se emplea el framework INET, que proporciona modelos de protocolos y tecnologías de red, así como Veins, una extensión orientada a la simulación de redes vehiculares. La combinación de OMNeT++, INET y Veins proporciona una plataforma adecuada para la evaluación de sistemas V2X, permitiendo obtener resultados representativos en escenarios vehiculares.

No obstante, el uso de un simulador completo implica un mayor coste computacional, especialmente cuando se desea realizar un gran número de pruebas, modificar parámetros de forma iterativa o comprobar rápidamente el efecto de distintas hipótesis del modelo analítico. Es por esto que el simulador propio de MATLAB es especialmente útil. Esta herramienta no pretende sustituir al simulador de referencia, sino actuar como una verificación intermedia entre el modelo analítico y la simulación completa.

Simulador basado en MATLAB

El simulador desarrollado en MATLAB reproduce de forma secuencial la lógica temporal considerada en el modelo analítico. Para ello, parte de la definición de un modelo de tráfico y de un conjunto de parámetros de transmisión, incluyendo el tamaño de los paquetes, el intervalo de generación, el RRI, el PDB y los parámetros necesarios para determinar el número de subcanales ocupados por cada paquete. A partir de estos datos se genera una secuencia de transmisiones, donde cada paquete tiene asociado un instante de generación, un tamaño y unos requisitos temporales.

Durante la simulación, la herramienta actualiza la reserva activa y evalúa si dicha reserva puede utilizarse para transmitir el paquete actual. Si la reserva se produce antes de que el paquete haya sido generado, se contabiliza como reserva descartada y se avanza hasta la siguiente oportunidad de transmisión. En caso contrario, se comprueba si existen condiciones que obliguen a realizar una reelección de recursos, ya sea por agotamiento del contador, por incumplimiento del requisito de latencia, por aumento del número de subcanales necesarios o por la combinación simultánea de tamaño y latencia.

A partir de este procedimiento, la herramienta obtiene las magnitudes principales asociadas al modelo de gestión de recursos y al modelo de latencia: probabilidades de reelección, probabilidad de descarte de reserva, distribución del número de subcanales

reservados y distribución de la latencia de transmisión. Por tanto, aunque se trata de una herramienta simplificada, resulta útil para comprobar que la cadena de Markov reproduce correctamente la lógica temporal considerada antes de comparar los resultados con el simulador de referencia.

4.1.2. Modelos de tráfico considerados

Para realizar la validación se ha definido un conjunto de modelos de tráfico que permiten aumentar progresivamente la complejidad del proceso de generación de paquetes. Estos modelos se han construido tomando como referencia las configuraciones de tráfico propuestas por 3GPP [5], que incluyen tanto tráfico periódico como tráfico aperiódico y contemplan distintos tamaños de mensaje y requisitos de latencia. No obstante, en este trabajo no se emplean directamente dichas configuraciones de forma cerrada, sino que se utilizan como punto de partida para definir un conjunto propio de modelos de validación adaptado a los mecanismos analíticos estudiados.

El objetivo de esta elección es disponer de modelos sencillos que permitan aislar el efecto de cada fenómeno considerado en la cadena de Markov. En primer lugar, se consideran modelos periódicos, en los que el intervalo entre generaciones consecutivas de paquetes es constante. Estos modelos permiten validar inicialmente el comportamiento básico del mecanismo SPS, incluyendo el mantenimiento de reservas, el decremento del contador de reelección y las reelecciones asociadas al tamaño del paquete. Posteriormente, se introducen modelos aperiódicos, en los que el intervalo entre generaciones consecutivas deja de ser constante y pasa a incluir una componente aleatoria. Estos modelos permiten validar los mecanismos adicionales incorporados en este trabajo, especialmente las reservas descartadas por ausencia de paquete generado y las reelecciones por latencia. Los modelos de tráfico considerados se resumen en la Tabla 1.

ID	Tipo de generación	Tamaño de paquete (Bytes)	Intervalo de generación (ms)	RRI (ms)	PDB (ms)
P1	Periódica	Patrón (190,190,190,190,300)	Constante (100)	100	100
P2	Periódica	Aleatorio (1200 con probabilidad 0.2 y 800 con probabilidad 0.8)	Constante (100)	100	100
P3	Periódica	Aleatoriamente uniforme {200, 400, 600, 800, 1000, 1200}	Constante (100)	100	100
AP1	Aperiódica	Constante (200)	Aleatorio (50+Exp(50))	50	50
AP2	Aperiódica	Aleatoriamente uniforme {200, 400, 600, 800, 1000, 1200}	Aleatorio (50+Exp(50))	50	50
AP3	Aperiódica	Constante (200)	Aleatorio (50+Exp(50))	100	50
AP4	Aperiódica	Aleatoriamente uniforme {200, 400, 600, 800, 1000, 1200}	Aleatorio (50+Exp(50))	100	50

Tabla 1. Modelos de tráfico para validación del modelo analítico.

El modelo P1 considera un patrón periódico de tamaños de mensaje formado por cuatro paquetes de 190 bytes y un paquete de 300 bytes. Este caso permite estudiar un tráfico periódico con variación determinista del tamaño de paquete. El modelo P2 introduce una variación aleatoria discreta, donde el tamaño del paquete puede ser de 1200 bytes con probabilidad 0,2 o de 800 bytes con probabilidad 0,8. Finalmente, el modelo P3 considera un tamaño de paquete aleatorio uniforme entre 200 y 1200 bytes, con incrementos de 200 bytes. De este modo, los tres modelos periódicos permiten analizar cómo afecta la variabilidad del tamaño de paquete a la probabilidad de reservar un determinado número de subcanales y a la probabilidad de reelección por tamaño.

En los modelos aperiódicos, el intervalo entre generaciones consecutivas no es constante, sino que se modela como una componente fija de 50 ms más una variable aleatoria exponencial de media 50 ms. Los modelos AP1 y AP2 consideran $RRI = PDB = 50$ ms. AP1 utiliza paquetes de tamaño constante de 200 bytes, por lo que permite estudiar únicamente el efecto de la generación aperiódica sobre el uso de las reservas. AP2 introduce además tamaño variable, considerando paquetes distribuidos uniformemente entre 200 y 1200 bytes con pasos de 200 bytes.

Por otra parte, los modelos AP3 y AP4 consideran $RRI = 2PDB = 100$ ms. Esta configuración permite analizar cómo cambia el equilibrio entre reservas descartadas y reelecciones por latencia cuando el intervalo de reserva aumenta. AP3 mantiene tamaño de paquete constante y AP4 añade tamaño de paquete variable, por lo que constituye el caso más completo dentro de los modelos considerados, al combinar generación aperiódica, tamaño variable, reservas descartadas y reelecciones por latencia y por tamaño.

Es importante destacar que estos modelos definen únicamente el patrón de generación de tráfico, el tamaño de los paquetes, el PDB y el RRI. Otros parámetros del sistema, como la numerología, el MCS o el ancho de banda, se configuran de forma independiente. De este modo, un mismo modelo de tráfico puede evaluarse bajo distintas configuraciones radio, permitiendo estudiar separadamente el efecto del tráfico y el efecto de los parámetros radio.

Para la comparación con el simulador de referencia se emplean las configuraciones disponibles para cada modelo de tráfico. Algunas combinaciones de modelo, numerología y MCS no se consideran porque, bajo la configuración de ancho de banda de 20 MHz utilizada en simulación, determinados tamaños de paquete no pueden alojarse en los

recursos disponibles. Por ello, la validación se realiza sobre el conjunto de simulaciones disponibles.

4.2. Validación de los modelos analíticos

4.2.1. Validación del modelo de gestión de recursos

El primer modelo que se valida es el modelo analítico de gestión de recursos presentado en el apartado 3.1. Este modelo permite obtener, a partir de la distribución estacionaria de la cadena de Markov, las probabilidades asociadas a los distintos mecanismos de reelección, la probabilidad de descarte de reserva y la distribución del número de subcanales reservados.

Para validar el modelo de gestión de recursos se comparan las magnitudes directamente asociadas a los estados y transiciones de la cadena de Markov. En particular, se analizan las probabilidades de reelección definidas en el apartado 3.1: reelección por contador (p_{CR}), reelección por latencia (p_{LR}), reelección por tamaño (p_{SR}), y reelección simultánea por tamaño y latencia (p_{SLR}). Con la suma de estas obtenemos además la probabilidad total de reelección, p_{rsl} .

Además, se evalúa la probabilidad de descarte de reserva, p_d . Junto a estas probabilidades, se valida también la distribución del número de subcanales reservados. Esta distribución indica la probabilidad de que la reserva activa esté asociada a un determinado número de subcanales, y permite comprobar si el modelo analítico reproduce correctamente la ocupación de recursos derivada de los tamaños de paquete, el MCS y la numerología. Como medida resumida de esta distribución se emplea también el número medio de subcanales reservados.

La validación se realiza inicialmente para una configuración radio de referencia, manteniendo fijos la numerología y el MCS (ver Tabla 2). Esta elección permite analizar con detalle el comportamiento del modelo analítico frente a los siete modelos de tráfico considerados sin introducir un número excesivo de figuras redundantes. En concreto, se realiza la validación inicial para la numerología 1 ($\mu = 1$) y MCS=13, ya que para ella se dispone de resultados del simulador de referencia para todos los modelos de tráfico. Además, desde el punto de vista del modelo de gestión de recursos, la numerología y el MCS no modifican la estructura de la cadena de Markov, sino que su efecto aparece de forma indirecta a través del número de subcanales necesarios para transmitir cada paquete. Por tanto, dos configuraciones radio que generen la misma distribución de subcanales requeridos proporcionan las mismas probabilidades de entrada a la cadena y,

en consecuencia, producen el mismo comportamiento desde el punto de vista de la gestión de recursos. Esta variabilidad en la distribución de subcanales requeridos puede observarse con detalle manteniendo una misma configuración radio, gracias a los distintos modelos de tráfico considerados. No obstante, tras mostrar los resultados para esta configuración concreta, el resto de las configuraciones disponibles se validan de forma global analizando el error absoluto del modelo analítico respecto a los simuladores. De esta forma, se realiza una validación completa para los modelos de tráfico y configuraciones radio disponibles.

Parámetros	Configuración
Separación entre subportadoras	30 kHz
Índice de la tabla MCS	1
Índice MCS	13 (16QAM-490)
Tamaño de subcanal	10 RBs
Tamaño del SCI de primera etapa	10 RBs y 2 símbolos OFDM

Tabla 2. Configuración inicial de validación.

La Figura 8 muestra la comparación de las probabilidades de reelección y descarte para los siete modelos de tráfico. En cada subfigura se representan conjuntamente los resultados del simulador MATLAB, del modelo analítico y del simulador de referencia. Esta representación permite comprobar en una única figura el comportamiento del modelo para tráfico periódico y aperiódico, así como para paquetes de tamaño constante y variable.

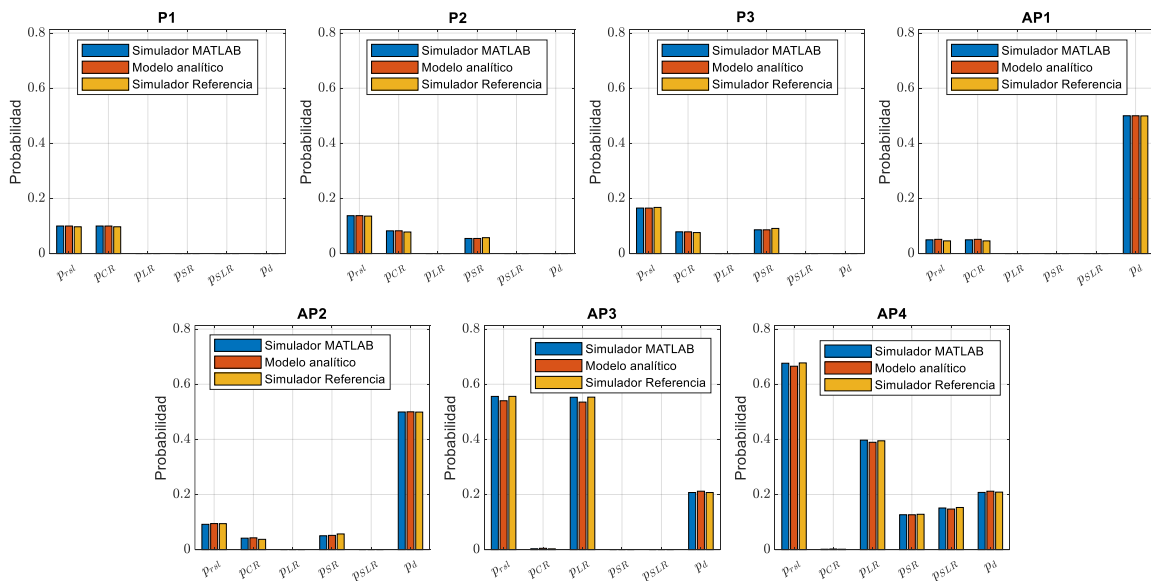


Figura 8. Probabilidades de reelección y descarte para todos los modelos de tráfico con $\mu = 1$ y MCS=13.

Los resultados muestran que el modelo analítico reproduce adecuadamente las probabilidades obtenidas mediante simulación. La coincidencia con el simulador MATLAB es especialmente relevante, ya que ambos comparten las mismas hipótesis de modelado y, por tanto, permite verificar la formulación matemática de la cadena de Markov. Además, la comparación con el simulador de referencia confirma que el modelo mantiene un comportamiento coherente en un entorno de simulación más completo. Las pequeñas diferencias que puedan aparecer pueden atribuirse principalmente a dos factores. Por un lado, el modelo analítico trabaja con probabilidades estacionarias y con una representación simplificada del proceso, mientras que el simulador de referencia genera secuencias temporales concretas e incorpora un mayor nivel de detalle. Por otro lado, las probabilidades obtenidas mediante simulación se estiman a partir de un número finito de muestras, por lo que es esperable que exista cierta variabilidad estadística en los resultados. Bajo las mismas hipótesis de modelado y con un número suficientemente elevado de muestras, los resultados simulados deberían converger hacia las probabilidades teóricas proporcionadas por el modelo analítico.

La Figura 9 muestra además la distribución del número de subcanales reservados para los mismos siete modelos de tráfico, lo cual corresponde a la ocupación de recursos en frecuencia.

En cada subfigura se compara la distribución obtenida mediante el simulador MATLAB, el modelo analítico y el simulador de referencia.

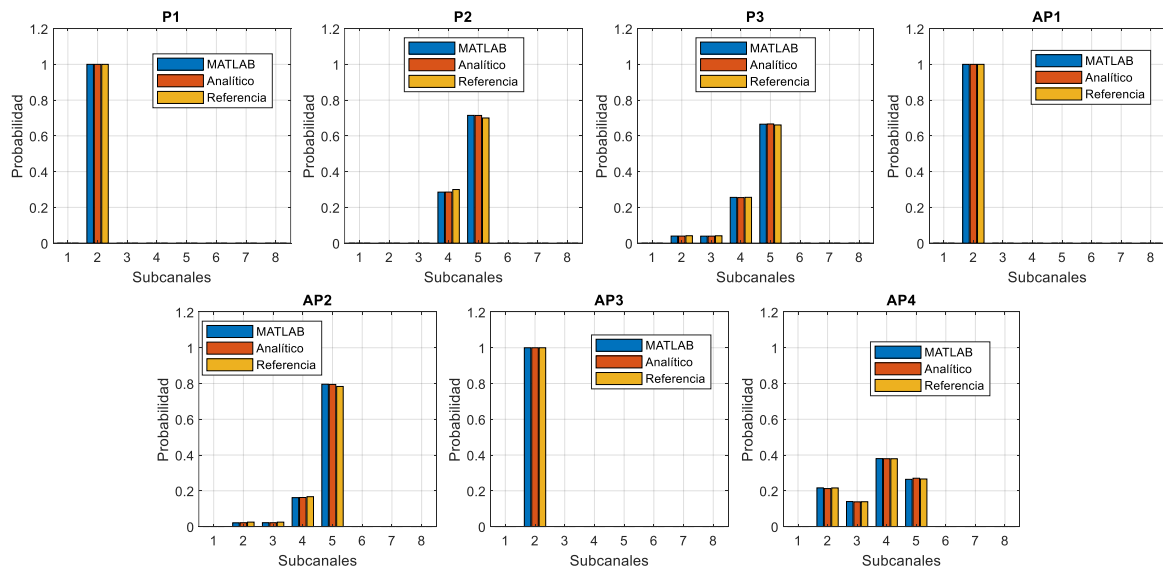


Figura 9. Distribución del número de subcanales reservados para todos los modelos de tráfico con $\mu = 1$ y MCS=13.

La comparación muestra que el modelo analítico reproduce correctamente la distribución de subcanales reservados. En los modelos de tamaño constante, la distribución se concentra en un único valor, mientras que en los modelos de tamaño variable la probabilidad se reparte entre distintos números de subcanales. Esta diferencia es capturada por el modelo analítico, lo que confirma que la cadena incorpora correctamente el efecto del tamaño de paquete sobre la selección de recursos.

Para cuantificar las diferencias observadas en las figuras anteriores se calcula el error absoluto entre el modelo analítico y cada simulador. Para una métrica escalar x , el error absoluto se define como:

$$e_x = |x_{ana} - x_{sim}| \quad (105)$$

donde x_{ana} representa el valor obtenido mediante el modelo analítico y x_{sim} el valor obtenido mediante simulación. En el caso de la distribución de subcanales, se calcula el error absoluto para cada posible número de subcanales reservados:

$$e_s = |p_{R,ana}(s) - p_{R,sim}(s)| \quad (106)$$

donde $p_{R,ana}(s)$ y $p_{R,sim}(s)$ representan, respectivamente, la probabilidad analítica y simulada de reservar s subcanales.

La Figura 10 resume el error absoluto entre el modelo analítico y el simulador MATLAB para el conjunto de probabilidades evaluadas, incluyendo las probabilidades de reelección, la probabilidad de descarte y la probabilidad de reservar s subcanales, $p_R(s)$. Estos errores se representan para los modelos de tráfico considerados y para distintas configuraciones de numerología y MCS.

De forma análoga, la Figura 11 muestra el error absoluto obtenido al comparar el modelo analítico con el simulador de referencia bajo las mismas condiciones. Esta representación permite evaluar de forma compacta si las diferencias entre el modelo analítico y las simulaciones se mantienen reducidas al variar tanto el modelo de tráfico como la configuración radio considerada.

La Figura 10 y la Figura 11 muestran valores mayores de error en el modelo de tráfico P2 en el caso de la distribución de subcanales respecto al simulador de referencia, y en el modelo AP3 en el caso de las probabilidades de reelección en ambos simuladores. Para explicar esto, por un lado, se observa en la Figura 9 cómo la distribución de subcanales de P2 es igual para el simulador de MATLAB y para el modelo analítico. Esto es esperable, ya que ambos parten de la misma distribución de subcanales requeridos por los paquetes, calculada con un script de MATLAB a partir de los valores de configuración radio. Ambos realizan un proceso probabilístico para obtener los resultados (aunque el simulador de MATLAB lo haga recogiendo muchas muestras aleatorias y el modelo analítico resolviendo una cadena de Markov). En cambio, el simulador de referencia reproduce un comportamiento más detallado, donde pueden aparecer factores no contemplados de forma analítica. Además, de forma general, es esperable obtener un menor error respecto al simulador de MATLAB ya que este realiza simulaciones con un número elevado de paquetes y proporciona valores más estables que los obtenidos mediante el simulador de referencia, cuyas simulaciones finitas tienen menor duración.

Por otro lado, al comparar el modelo analítico con ambos simuladores, se observa cómo la probabilidad de reelección por latencia y, en consecuencia, la probabilidad total de reelección, presentan un mayor error. Esto es esperable teniendo en cuenta la aproximación que hace la cadena de Markov respecto a la probabilidad de reelección por latencia, como se ha explicado en el apartado 3.1.3. La probabilidad de reelección por latencia no es la misma para cada paquete, sino que esta probabilidad introduce una memoria temporal, la cual se ha simplificado en la cadena. Es normal, por tanto, que en los modelos AP3 y AP4, donde aparecen reelecciones por latencia, estas presenten un mayor error. Además, el modelo AP3 tiene una menor probabilidad de reelección global, lo que se traduce en un mayor número de paquetes consecutivos sin reelección, o dicho de otra forma, mayor variabilidad temporal en las probabilidades de transición. Por ello, se obtiene un mayor error en AP3 fruto de estas aproximaciones.

No obstante, las diferencias entre el modelo analítico y los simuladores son reducidas en el conjunto de configuraciones evaluadas. Concretamente, el error absoluto máximo observado se mantiene en torno a $2 \cdot 10^{-2}$, lo que indica una buena correspondencia entre las probabilidades obtenidas analíticamente y las estimadas mediante simulación.

4.2.2. Validación del modelo de distribución de la latencia

Una vez validado el modelo de gestión de recursos, se procede a validar el modelo analítico de distribución de la latencia desarrollado en el apartado 3.2. A diferencia del modelo anterior, en este caso la magnitud de interés no es un conjunto de probabilidades asociadas a los estados de la cadena de Markov, sino la distribución de probabilidad de la latencia experimentada por los paquetes transmitidos. Por tanto, la validación se realiza comparando la función de densidad obtenida analíticamente con la distribución obtenida mediante simulación.

En los modelos considerados, esta magnitud queda determinada por la relación entre el proceso de generación de paquetes, la periodicidad de las reservas, el Packet Delay Budget y la posible aparición de reselecciones de recursos. Por este motivo, la validación se organiza siguiendo los tres comportamientos teóricos identificados en el desarrollo analítico. En primer lugar, se consideran los modelos periódicos, para los que la distribución de la latencia es uniforme dentro del intervalo permitido. En segundo lugar, se analizan los modelos aperiódicos con $RRI = PDB$, en los que la operación módulo asociada a la generación aperiódica también conduce a una distribución uniforme. Finalmente, se consideran los modelos aperiódicos con $RRI = 2PDB$, donde la rama sin reselección da lugar a una distribución creciente de la latencia.

El simulador MATLAB permite obtener una estimación precisa de la distribución de la latencia, ya que puede ejecutarse con un número elevado de paquetes y bajo coste computacional. Por este motivo, en las figuras se representa su distribución mediante un histograma normalizado. Sobre dicho histograma se superponen la curva obtenida mediante el modelo analítico y la distribución obtenida a partir del simulador de referencia.

El simulador de referencia permite contrastar los resultados en un entorno de simulación más completo. No obstante, debido a su mayor coste computacional, las simulaciones disponibles para esta herramienta contienen un número menor de muestras que las generadas mediante el simulador MATLAB. Como consecuencia, la distribución estimada a partir del simulador de referencia puede presentar ligeras oscilaciones, especialmente en aquellos casos en los que la distribución esperada es prácticamente uniforme. Este efecto se aprecia con mayor claridad en los modelos periódicos, ya que en ellos la latencia se distribuye en un intervalo de 100 ms, frente al intervalo de 50 ms de los modelos aperiódicos considerados. Por tanto, para un número de muestras

comparable, la estimación estadística dispone de menos muestras por intervalo de latencia, lo que hace que la curva del simulador de referencia presente mayores fluctuaciones. Estas oscilaciones no deben interpretarse como una discrepancia del modelo analítico, sino como una consecuencia del menor número de muestras disponible en la simulación de referencia.

La Figura 12 muestra la distribución de la latencia para los siete modelos de tráfico considerados, utilizando como configuración de referencia la numerología $\mu = 1$ y $MCS = 13$. Cada subfigura representa el histograma normalizado obtenido mediante el simulador MATLAB, la curva del modelo analítico y la curva correspondiente al simulador de referencia.

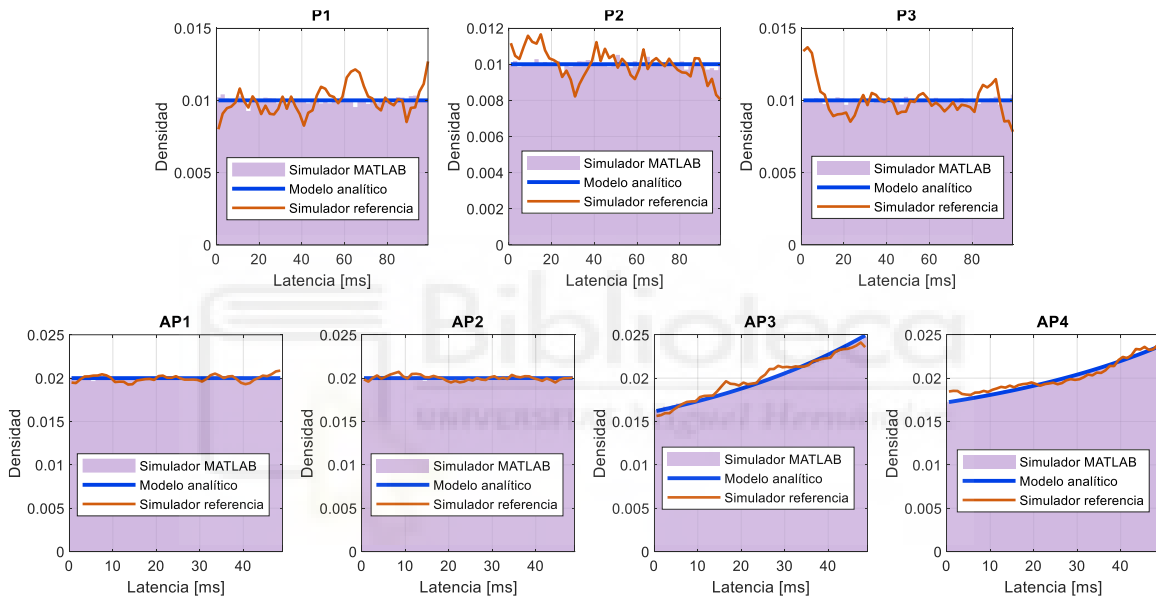


Figura 12. Distribución de la latencia para todos los modelos de tráfico con $\mu = 1$ y $MCS = 13$.

Los resultados muestran que el modelo analítico reproduce correctamente la forma de la distribución de latencia en los distintos casos considerados. En los modelos periódicos, P1, P2 y P3, la distribución obtenida es uniforme, de acuerdo con lo previsto por el modelo analítico. En estos casos, la generación de paquetes está sincronizada con la periodicidad de las reservas, por lo que la latencia queda distribuida de forma uniforme dentro del intervalo permitido. Aunque la curva del simulador de referencia presenta mayores oscilaciones en estos modelos, estas variaciones se deben al menor número de muestras disponible por intervalo de latencia.

En los modelos aperiódicos AP1 y AP2, donde se cumple $RRI = PDB$, también se observa una distribución uniforme. Este resultado valida la propiedad analítica descrita

en el apartado 3.2, según la cual la generación aperiódica de paquetes no modifica la uniformidad de la latencia cuando el intervalo de reserva coincide con el Packet Delay Budget. Por tanto, aunque el proceso de generación de paquetes contiene una componente aleatoria, la latencia resultante mantiene una distribución uniforme en el intervalo considerado.

Por último, los modelos AP3 y AP4 presentan un comportamiento diferente, ya que en estos casos se cumple $RRI = 2PDB$. La distribución deja de ser uniforme y adquiere una forma creciente, de manera que las latencias más altas presentan mayor probabilidad que las latencias bajas. Este comportamiento coincide con el desarrollo analítico realizado previamente en el apartado 3.2, donde se obtuvo que la rama sin reelección introduce una componente creciente en la función de densidad. La coincidencia entre la curva analítica y las distribuciones obtenidas mediante simulación confirma que el modelo reproduce correctamente el cambio de forma asociado a esta configuración.

Para cuantificar las diferencias entre el modelo analítico y las simulaciones se calcula el error absoluto punto a punto entre las funciones de densidad evaluadas en los mismos valores de latencia. Para un punto de latencia l_i , el error absoluto se define como:

$$e_i = |f_{L,ana}(l_i) - f_{L,sim}(l_i)| \quad (107)$$

donde $f_{L,ana}(l_i)$ representa el valor de la función de densidad obtenido mediante el modelo analítico y $f_{L,sim}(l_i)$ el valor estimado mediante simulación. A partir de este error punto a punto se calcula el error absoluto medio:

$$MAE = \frac{1}{N} \sum_{i=1}^N |f_{L,ana}(l_i) - f_{L,sim}(l_i)| \quad (108)$$

donde N es el número de puntos de latencia considerados. Esta métrica permite resumir en un único valor la diferencia media punto a punto entre la distribución analítica y la distribución simulada. Además, también se considera el error absoluto máximo:

$$e_{\max} = \max_i |f_{L,ana}(l_i) - f_{L,sim}(l_i)| \quad (109)$$

Esta métrica permite identificar la mayor desviación puntual entre ambas distribuciones. A continuación, se realiza una evaluación global para el resto de configuraciones de

numerología y MCS con el fin de comprobar si la validez del modelo de latencia se mantiene al modificar la configuración radio. La Figura 13 muestra estos errores (error absoluto medio y máximo) para distintos modelos de tráfico y configuraciones, tanto al comparar el modelo analítico con el simulador MATLAB como al compararlo con el simulador de referencia.

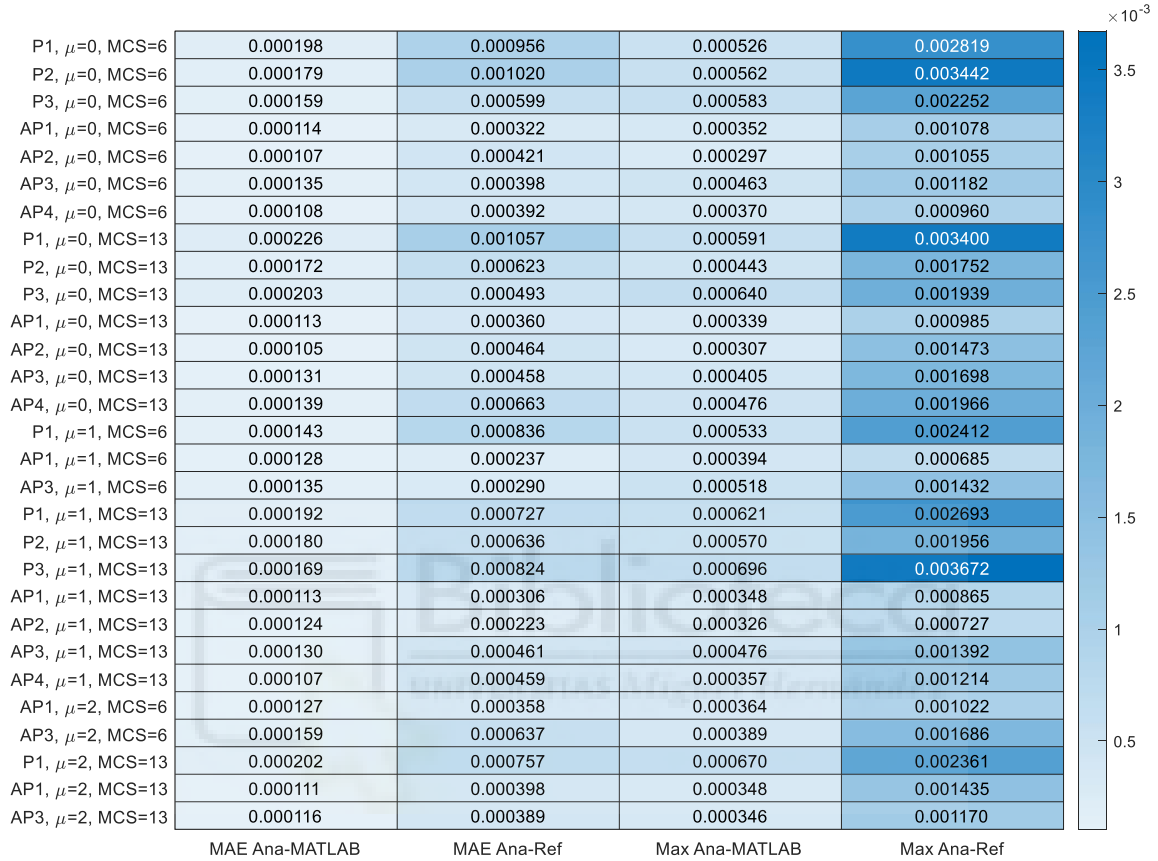


Figura 13. Error absoluto medio (MAE) y error absoluto máximo del modelo analítico de latencia respecto a ambos simuladores para las configuraciones evaluadas.

Los resultados confirman que el modelo analítico mantiene una buena correspondencia con las simulaciones en el conjunto de configuraciones consideradas. En todos los casos, el error absoluto máximo se mantiene en el orden de 10^{-3} . Los valores de error más altos se observan en la columna de la derecha de la Figura 13, la cual representa el error absoluto máximo al comparar el modelo analítico con el simulador de referencia. Como ya se ha explicado, este error puede atribuirse principalmente al menor número de muestras disponible en el simulador de referencia, debido al mayor coste computacional de dicha herramienta. Este efecto es especialmente visible en los modelos periódicos, donde la distribución se reparte sobre un intervalo de latencia mayor (100 ms frente a los 50 ms de los modelos de tráfico aperiódicos). Al existir un mayor rango de latencia, las muestras obtenidas por intervalo de latencia son menores, lo que da lugar a una curva

menos estable. Estas desviaciones puntuales mayores aumentan directamente este error máximo, como se observa también en la Figura 12.

A pesar de esto, incluso en este caso el error sigue siendo muy reducido, por lo que puede concluirse que los resultados obtenidos validan el modelo analítico de distribución de la latencia. El modelo reproduce correctamente la distribución uniforme asociada a los modelos periódicos y a los modelos aperiódicos con $RRI = PDB$, así como la distribución creciente que aparece en los modelos aperiódicos con $RRI = 2PDB$. Por tanto, la expresión analítica desarrollada permite caracterizar la latencia de forma correcta, utilizando como entrada las probabilidades obtenidas previamente mediante el modelo de gestión de recursos.

4.2.3. Validación del modelo de pérdidas entre recepciones correctas

El último modelo que se valida es el modelo analítico de pérdidas consecutivas entre recepciones correctas, presentado en el apartado 3.3.

El modelo analítico utilizado para esta métrica tiene una estructura compacta basada en tres estados: recepción correcta, pérdida no asociada a colisión y pérdida por colisión. A partir de esta cadena se obtiene la distribución $P(N_{\text{loss}})$, empleando como entradas la probabilidad de reelección de recursos calculada mediante el modelo de gestión de recursos y las probabilidades de pérdida asociadas a half-duplex, sensado, propagación y colisión. Esta validación se realiza comparando el simulador de referencia con el modelo analítico, ya que el simulador de MATLAB desarrollado en este trabajo únicamente simula transmisiones y no tiene la capacidad de simular si un paquete se recibe correctamente o se pierde.

Como se muestra en el desarrollo de este modelo, la distribución de pérdidas consecutivas depende principalmente de las probabilidades de pérdida de paquetes, las cuales dependen de la distancia entre transmisor y receptor. Por tanto, en primer lugar se representa la distribución $P(N_{\text{loss}})$ para tres rangos de distancia representativos: una distancia corta, una distancia media y una distancia elevada. Se han seleccionado estos tres rangos porque representan condiciones de enlace diferenciadas: una zona de alta fiabilidad a corta distancia, una zona intermedia donde comienzan a aparecer pérdidas con mayor frecuencia y una zona de distancia elevada donde la degradación del enlace aumenta la probabilidad de ráfagas de pérdidas. Esta comparación permite observar cómo cambia la

forma de la distribución a medida que aumenta la separación entre transmisor y receptor. La Figura 14 representa esta distribución para los tres modelos de tráfico periódicos con $\mu=1$ y MCS=13 para tres rangos de distancia distintos.

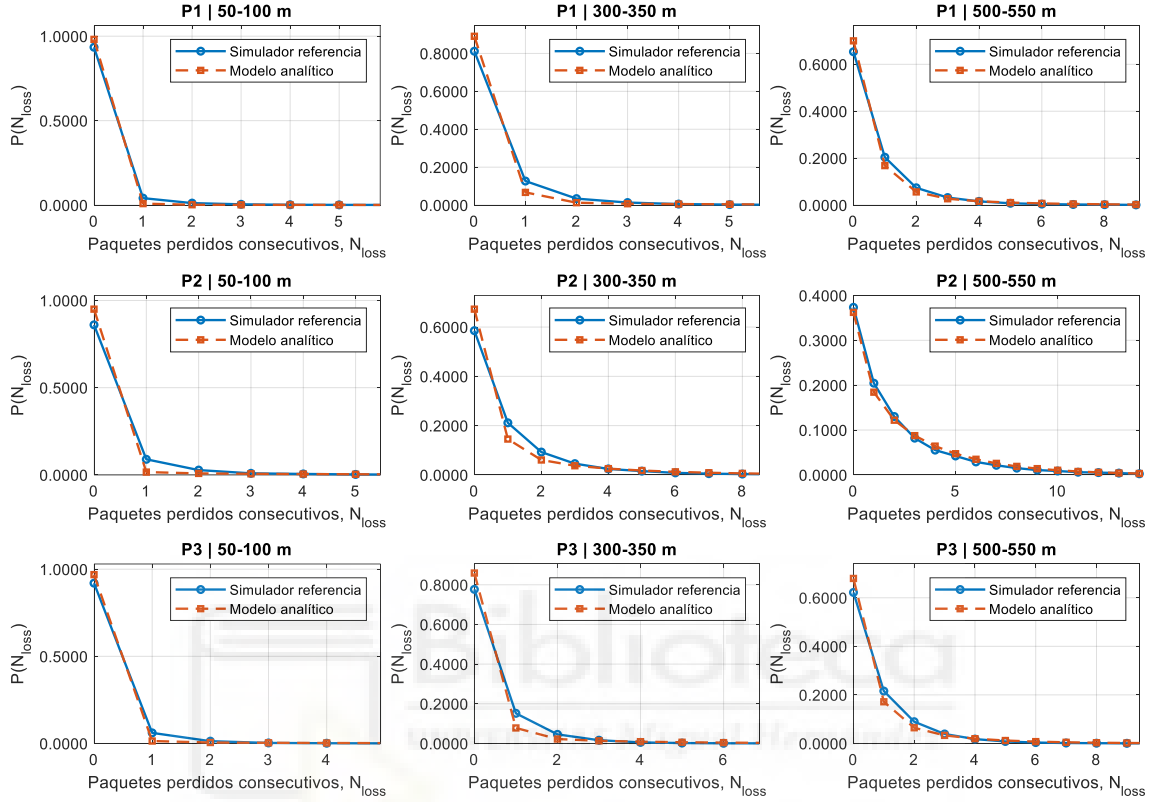


Figura 14. Distribución de pérdidas consecutivas para los modelos de tráfico periódicos.

Para distancias reducidas, la mayor parte de la probabilidad se concentra en valores bajos de N_{loss} , especialmente en $N_{\text{loss}} = 0$, lo que indica que las recepciones correctas consecutivas son frecuentes. A medida que aumenta la distancia, la probabilidad de recepción correcta disminuye y la distribución se desplaza progresivamente hacia valores mayores de N_{loss} , reflejando una mayor probabilidad de sufrir ráfagas de pérdidas consecutivas.

Las diferencias más visibles se producen en algunos valores bajos de N_{loss} , especialmente en torno a $N_{\text{loss}} = 1$. Este comportamiento puede explicarse porque dichos valores son especialmente sensibles al equilibrio entre pérdidas aisladas y pérdidas persistentes. En un esquema SPS periódico, los recursos pueden mantenerse durante varias transmisiones consecutivas, lo que introduce correlación temporal en las pérdidas. Dado que el modelo propuesto utiliza una cadena simplificada de tres estados, en la que las pérdidas por half-duplex, sensado y propagación se agrupan en un único estado de pérdida no asociada a

colisión, algunas transiciones cortas entre recepción correcta, pérdida aislada y nueva recepción correcta pueden no reproducirse exactamente, aunque la forma global de la distribución se aproxima adecuadamente. En la Figura 15 se muestra esta distribución para los modelos AP1 y AP2 utilizando la misma configuración.

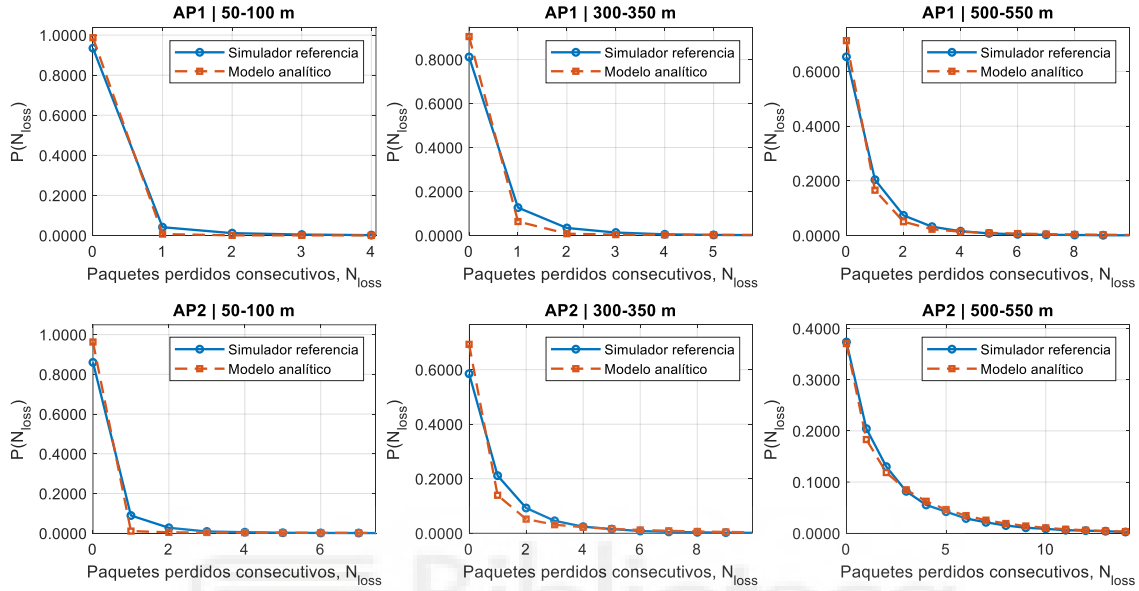


Figura 15. Distribución de pérdidas consecutivas para los modelos de tráfico AP1 y AP2.

En estos modelos de tráfico, correspondientes a tráfico aperiódico con $RRI = PDB$, el ajuste entre el modelo analítico y el simulador de referencia es similar al observado en los modelos periódicos. De nuevo, las diferencias más visibles aparecen en torno a $N_{loss} = 1$, al tratarse del valor más sensible al equilibrio entre pérdidas aisladas y pérdidas persistentes. Por último, la Figura 16 muestra esta distribución para los modelos aperiódicos AP3 y AP4 utilizando la misma configuración.

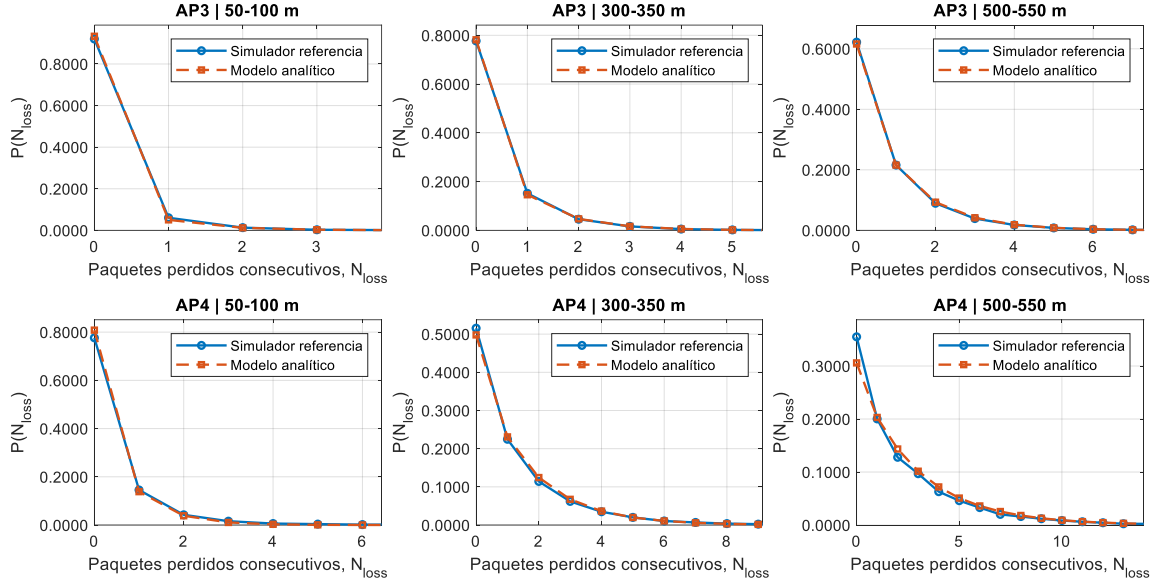


Figura 16. Distribución de pérdidas consecutivas para los modelos de tráfico AP3 y AP4.

Los modelos AP3 y AP4 representan los casos aperiódicos con $RRI = 2PDB$, por lo que incorporan reselecciones por latencia y reservas descartadas por ausencia de paquete generado. Además, AP4 constituye el caso más completo de los analizados, al añadir también tamaño de paquete variable y posibles reselecciones por tamaño. Aunque a priori presentan una dinámica de gestión de recursos más compleja, son también los casos donde se observa un ajuste especialmente bueno. Esto puede explicarse porque la mayor presencia de reselecciones introduce una renovación más frecuente del patrón de recursos, reduciendo la persistencia temporal de determinadas situaciones de pérdida. Como consecuencia, la cadena simplificada de tres estados resulta adecuada para capturar la tendencia principal de las ráfagas de pérdidas.

Una vez analizada la distribución para varios rangos de distancia, se estudia la misma métrica desde un segundo punto de vista. En lugar de representar la distribución completa para distancias concretas, se representa la probabilidad $P(N_{\text{loss}} = n)$ en función de la distancia para distintos valores de n . Esta representación permite comprobar cómo evoluciona con la distancia la probabilidad de sufrir pérdidas aisladas o ráfagas de distinta longitud. En las Figuras 17-23 se muestra la probabilidad en función de la distancia de tener n pérdidas consecutivas, para los siete modelos de tráfico usando numerología 1 y MCS=13. Concretamente, se representan los valores de $n=\{0,1,3,5,10,15\}$, para tener así una visión global de cómo se distribuye la probabilidad en gráficas de un tamaño manejable.

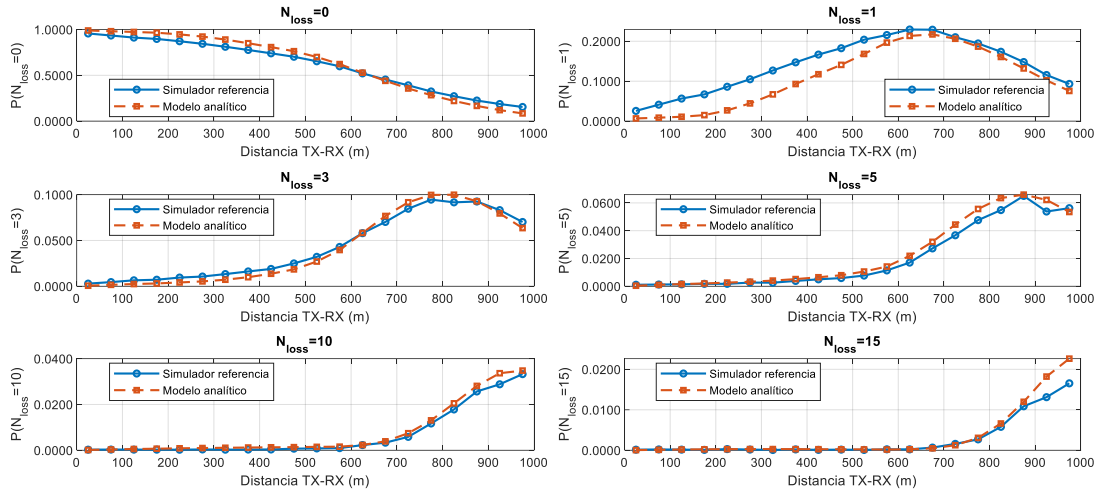


Figura 17. Probabilidad de perder N_{loss} paquetes consecutivos en función de la distancia para el modelo P1.

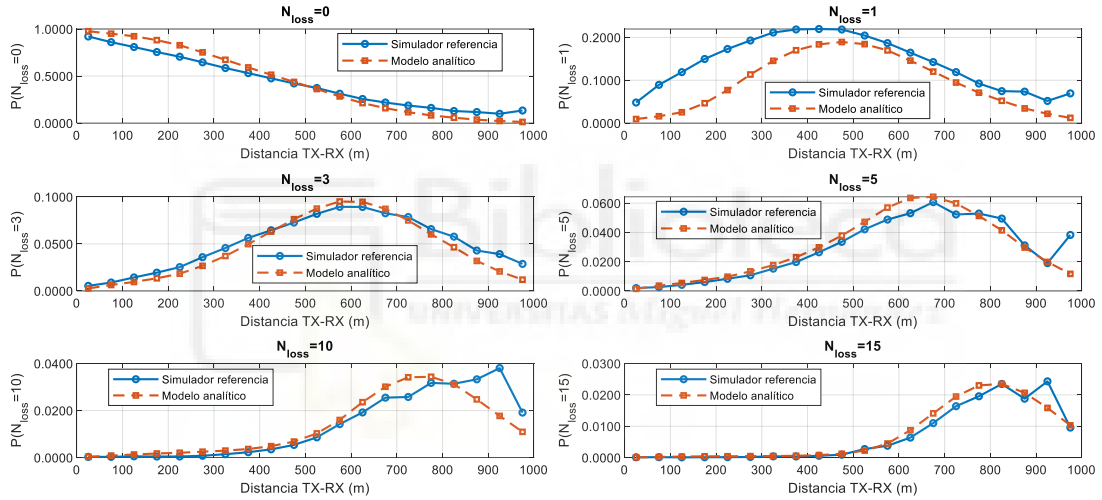


Figura 18. Probabilidad de perder N_{loss} paquetes consecutivos en función de la distancia para el modelo P2.

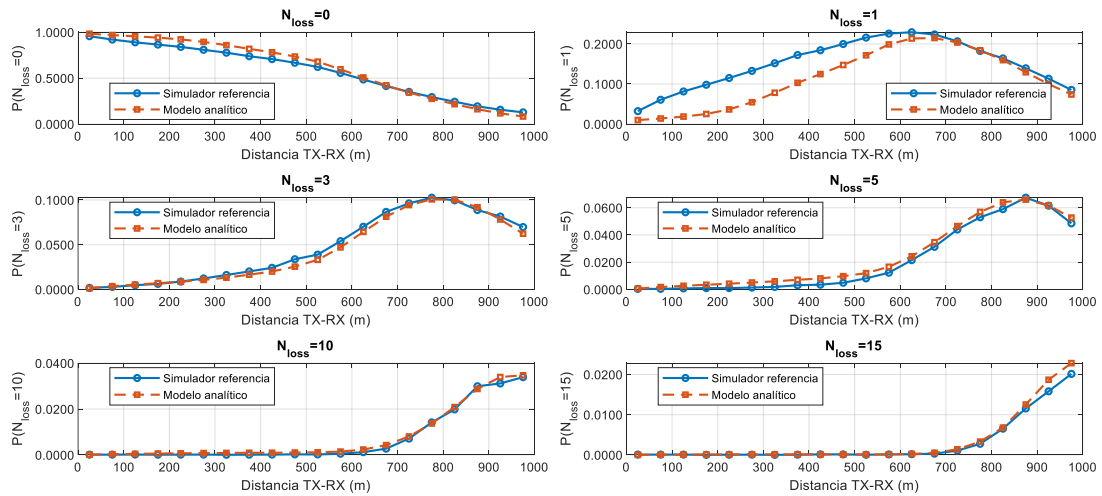


Figura 19. Probabilidad de perder N_{loss} paquetes consecutivos en función de la distancia para el modelo P3.

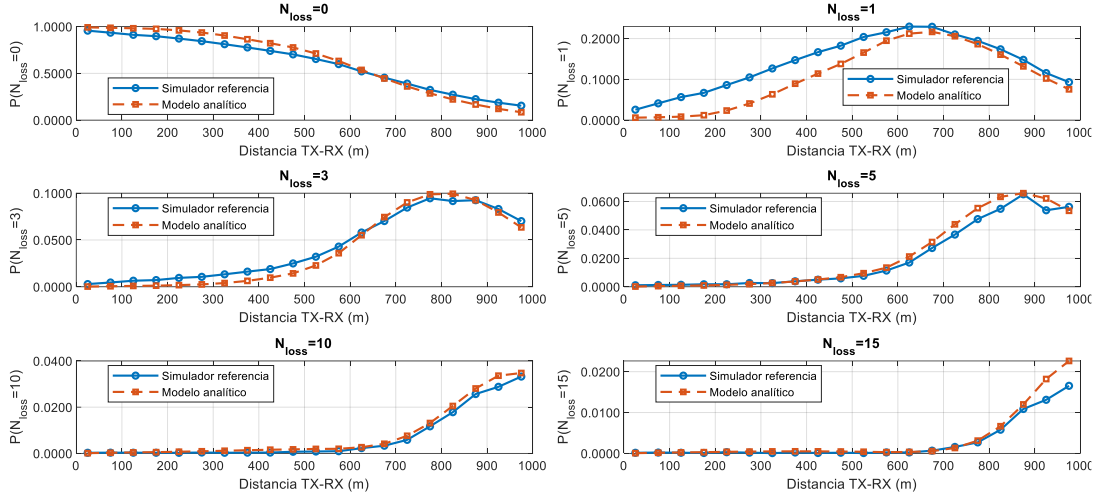


Figura 20. Probabilidad de perder N_{loss} paquetes consecutivos en función de la distancia para el modelo AP1.

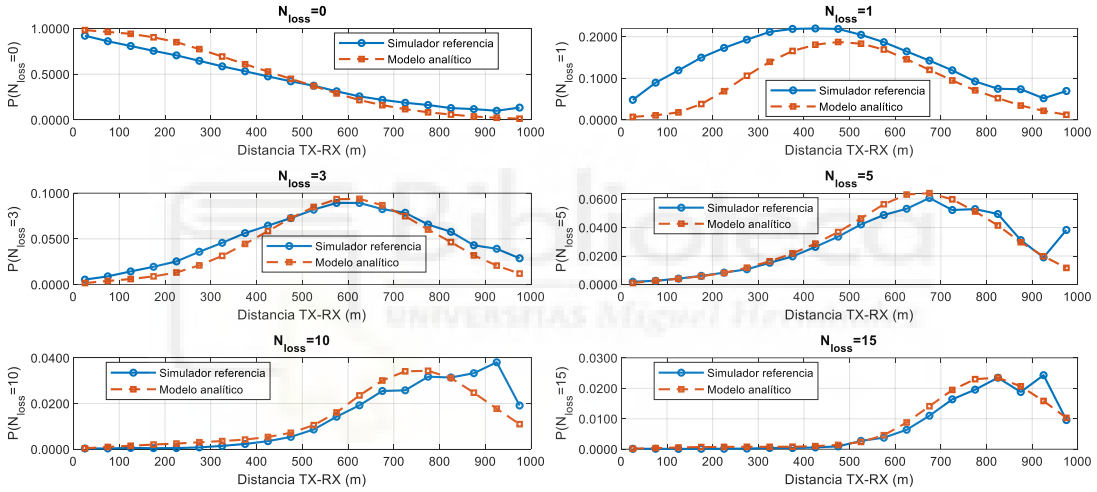


Figura 21. Probabilidad de perder N_{loss} paquetes consecutivos en función de la distancia para el modelo AP2.

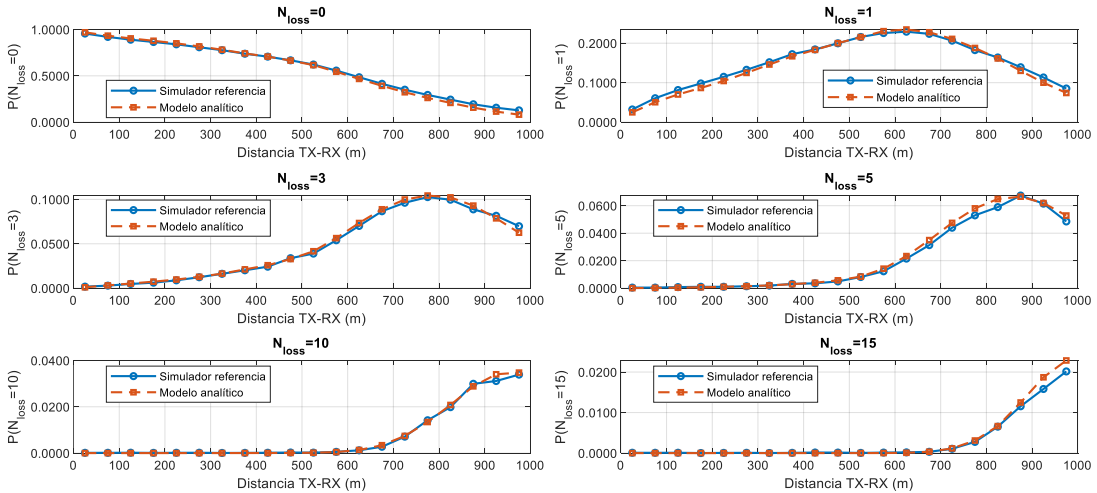


Figura 22. Probabilidad de perder N_{loss} paquetes consecutivos en función de la distancia para el modelo AP3.

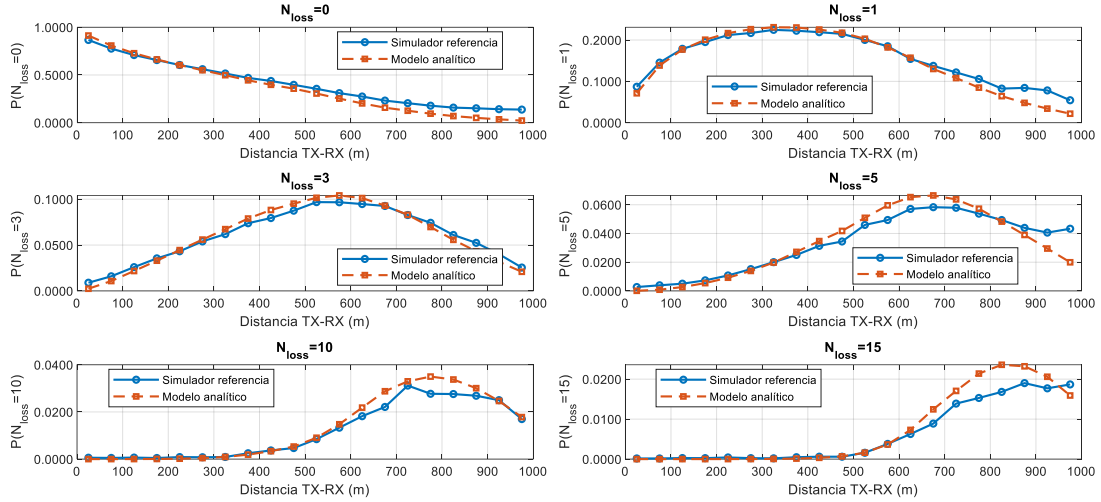


Figura 23. Probabilidad de perder N_{loss} paquetes consecutivos en función de la distancia para el modelo AP4.

Se observa, de forma general para todos los modelos, que las curvas no crecen de forma monótona con la distancia, sino que presentan un máximo en una zona concreta. Este comportamiento se debe a que cada curva representa la probabilidad de obtener exactamente un número determinado de pérdidas consecutivas entre dos recepciones correctas. A distancias reducidas, la probabilidad de pérdida es baja y, por tanto, también es baja la probabilidad de observar ráfagas de pérdidas. A medida que aumenta la distancia, aparecen más pérdidas y crece la probabilidad de obtener secuencias cortas o medias de paquetes perdidos. Sin embargo, cuando la distancia es elevada, la fiabilidad del enlace disminuye de forma más acusada y las ráfagas tienden a prolongarse. En ese caso, deja de ser probable que la secuencia termine justo después de n pérdidas, ya que la siguiente recepción correcta es menos probable, y la masa de probabilidad se desplaza hacia valores mayores de N_{loss} . Por este motivo, la probabilidad asociada a un valor concreto de N_{loss} vuelve a disminuir para distancias altas. Además, se aprecia que los máximos de las curvas se desplazan progresivamente hacia distancias más elevadas conforme aumenta N_{loss} , lo que refleja que las ráfagas largas son más probables cuando la separación transmisor-receptor es mayor. El modelo analítico reproduce adecuadamente esta evolución, capturando tanto la forma general de las curvas como la zona de distancia en la que cada valor de N_{loss} alcanza mayor probabilidad.

Finalmente, para completar la validación, se calcula el error absoluto medio entre el modelo analítico y el simulador de referencia para distintos modelos de tráfico y configuraciones radio, con el fin de realizar la validación con todos los resultados disponibles. El error se evalúa sobre los valores de N_{loss} considerados anteriormente. La métrica utilizada es el error absoluto medio, definido como:

$$MAE = \frac{1}{N_d} \sum_{i=1}^{N_d} |P_{\text{ana}}(N_{\text{loss}} = n, d_i) - P_{\text{sim}}(N_{\text{loss}} = n, d_i)| \quad (110)$$

donde $P_{\text{ana}}(N_{\text{loss}} = n, d_i)$ representa la probabilidad obtenida mediante el modelo analítico para un valor concreto de N_{loss} y una distancia d_i , mientras que $P_{\text{sim}}(N_{\text{loss}} = n, d_i)$ representa la probabilidad obtenida mediante el simulador de referencia. N_d es el número de rangos de distancia considerados en el cálculo. Esta métrica permite obtener una medida agregada del ajuste sin que el resultado quede dominado por desviaciones puntuales en un único punto de la curva. Adicionalmente, se calcula un MAE global promediando el error sobre todos los valores seleccionados de N_{loss} , de forma que se obtiene una medida resumida del ajuste general del modelo para cada configuración.

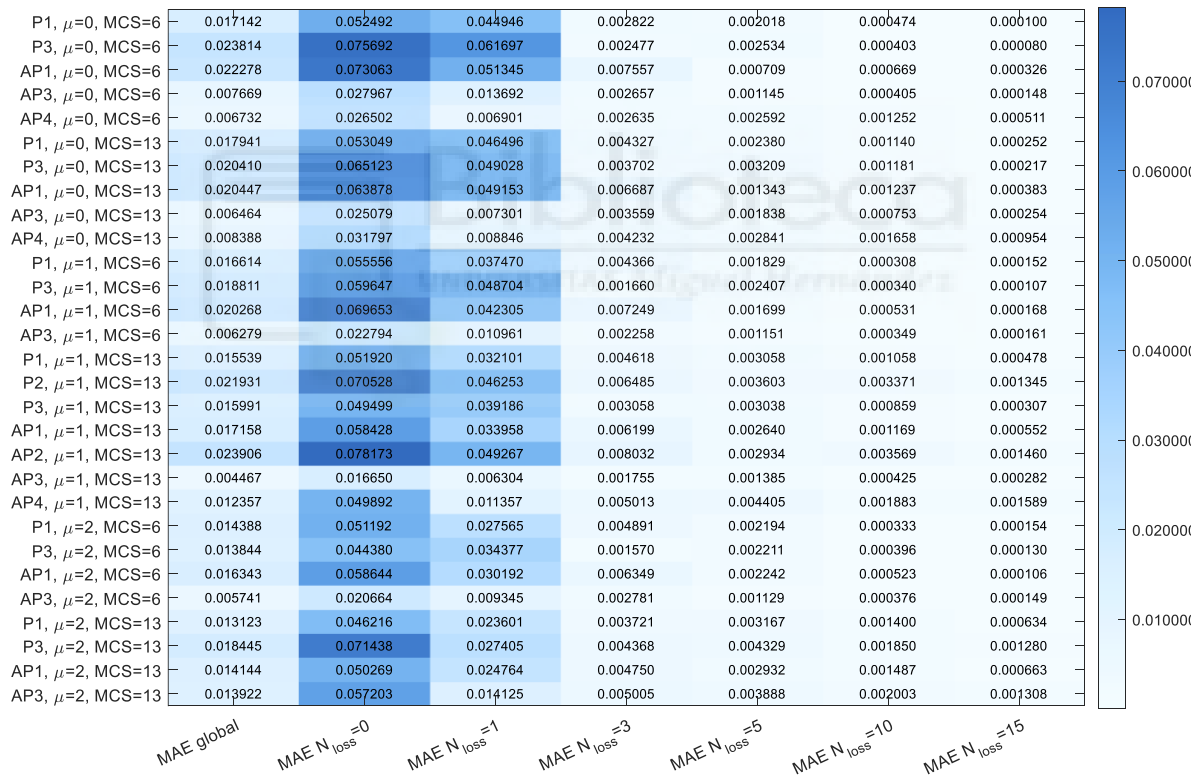
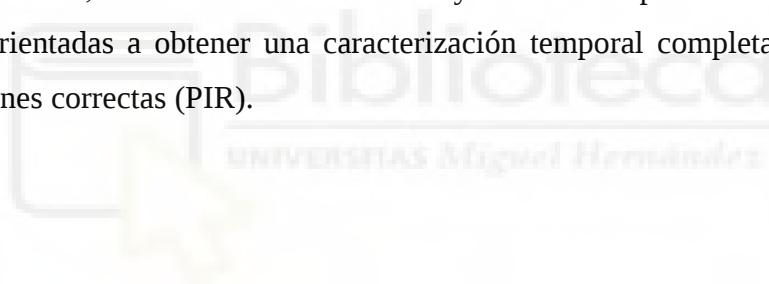


Figura 24. Error absoluto medio del Modelo Analítico de Pérdidas consecutivas respecto al simulador de referencia para los modelos y configuraciones evaluadas.

Los resultados de la Figura 24 muestran que el error medio se mantiene en valores reducidos para la mayoría de modelos de tráfico y configuraciones. Las mayores diferencias tienden a aparecer en modelos con menor probabilidad de reelección y, especialmente, para $N_{\text{loss}} = 1$, donde la distribución es más sensible a la correlación temporal de las pérdidas. En cambio, los modelos aperiódicos con mayor presencia de

reselecciones y descartes presentan, en general, un ajuste más favorable. Esto confirma la interpretación observada anteriormente: cuando el patrón de recursos se renueva con mayor frecuencia, la aproximación basada en tres estados representa mejor la evolución de las pérdidas consecutivas.

Debe tenerse en cuenta que este modelo es más compacto que los modelos de gestión de recursos y latencia validados anteriormente. En concreto, la cadena empleada agrupa las pérdidas por half-duplex, sensado y propagación en un único estado de pérdida no asociada a colisión, e introduce la persistencia de colisión mediante una probabilidad efectiva. Además, las probabilidades de pérdida utilizadas como entrada proceden de un modelo analítico auxiliar, por lo que las diferencias observadas pueden incluir tanto el error propio de la cadena de pérdidas consecutivas como el error acumulado asociado a dichas probabilidades de entrada. Este modelo se propone, por tanto, como una opción simplificada capaz de obtener resultados aproximados con un bajo coste computacional y reproducir correctamente la tendencia de las pérdidas consecutivas en los escenarios analizados. Además, esta formulación constituye una base para futuras mejoras y extensiones orientadas a obtener una caracterización temporal completa del intervalo entre recepciones correctas (PIR).



5. CONCLUSIONES Y LÍNEAS FUTURAS

En este trabajo se han diseñado y validado modelos analíticos para la gestión de recursos en 5G NR V2X Modo 2 mediante cadenas de Markov y expresiones analíticas. Los modelos desarrollados permiten caracterizar eventos que tienen un impacto directo sobre el rendimiento del sistema, como las reselecciones por contador, por tamaño de paquete, por restricciones de latencia y las reservas descartadas por ausencia de paquete generado. La incorporación de estos fenómenos resulta especialmente importante en escenarios con tráfico aperiódico y tamaño de paquete variable, donde el comportamiento del mecanismo SPS deja de ser puramente periódico y aparecen situaciones que no pueden describirse adecuadamente mediante modelos más simples.

Los resultados obtenidos muestran que el modelo analítico de gestión de recursos reproduce correctamente el comportamiento observado en simulación para las métricas principales estudiadas. Esto valida la utilidad de la cadena de Markov propuesta como herramienta para estimar probabilidades de reSelección, descarte y ocupación de subcanales sin necesidad de ejecutar una simulación detallada paquete a paquete.

También se ha demostrado que el modelo de gestión de recursos puede utilizarse como base para caracterizar métricas de rendimiento adicionales. En el caso de la latencia, las expresiones obtenidas permiten explicar de forma clara la diferencia entre los distintos modelos de tráfico. Los modelos periódicos y los modelos aperiódicos con $RRI = PDB$ presentan una distribución uniforme de la latencia, mientras que los modelos aperiódicos con $RRI = 2PDB$ generan una distribución creciente. La validación realizada confirma que el modelo analítico reproduce correctamente estas formas de distribución, lo que demuestra su utilidad para estudiar la latencia sin recurrir a simulaciones temporales extensas.

Además de la latencia, el trabajo ha permitido avanzar hacia una caracterización de la continuidad temporal de la recepción mediante el estudio de las pérdidas consecutivas entre recepciones correctas. La validación del modelo de pérdidas consecutivas muestra que, aun utilizando una cadena compacta de tres estados, es posible capturar la tendencia principal de las ráfagas de pérdidas. El modelo reproduce correctamente el desplazamiento de la probabilidad hacia valores mayores de N_{loss} conforme aumenta la distancia entre transmisor y receptor. Este resultado confirma que el modelo no solo aproxima valores medios, sino que también reproduce tendencias temporales relevantes para la fiabilidad percibida por el receptor.

Desde un punto de vista global, la principal ventaja de los modelos desarrollados es que proporcionan una herramienta de análisis rápida e interpretable. Frente a una simulación completa, que puede requerir tiempos elevados de ejecución y una configuración detallada de múltiples parámetros, estos modelos permiten evaluar de forma directa el efecto de distintos patrones de tráfico, valores de *RRI*, restricciones de latencia, tamaños de paquete y configuraciones radio. Esto resulta especialmente útil en fases iniciales de diseño, comparación de configuraciones o análisis paramétrico, donde no siempre es necesario ejecutar simulaciones completas para cada caso de estudio.

No obstante, estos modelos analíticos presentan ciertas limitaciones asociadas al ámbito de aplicación considerado, ya que no incorporan otros procedimientos avanzados del sistema, como la *re-evaluation*, la *preemption*, el control de congestión, las retransmisiones HARQ o configuraciones con adaptación dinámica del MCS. La inclusión de estos mecanismos permitiría aumentar el nivel de detalle de los modelos y constituye una línea de trabajo futuro para extender la formulación desarrollada hacia escenarios V2X más completos.

Además, otra línea de trabajo futuro especialmente directa sería mejorar el modelo de pérdidas consecutivas entre recepciones correctas. En este trabajo se ha optado por una cadena compacta de tres estados, capaz de reproducir la tendencia principal de las ráfagas de pérdidas con bajo coste computacional. Sin embargo, este modelo podría ampliarse incorporando un mayor número de estados y condiciones de transición más detalladas para representar de forma más precisa las distintas causas de pérdida, su persistencia temporal y su relación con el proceso de reelección de recursos. Esta extensión permitiría reducir las aproximaciones introducidas en el modelo actual y mejorar su ajuste frente a simulaciones más completas.

Esta mejora resulta especialmente interesante como paso previo hacia una caracterización analítica del *Packet Inter-Reception* (PIR). El modelo desarrollado permite estimar el número de paquetes perdidos entre dos recepciones correctas, pero una extensión futura podría combinar esta información con un modelo temporal de generación y transmisión de paquetes para obtener el intervalo de tiempo entre recepciones correctas. De esta forma, sería posible avanzar desde una caracterización basada en el número de pérdidas consecutivas hacia una caracterización temporal completa del PIR, métrica de gran interés para evaluar la continuidad de la recepción en comunicaciones V2X.

Por último, otra línea de trabajo futuro consistiría en aplicar los modelos desarrollados a servicios V2X más realistas, como el *Collective Perception Service* (CPS). Para ello,

habría que caracterizar primero el tráfico que genera este servicio, incluyendo los tamaños de paquete, los intervalos de generación y los requisitos de latencia. Con esos datos, se podrían aplicar los modelos analíticos y comparar sus resultados con los obtenidos mediante una simulación completa, comprobando así si siguen ofreciendo un buen ajuste.



6. REFERENCIAS

- [1] S. Singh, “Critical reasons for crashes investigated in the National Motor Vehicle Crash Causation Survey,” Traffic Safety Facts Crash•Stats, Rep. DOT HS 812 506, National Highway Traffic Safety Administration, Washington, DC, USA, Mar. 2018.
- [2] P. Viktor and G. Kiss, “Sensors in Self-Driving Vehicles: A Detailed Literature Review and New Trends,” Sensors, vol. 26, no. 7, Art. no. 2153, 2026, doi: 10.3390/s26072153.
- [3] ETSI, “5G; Service requirements for enhanced V2X scenarios (3GPP TS 22.186 version 16.2.0 Release 16),” ETSI TS 122 186 V16.2.0, Nov. 2020.
- [4] 5G Automotive Association, “C-V2X use cases and service level requirements, Vols. I–III,” Technical Report, Jan. 2025. Available: [C-V2X Use Cases and Service Level Requirements](#).
- [5] M. H. C. Garcia, A. Molina-Galan, M. Boban, J. Gozalvez, B. Coll-Perales, T. Şahin, and A. Kousaridas, “A tutorial on 5G NR V2X communications,” IEEE Communications Surveys & Tutorials, vol. 23, no. 3, pp. 1972–2026, Third Quarter 2021, doi: 10.1109/COMST.2021.3057017.
- [6] 3GPP, “Study on enhancement of 3GPP support for 5G V2X services,” 3GPP TR 22.886 V16.2.0, Release 16, Dec. 2018.
- [7] M. Sepulcre, L. Lusvarghi, S. Alacid, M. Cuenca y J. Gozalvez, “Analytical Models of the Performance of 5G NR V2X Mode 2 Vehicular Communications”, bajo revisión.
- [8] ETSI, “5G; NR; NR and NG-RAN overall description; Stage 2 (3GPP TS 38.300 version 16.4.0 Release 16),” ETSI TS 138 300 V16.4.0, Jan. 2021.
- [9] ETSI, “5G; NR; Physical channels and modulation (3GPP TS 38.211 version 16.6.0 Release 16),” ETSI TS 138 211 V16.6.0, Aug. 2021.
- [10] ETSI, “5G; NR; Physical layer procedures for control (3GPP TS 38.213 version 16.6.0 Release 16),” ETSI TS 138 213 V16.6.0, Aug. 2021.
- [11] ETSI, “5G; NR; Physical layer procedures for data (3GPP TS 38.214 version 16.6.0 Release 16),” ETSI TS 138 214 V16.6.0, Aug. 2021.
- [12] ETSI, “5G; NR; Physical layer measurements (3GPP TS 38.215 version 16.6.0 Release 16),” ETSI TS 138 215 V16.6.0, Apr. 2023.
- [13] ETSI, “5G; NR; Medium Access Control (MAC) protocol specification (3GPP TS 38.321 version 16.6.0 Release 16),” ETSI TS 138 321 V16.6.0, Oct. 2021.

- [14] R. Molina-Masegosa, J. Gozalvez, and M. Sepulcre, "Comparison of IEEE 802.11p and LTE-V2X: An evaluation with periodic and aperiodic messages of constant and variable size," *IEEE Access*, vol. 8, pp. 121526–121548, 2020, doi: 10.1109/ACCESS.2020.3007115.
- [15] M. Gonzalez-Martin, M. Sepulcre, R. Molina-Masegosa, and J. Gozalvez, "Analytical models of the performance of C-V2X mode 4 vehicular communications," *IEEE Transactions on Vehicular Technology*, vol. 68, no. 2, pp. 1155–1166, Feb. 2019, doi: 10.1109/TVT.2018.2888704.
- [16] N. G. P. Wijesiri, J. Haapola, and T. Samarasinghe, "A discrete-time Markov chain based comparison of the MAC layer performance of C-V2X mode 4 and IEEE 802.11p," *IEEE Transactions on Communications*, vol. 69, no. 4, pp. 2505–2517, Apr. 2021, doi: 10.1109/TCOMM.2020.3044340.
- [17] M. N. Sial, Y. Deng, J. Ahmed, A. Nallanathan, and M. Dohler, "Stochastic geometry modeling of cellular V2X communication over shared channels," *IEEE Transactions on Vehicular Technology*, vol. 68, no. 12, pp. 11873–11887, Dec. 2019, doi: 10.1109/TVT.2019.2945481.
- [18] X. Gu, J. Peng, L. Cai, X. Zhang, and Z. Huang, "Markov analysis of C-V2X resource reservation for vehicle platooning," in *Proc. IEEE 95th Vehicular Technology Conference (VTC2022-Spring)*, 2022, pp. 1–5, doi: 10.1109/VTC2022-Spring54318.2022.9860899.
- [19] C. Brady, L. Cao, and S. Roy, "Modeling of NR C-V2X mode 2 throughput," in *Proc. IEEE International Workshop Technical Committee on Communications Quality and Reliability (CQR)*, 2022, pp. 19–24, doi: 10.1109/CQR54764.2022.9918559.
- [20] M. M. Saad, M. A. Tariq, A. Siddiqua, and D. Kim, "Analytical models for distributed resource allocation in NR-V2X mode 2," *IEEE Internet of Things Journal*, early access, 2023, doi: 10.1109/JIOT.2023.3328885.
- [21] M. Elsharief, S. R. Sabuj, N. Jeong, and H.-S. Jo, "Advancing NR-V2X: Evaluating mode 2 performance through analytical modeling and simulations," *IEEE Access*, vol. 13, pp. 141636–141650, 2025, doi: 10.1109/ACCESS.2025.3596946.
- [22] D. Bankov, E. Khorov, A. Krasilov, and A. Otmakhov, "Analytical model of 5G V2X mode 2 for sporadic traffic," *IEEE Wireless Communications Letters*, vol. 12, no. 8, pp. 1449–1453, Aug. 2023, doi: 10.1109/LWC.2023.3278181.
- [23] L. Lusvarghi, M. L. Merani, F. Pasquinnucci, and M. Andreani, "Dynamic scheduling in NR-V2X mode 2: Probability of successful packet delivery and packet inter-reception,"

IEEE Transactions on Vehicular Technology, vol. 74, no. 12, pp. 18353–18368, Dec. 2025, doi: 10.1109/TVT.2025.3583092.

[24] S. Zhu and S. Lin, “Analytical model of NR-V2X mode 2 with re-evaluation mechanism,” arXiv:2510.27108, Oct. 2025.

[25] M. R. Wójcik, “Notes on scale-invariance and base-invariance for Benford’s Law,” arXiv:1307.3620, 2013.

[26] A. Berger and T. P. Hill, “A basic theory of Benford’s Law,” *Probability Surveys*, vol. 8, pp. 1–126, 2011.

[27] R. Verbelen, “Phase-type distributions & mixtures of Erlangs: A study of theoretical concepts, calibration techniques & actuarial applications,” M.S. thesis, Faculty of Science, University of Leuven, Leuven, Belgium, 2013. Available: [Verbelen, R. \(2013\). Phase-type distributions & mixtures of Erlangs. A study of theoretical concepts, calibration techniques & actuarial applications.pdf.](#)

