

Infocommunications Journal

A PUBLICATION OF THE SCIENTIFIC ASSOCIATION FOR INFOCOMMUNICATIONS (HTE)

September 2025

Volume XVII

Number 3

ISSN 2061-2079

Authors, co-authors of the September 2025 issue	1
<i>PAPERS FROM OPEN CALL</i>	
Spring: Theory and an Efficient Heuristic for Programmable Packet Scheduling with SP-PIFO	<i>Balázs Vass</i> 2
Challenges of DNS in the Post-Quantum Era: Improving Security with Post-Quantum TLS	<i>Abdullah Aydeger, Sanzida Hoque, and Engin Zeydan</i> 11
OSINT Based Recognized Cyber Picture	<i>Gergo Gyebnar</i> 22
Cost Model of Mode Division Multiplexing Bidirectional Transmission System	<i>Ahmed S. Mohamed, and Eszter Udvary</i> 30
Improved Outage Probability using Multi-IRS-Assisted MIMO Wireless Communications in Cellular Blockage Scenarios	<i>Suresh Penchala, Shravan Kumar Bandari, and V.V. Mani</i> 37
Improving the Clipping-based Active Constellation Extension PAPR Reduction Technique for FBMC Systems	<i>Zahraa Tagelsir, and Zsolt Kollár</i> 45
Convolution Assisted Polar Encoder With Flexible Iterative Decoding For Forward Error Correction In Wireless Communication Incorporated In Fpga	<i>T. Ranjitha Devi, and Dr. C. Kamalnathan</i> 54
Improving Technical Reviews: Metrics, Conditions for Quality, and Linguistic Considerations to Minimize Errors	<i>Gabriella Tóth</i> 62
Attitude-driven Simultaneous Online Auctions for Parking Spaces	<i>Levente Alekszejenkó and Tadeusz Dobrowiecki</i> 73
Model-based mutation testing for Finite State Machine specifications with MTR	<i>Gábor Árpád Németh</i> 84
<i>CALL FOR PAPER / PARTICIPATION</i>	
ICIN 2026 / 29th Conference on Innovation in Clouds, Internet and Networks Athens, Greece	93
<i>ADDITIONAL</i>	
Guidelines for our Authors	92

Technically Co-Sponsored by



Editorial Board

Editor-in-Chief: PÁL VARGA, Budapest University of Technology and Economics (BME), Hungary

Associate Editor-in-Chief: LÁSZLÓ BACSÁRDI, Budapest University of Technology and Economics (BME), Hungary

Associate Editor-in-Chief: JÓZSEF BÍRÓ, Budapest University of Technology and Economics (BME), Hungary

Area Editor – Quantum Communications: ESZTER UDVARY, Budapest University of Technology and Economics (BME), Hungary

Area Editor – Cognitive Infocommunications: PÉTER BARANYI, Corvinus University of Budapest, Hungary

Area Editor – Radio Communications: LAJOS NAGY, Budapest University of Technology and Economics (BME), Hungary

Area Editor – Networks and Security: GERGELY BICZÓK, Budapest University of Technology and Economics (BME), Hungary

JAVIER ARACIL, Universidad Autónoma de Madrid, Spain

LUIGI ATZORI, University of Cagliari, Italy

VESNA CRNOJEVIĆ-BENGIN, University of Novi Sad, Serbia

KÁROLY FARKAS, Budapest University of Technology and Economics (BME), Hungary

VIKTORIA FODOR, KTH, Royal Institute of Technology, Stockholm, Sweden

JAIME GALÁN-JIMÉNEZ, University of Extremadura, Spain

Molka GHARBAOUI, Sant'Anna School of Advanced Studies, Italy

EROL GELENBE, Institute of Theoretical and Applied Informatics Polish Academy of Sciences, Gliwice, Poland

ISTVÁN GÓDOR, Ericsson Hungary Ltd., Budapest, Hungary

CHRISTIAN GÜTL, Graz University of Technology, Austria

ANDRÁS HAJDU, University of Debrecen, Hungary

LAJOS HANZO, University of Southampton, UK

THOMAS HEISTRACHER, Salzburg University of Applied Sciences, Austria

ATTILA HILT, Nokia Networks, Budapest, Hungary

DAVID HÄSTBACKA, Tampere University, Finland

JUKKA HUHTAMÄKI, Tampere University of Technology, Finland

SÁNDOR IMRE, Budapest University of Technology and Economics (BME), Hungary

ANDRZEJ JAJSZCZYK, AGH University of Science and Technology, Krakow, Poland

GÁBOR JÁRÓ, Nokia Networks, Budapest, Hungary

MARTIN KLIMO, University of Zilina, Slovakia

ANDREY KOUCHERYAVY, St. Petersburg State University of Telecommunications, Russia

LEVENTE KOVÁCS, Óbuda University, Budapest, Hungary

MAJA MATIJASEVIC, University of Zagreb, Croatia

OSCAR MAYORA, FBK, Trento, Italy

MIKLÓS MOLNÁR, University of Montpellier, France

SZILVIA NAGY, Széchenyi István University of Győr, Hungary

PÉTER ODRY, VTS Subotica, Serbia

JAUELICE DE OLIVEIRA, Drexel University, Philadelphia, USA

MICHAL PIORO, Warsaw University of Technology, Poland

GHEORGHE SEBESTYÉN, Technical University Cluj-Napoca, Romania

BURKHARD STILLER, University of Zürich, Switzerland

CSABA A. SZABÓ, Budapest University of Technology and Economics (BME), Hungary

GÉZA SZABÓ, Ericsson Hungary Ltd., Budapest, Hungary

LÁSZLÓ ZSOLT SZABÓ, Sapientia University, Tirgu Mures, Romania

TAMÁS SZIRÁNYI, Institute for Computer Science and Control, Budapest, Hungary

JÁNOS SZTRIK, University of Debrecen, Hungary

DAMLA TURGUT, University of Central Florida, USA

SCOTT VALCOURT, Roux Institute, Northeastern University, Boston, USA

JÓZSEF VARGA, Nokia Bell Labs, Budapest, Hungary

ROLLAND VIDA, Budapest University of Technology and Economics (BME), Hungary

JINSONG WU, Bell Labs Shanghai, China

KE XIONG, Beijing Jiaotong University, China

GERGELY ZÁRUBA, University of Texas at Arlington, USA

Indexing information

Infocommunications Journal is covered by Inspec, Compendex and Scopus.

Infocommunications Journal is also included in the Thomson Reuters – Web of Science™ Core Collection, Emerging Sources Citation Index (ESCI)

Infocommunications Journal

Technically co-sponsored by IEEE Communications Society and IEEE Hungary Section

Supporters

FERENC VÁGUJHELYI – president, Scientific Association for Infocommunications (HTE)

The publication was produced with the support of the Hungarian Academy of Sciences and the NMHH



National Media and Infocommunications Authority | Hungary

Editorial Office (Subscription and Advertisements):

Scientific Association for Infocommunications

H-1051 Budapest, Bajcsy-Zsilinszky str. 12, Room: 502

Phone: +36 1 353 1027 • E-mail: info@hte.hu • Web: www.hte.hu

Articles can be sent also to the following address:

Budapest University of Technology and Economics

Department of Telecommunications and Media Informatics

Phone: +36 1 463 4189 • E-mail: pvarga@tmit.bme.hu

Subscription rates for foreign subscribers: 4 issues 13.700 HUF + postage

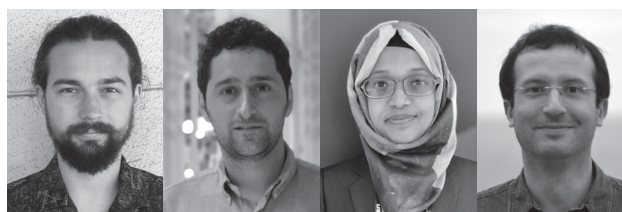
Publisher: PÉTER NAGY

HU ISSN 2061-2079 • Layout: PLAZMA DS • Printed by: FOM Media

www.infocommunications.hu

Authors, co-authors of the September 2025 issue

Balázs Vass, Abdullah Aydeger,
Sanzida Hoque, Engin Zeydan,
Gergő Gyebnár, Ahmed S. Mohamed,
Eszter Udvary, Suresh Penchala,
Shravan Kumar Bandari,
Venkata Mani Vakamulla,
Zahraa Tagelsir, Zsolt Kollár,
T. Ranjitha Devi, C. Kamalanathan,
Gabriella Tóth, Levente Alekszejenkó,
Tadeusz Dobrowiecki,
Gábor Árpád Németh



Spring: Theory and an Efficient Heuristic for Programmable Packet Scheduling with SP-PIFO

Balázs Vass, *Member, IEEE*

Abstract—Theoretical hardware model Push-In First-Out (PIFO) is used for programmable packet scheduling, allowing for flexible and dynamic reconfiguration of scheduling policies. SP-PIFO (Strict Priority-PIFO), on the other hand, is a practical emulation of PIFO that can be easily implemented using standard P4 switches. The efficiency of SP-PIFO relies on a heuristic called Push-Up/Push-Down (PUPD), which dynamically adjusts the mapping of input packets to a fixed set of priority queues in order to minimize scheduling errors compared to an ideal PIFO. This paper presents the first formal analysis of the PUPD algorithm. Our analysis shows that as more priority queues are added to the system, the ability of PUPD to emulate an optimal PIFO model decreases linearly. Based on this finding, we propose an optimal offline scheme that can determine the optimal SP-PIFO configuration in polynomial time, given a stochastic model of the input. Additionally, we introduce an online heuristic called Spring, which aims to approximate the offline optimum without requiring a stochastic input model. Our simulations demonstrate that Spring can improve the performance of SP-PIFO by a factor of 2x in certain configurations.

Index Terms—programmable packet scheduling, P4, SP-PIFO, competitive analysis, polynomial algorithm, online heuristic

1. INTRODUCTION

TRADITIONALLY, fixed-function switches implement specific network protocols engraved into the hardware. More recently, programmable switches have emerged, allowing network operators to refine on-the-fly the packet processing functionality, including header parsing and forwarding policies, using a high-level programming language like P4 [2]. However, packet scheduling in P4 switches has remained mostly fixed until recently. Push-In First-Out (PIFO) was the first hardware abstraction that, theoretically, enabled programming new scheduling algorithms into a switch without changing the hardware layout [3], [4]. In PIFO, each packet is assigned a *rank*, and the switch maintains packets in the hardware queues sorted by their rank (see Fig. 1). Then, different scheduling policies, like SRPT (Shortest Remaining Processing Time) or STFQ (Start-Time Fair Queueing) [4], can be implemented on top of PIFO by changing the way ranks are assigned to packets.

Lacking a viable hardware realization, PIFO had been considered mostly a theoretical possibility until the appearance of Strict Priority PIFO (SP-PIFO, [5]). SP-PIFO approximates PIFO by maintaining a set of strict priority (SP) queues

B. Vass is with High Speed Networks Laboratory (HSNLab), Department of Telecommunications and Media Informatics, Faculty of Electrical Engineering and Informatics (VIK), Budapest University of Technology and Economics, and Faculty of Mathematics and Computer Science, Babes-Bolyai University, Cluj Napoca, Romania. (E-mail: balazs.vass@tmit.bme.hu.) An earlier version of this paper appeared on IEEE INFOCOM Workshop on Networking Algorithms (WNA) 2022 [1].

DOI: 10.36244/ICJ.2025.3.1

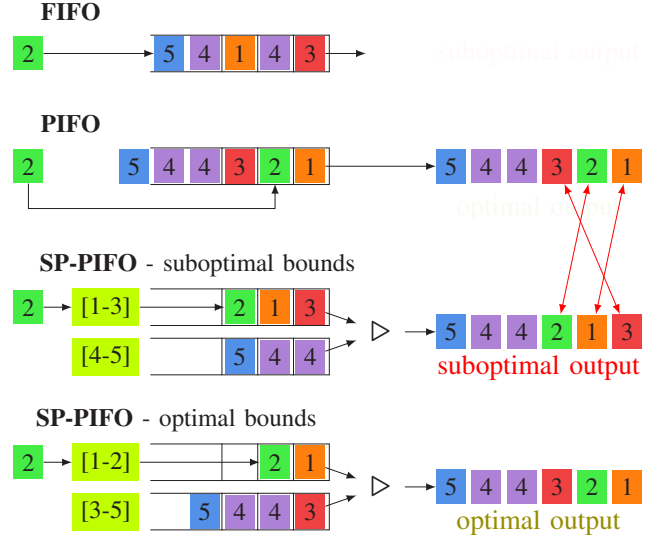


Fig. 1: Programmable scheduling: FIFO, PIFO, and SP-PIFO. The packets arrived in the order they are enqueued in FIFO (right to left). Trivially, FIFO does not sort, while PIFO releases the smallest ranked packets first. To avoid unnecessary rank inversions, the queue bounds of SP-PIFO must be optimised.

and dynamically adapting the mapping between packet ranks and SP queues. The objective is to minimize the number *inversions*, where inversions connote the event that a high-rank (i.e., ‘low-importance’) packet precedes a low-rank (i.e., ‘very important’) packet during dequeuing. As such, the inversion count effectively measures the rate of ‘scheduling mistakes’ relative to an ideal PIFO implementation (see Fig.1 again). SP-PIFO maps packets to queues by maintaining a *queue bound* for each queue, which determines the smallest packet rank that can be assigned to the queue. Then, the mapping is adapted by dynamically adjusting the queue bounds in concert with the ingress traffic ranks using the Push-Up/Push-Down (PUPD) heuristic (see later). PUPD can be implemented in P4 and thus can run inside the data plane at line rate.

The efficiency of SP-PIFO is ultimately contingent on the capacity of PUPD to adapt the way packets are assigned to queues *quickly* and *efficiently*. Unfortunately, as of now, no thorough formal analysis of PUPD is available, nor is it known whether more efficient algorithms could be defined to drive SP-PIFO. Our primary goal in this study is to fill this gap.

The main contributions are as follows. After a quick recap on programmable scheduling (Sec. 2), we present the first *competitive analysis of PUPD* (Sec. 3). Our results are mostly

negative: we show that the number of inversions produced by PUPD can be n times worse than that of a hypothetical ‘clairvoyant’ optimal bound-adaptation scheme, where n is the number of SP queues. This result suggests that the capacity of PUPD to adapt to yet unknown future ranks is limited, compared to an ideal bound-adaptation algorithm that ‘knows the future’, and the optimality gap increases linearly with the number of queues in SP-PIFO.

Driven by this observation, our next goal is to define an optimal algorithm. Our result in this context is a *polynomial-time offline algorithm*, which, given the probability distribution of ranks in incoming packets, computes a set of optimal queue bounds to minimize the number of inversions (Sec. 4).

Given that the rank distribution is hard to fix offline, our third contribution is a *fast online algorithm* that dynamically ‘learns’ the rank distribution and adaptively adjusts the SP-PIFO queue bounds following the learned distribution (Sec. 5).¹ Our evaluations (Sec. 6) suggest that the new algorithm can approximate the optimal queue bounds more efficiently than PUPD. We conclude the paper in Sec. 8.

2. BACKGROUND

Suppose input packet ranks are taken from the interval $[1, k]$ and assume we are given n SP queues $Q[i] : i \in [1, n]$. The main idea in SP-PIFO is to maintain a queue bound q_i with respect to each queue so that all packets with rank from q_i to $q_{i+1} - 1$ are assigned to $Q[i]$ (assuming an imaginary q_{n+1} equaling $k + 1$), and adapt the bounds q_i as follows:

- At the beginning, all queue bounds are 0: $\forall i \in [1, n] : q_i = 0$.
- An incoming packet with rank r is enqueued into queue $Q[i]$ if $r \geq q_i$ and i is maximal among the queues with this property, and q_i is immediately set to r (*push-up*).
- If no such i exists (i.e., $p < q_1$), then r is enqueued to $Q[1]$ and all queue bounds are decreased by $q_1 - r$ (*push-down*).
- SP queues are drained in strict priority order: packets from $Q[1]$ are dequeued first; if $Q[1]$ is empty, then $Q[2]$ is dequeued next, etc.

PUPD is appealing for a number of reasons. First, it closely emulates the ideal PIFO behaviour in that it tends to assign low-rank (i.e., ‘important’) packets to high-priority queues and high-rank (i.e., ‘low importance’) packets to low-priority queues so that the packet sequence leaving the SP queues is approximately sorted by rank. Second, PUPD can be implemented in P4 entirely, maintaining the queue bounds in P4 registers. Hence, bound adaptation in PUPD occurs after processing each packet, at line rate. Third, PUPD can dynamically adapt the bounds to even potentially rapidly changing input rank patterns. Experimental evaluations show that PUPD can produce consistently good performance under a wide range of operational conditions [5].

The efficiency of SP-PIFO to emulate an optimal PIFO model is contingent on its capacity to map low-rank packets

to high-priority queues consistently, and this is ultimately dependent on the efficiency of PUPD to adapt queue bounds to prevent ‘scheduling mistakes’. Here, any deviation of the output sequence of SP-PIFO from an ideal PIFO sequence, which is strictly sorted by packet rank, is considered an error.

3. A COMPETITIVE ANALYSIS OF PUPD

Due to the limited lookahead available in the P4 data plane pipeline, SP-PIFO bound-adaptation is severely hampered by the unpredictability of the ranks of future packets. Competitive analysis is a methodology to analyze such *online* algorithms by comparing the performance to an optimal *offline* scheme that can access the entire sequence of future requests in advance. The *competitive ratio* is defined as the quotient of the worst-case error produced by the online and the offline algorithm over an adversarial input. The smaller the competitive ratio, the less the online algorithm is hampered by the unavailability of future requests and the closer the algorithm is to ‘optimality’. Contrariwise, a relatively huge competitive ratio suggests that the performance penalty of the online setting can be prohibitive. The results below suggest that for PUPD, this is indeed the case in the case of *deterministic* and *stochastic* input packet sequences alike.

A. Deterministic competitive ratio

Given a *deterministic packet sequence* S , we measure the total error over S as the number of inversions via a simplified metric that enumerates only the intra-queue inversions, defined as follows:

$$U_{\text{det}}(S) := \#(\{a, b\} \subseteq S : a < b \text{ and} \\ a \text{ and } b \text{ enqueued to the same } Q[i] \text{ and} \\ a \text{ is the next packet enqueued in } Q[i] \text{ after } b) . \quad (1)$$

We will show that there is a family of packet sequences S , for which PUPD produces n times more inversions than the optimal offline algorithm in terms of the error function $U_{\text{det}}(S)$, where n is the number of SP queues. Intuitively, PUPD is getting further from the optimum as we add more queues, not being able to fully utilize the additional flexibility yielded by a higher number of queues. Unfortunately, this is the case even if the number of distinct ranks k is just one more than n (for $k = n$, the problem can be solved trivially by assigning each rank to a separate queue, ordered priority-wise; PUPD is intuitively close to this solution, as no push-down can happen).

Theorem 1: Given n queues and $k \geq n + 1$ input ranks, the competitive ratio of PUPD is at least n in terms of error function (1).

Proof: In our construction, we will set the maximum rank to $k := n + 1$. We will use an arbitrary positive integer l . Let the input rank sequence be as follows (with the first packet to arrive on the right):

$$S = l \times [n, \dots, 2, 1, 2, \dots, n, n + 1] .$$

Here, given a packet sequence $[Seq]$, sequences $l \times [Seq]$ denotes the l -time repetition of $[Seq]$.

¹Compared to the former version [1], the paper now includes a proof of (exponential) stability of a theoretical algorithm in a somewhat relaxed problem setting. The algorithm is very close to our offered Spring heuristic.

With this, the queues built by the PUPD are as follows (first packet enqueued to and dequeued from right):

$$Q_{\text{PUPD}}[i] = l \times [i, i+1], \quad \forall i \in \{1, \dots, n\}.$$

This means a number of l inversions in each queue, that sums up to nl for the PUPD.

On the other hand, by setting constant offline queue bounds to $\underline{q} = [q_1 = 2, q_2 = 3, \dots, q_n = n+1]$, the packet sequences built in the queues are:

$$\begin{aligned} Q_{\text{offline}}[1] &= l \times [2, 1, 2], \\ Q_{\text{offline}}[i] &= (2l) \times [i+1], \quad \forall i \in \{2, \dots, n-1\}, \\ Q_{\text{offline}}[n] &= l \times [n+1]. \end{aligned}$$

In this offline setting, we have l inversions in the first queue, and none in the other queues.

The number of inversions made by the PUPD divided by those made by the offline algorithm is n for any value of l . This means that the competitive ratio of the PUPD is not better than n on the number of inversions, i.e., error function (1).

For any higher possible maximum rank value $k' > k$, the claimed lower bound on the competitive ratio will automatically apply since the same input sequence S is feasible for k' too. ■

B. Randomized competitive ratio

The above competitive analysis was specified in terms of a deterministic adversarial packet sequence S . However, assuming an adversary can precisely control the succession of packet ranks as received by a P4 switch is somewhat unrealistic. Below, we turn to a weaker adversary model, where the adversary can choose the probability distribution of the packet ranks $\mathcal{P}$, but the exact succession of packet ranks S is not under control. Assuming that packet rank probabilities $\mathcal{P}$ across the input sequence are *i.i.d.*, the objective is to minimize the expected number of inversions in terms of $\mathcal{P}$ can be expressed as:

$$U_{\text{stoch}}(\mathcal{P}) = \mathbf{E}_{\mathcal{P}}(U_{\text{det}}(S)). \quad (2)$$

Theorem 2: Given n queues, $k \geq n+1$ input ranks, and a stationary packet rank probability distribution $\mathcal{P}$, the competitive ratio of PUPD is not better than n , in terms of error function (2).

Proof: Just as in the proof of Thm. 1, first, we will set $k := n+1$. We divide the input rank sequence into *phases*, where phase i starts when the i^{th} rank-1 packet arrives. We construct a packet rank distribution $\mathcal{P}$ such that, in each phase, PUPD makes n inversions *with high probability* (WHP), while with $\underline{q}_{\text{offline}} = [1, 3, 4, 5, \dots, n+1]$, 1 inversion per phase incurs WHP. We note that, in the proof of Thm. 1 (forged for the deterministic setting), S can be divided similarly.

We can observe that the arrival time T_i of the next rank- i packet has a $\text{Geo}(p_i)$ distribution, where p_i denotes the probability that the next packet will be a rank- i packet. Thus, $\mathbf{E}(T_i) = 1/p_i$. Let $p_1 = \epsilon_1$, $p_{n+1} = \epsilon_{n+1}$, and $p_i = \frac{1-\epsilon_1-\epsilon_{n+1}}{n-1}$ for $i \in \{2, \dots, n\}$. Suppose both $\frac{1/\epsilon_{n+1}}{n^2} > C$ and $\frac{1/\epsilon_1}{1/\epsilon_{n+1}} > C$,

for some sufficiently large constant C . With this, a rank $i \in \{2, \dots, n\}$ packet is much more likely to arrive than a rank- $(n+1)$ one, that is, on its turn, much more likely to arrive than a rank-1 one².

After an initial period, when, finally, a monotone decreasing subsequence $(n+1), n, \dots, 2$ can be chosen out of the packets arrived so far, the bounds of the PUPD are set to $\underline{q}_{\text{PUPD}} = [2, \dots, n+1]$. On average, this happens after less than $n^2 + 1/\epsilon_{n+1}$ packets. The probability of a rank-1 packet arriving in this time frame is very low. Eventually, a rank-1 packet will arrive, marking the start of a *phase*, and causing a push-down (setting $\underline{q}_{\text{PUPD}}$ to $[1, \dots, n]$), and an inversion in $Q[1]$. WHP, this is followed by an inversion in each other queue $Q[i]$ because of enqueueing a newly arriving rank- i packet in it. Once, a rank- $(n+1)$ packet will arrive, setting q_n to $(n+1)$, and initiating a series of push-ups, WHP leading back to $\underline{q}_{\text{PUPD}} = [2, \dots, n+1]$. Then, the arrival of a rank-1 packet will start a new phase. It is easy to see that PUPD commits n inversions in each phase, WHP.

On the other hand, static setting with queue bounds $\underline{q}_{\text{offline}} = [1, 3, 4, 5, \dots, n+1]$ makes at most 1 inversion per phase, that incurs when the rank-1 packet arrives. This completes the proof for $k \geq n+1$.

The lower bounds on competitiveness ratios for any possible maximum rank value $k' > n+1$ will automatically apply since the same input sequences as used above are feasible for k' too. ■

Note that while the *phases* of the input sequence defined in the proof above may be long, inducing low rates of inversions, if the parameters of the sequence are pushed to the limits; concentrating here on a ‘very low’ arrival rate p_1 , and a comparably ‘much lower’ p_{n+1} . Setting p_1 and p_{n+1} to less extreme values can push the inversion rates to considerable levels, although such settings might yield a somewhat smaller lower bound on the stochastic competitiveness ratio of PUPD. Conducting such a nuanced stochastic competitiveness analysis is beyond the scope of this study, and could be part of a follow-up work.

4. STOCHASTIC OFFLINE OPTIMUM

We have seen that, even in the stochastic model, PUPD can perform n times worse compared to a simple setting with static queue bounds. Below, we will show that the expected cost of inversions defined by (2) for constant queue bounds can be minimized in polynomial time. The first observation is that, given n queues and k ranks, with bounds fixed at $\underline{q} = [q_1 = 1, q_2, q_3, \dots, q_n]$ extended with an imaginary queue bound $q_{n+1} = k+1$, the expected error is:

$$U_{\text{stoch}}^{\text{const}}(\mathcal{P}) = \sum_{i=1}^n \left(P_i \sum_{q_i \leq a < b < q_{i+1}} \frac{p_b p_a}{P_i^2} \right), \quad (3)$$

where $P_i = \sum_{j=q_i}^{q_{i+1}-1} p_j$ denotes the probability that the next packet will be enqueued to queue i . Eq. (3) holds since the following. In the fraction, we multiply two probabilities: 1) the

²More formally, when we say here some event happens WHP, we claim they happen with probability 1 supposing $C \rightarrow +\infty$.

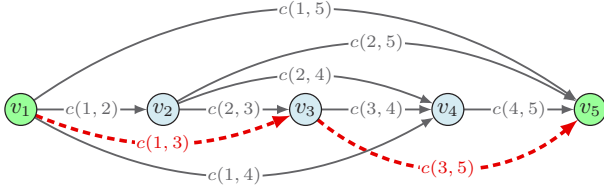


Fig. 2: Example of a DAG representing a possible configuration for $k = 4$. The highlighted path encodes the bounds for a 2-queue SP-PIFO setting ($n = 2$), with the bounds being $q_1 = 1$ and $q_2 = 3$.

probability of the rank of the latest enqueued packet being b is p_b/P_i , and 2) the probability of the next rank in the queue being a (i.e., p_a/P_i). The sum of these values for all the possibly inverted rank pairs in queue i (for b and a s.t. $q_i \leq a < b < q_{i+1}$), when multiplied with the incoming packet intensity P_i and summed over all i , yields (3).

Theorem 3: Let $\mathcal{P}$ the probability distribution of packet ranks and suppose that $\mathcal{P}$ is stationary. Then, the total expected inversion cost in terms of (3) can be minimized in $O(k^2n)$ time.

Proof: We construct a directed acyclic graph $D(V, A, c)$, where the set of nodes $V = \{v_1, \dots, v_{k+1}\}$ corresponds to the set of possible rank values $\{1, \dots, k\}$ associated with an auxiliary node v_{k+1} , the arc set A stands for arcs $(v_i, v_j) : 1 \leq i < j \leq k+1$, and cost $c(v_i, v_j)$ of an arc is the expected error in the queue it represents whenever exactly those packets are enqueued in queue i that have ranks from $\{i, \dots, j-1\}$ (see Fig. 2).³

We claim that, for all arcs $(v_i, v_j) \in A$, their costs $c(v_i, v_j)$ (measured in terms of (3)) can be determined in $O(k^2)$ total time. For this, given a fixed lower bound a , we can determine $c(a, a+1)$, $c(a, a+2)$, $\dots$, $c(a, k+1)$, each after, using $O(1)$ basic arithmetic operations per upper bound, where, for $i \geq a+2$, $c(a, i)$ is calculated using the previously calculated cost $c(a, i-1)$ and sum $\sum_{j=a}^{i-1} p_j$. Since there are $O(k^2)$ lower bound–upper bound pairs, the complexity follows. In conclusion, $D(V, A, c)$ can be determined in $O(k^2)$.

Next, we argue that queue bounds set to the node indexes appearing in a shortest $v_1 - v_{k+1}$ path with (at most) n arcs are optimal. Luckily, such a shortest path can be computed in $O(k^2n)$, e.g., with the help of a slight variation of the Bellman-Ford algorithm [6], see Algorithm 1. The proposition we take advantage of here is that, for any input graph, after i repetitions of the outermost for loop of the Bellman-Ford algorithm, the computed $s-t$ path is a shortest $s-t$ path with at most i edges, if there is an $s-t$ path with at most i edges at all (otherwise it just returns an error and an infinite cost). Applied to our case, a shortest $v_1 - v_{k+1}$ path in $D(V, A, c)$ constructed by the Bellman-Ford algorithm in n iterations of the outer for loop suits our needs. Since each iteration takes $O(k^2)$, the runtime complexity of the algorithm follows. ■

We note that the above algorithm works for any conservative cost function c . Thus, the minimization of the expected error

³If the expected error is chosen to be measured as defined in (3), then $c(a, b)$ equals $P \sum_{a \leq c < d < b} \frac{p_c p_d}{P^2}$, where $P = \sum_{j=a}^{b-1} p_j$.

Algorithm 1: Modified Bellman-Ford algorithm for the stochastic offline optimum

```

Input:  $D(V, A, c)$ 
for  $v \in V$  do
    1 |  $\text{cost}[v] := \infty$ ;  $\text{predec}[v] := \text{null}$ 
end
2  $\text{cost}[v] := 0$ 
for  $i = 1..n$  do
    for  $(v_a, v_b) \in A$  do
        3 | if  $\text{cost}[v_a] + c(v_a, v_b) < \text{cost}[v_b]$  then
            4 |  $\text{cost}[v_b] := \text{cost}[v_a] + c(v_a, v_b)$ 
               $\text{predec}[v_b] := v_a$ 
        end
    end
end
return  $v_1 - v_{k+1}$  path built from list  $\text{predec}$  starting at  $v_{k+1}$ 
    
```

can be done in polynomial time for any cost function c' that is conservative (e.g., non-negative), and polynomially computable. However, the computational complexity for the general case is infeasibly high. To this end, we mention that some related cost formulations meet the convex or concave *Monge* properties [7]. Both cases yield low optimization complexities: [7] shows that finding a cheapest n -link path in a complete DAG with the cost function fulfilling the concave Monge property can be done in $O(k\sqrt{n \log k})$. Then, [8] offers an algorithm that solves the same problem in $k2^{O(\sqrt{\log n \log \log k})}$, if $n = \Omega(k)$. Note that these complexities do not include determining the necessary cost values.

Finally, in the scope of this paragraph, revisiting the somewhat unrealistic adversary that can precisely control the succession of packet ranks as received by the P4 switch, one can see that with *deterministic* packet sequences, an offline stochastic optimum could be tricked to commit n times more inversions than an offline (deterministic) optimum, e.g., on the following sequence: $S = l \times [2n, 2n-1] + \dots + l \times [4, 3] + l \times [2, 1]$ (first packet to arrive on the right, $+$ means concatenation). Note that to lead astray a stochastic offline optimum this heavily, we needed $k = 2n$ ranks, almost twice as many as for PUPD. In fact, a more nuanced claim would be that using $k' \in \{n+1, \dots, 2n\}$ ranks, a stochastic optimum could commit $(k'-n)$ times more inversion than the offline optimum. For this, a possible input sequence S' could be similar to the above S , but some of the neighboring ranks in S had to be mapped together. Thus, in the deterministic setting, a stochastic optimum seems to degrade more gracefully compared to the PUPD in terms of the number of tanks k . Conducting a more nuanced deterministic competitiveness analysis for the stochastic optimum is beyond the scope of this study.

5. APPROXIMATING THE OPTIMAL STATIC BOUNDS ONLINE IN CONSTRAINED SPACE

Unfortunately, Algorithm 1 cannot be implemented in real P4 switches. First, an offline algorithm needs the rank distribution $\mathcal{P}$ that is not available in a switch. We solve this problem by learning the rank distribution *online*. Second, P4 switches do not have enough stateful memory to learn the empirical packet rank distribution as this would require $\Theta(k)$ space.

We solve this problem by offering an alternative algorithm that needs only $\Theta(n)$ memory, i.e., the space requirement is proportional to the number of queues, not the number of ranks (which is usually much larger). Of course, shrinking the algorithm's memory footprint results in a loss of optimality. In the evaluations (Sec. 6), we will show that the price we pay for reducing the memory is not prohibitive.

Consider a simplified error function, where the objective is to minimize the maximum per-queue error, instead of the sum of errors:

$$U_{\text{stoch}}^{\max}(\mathcal{P}) = \max_{i=1,\dots,n} \left(P_i \sum_{q_i \leq a < b < q_{i+1}} \frac{p_b p_a}{P_i^2} \right). \quad (4)$$

One can observe that minimizing (4) is a special case of the *sequence partitioning problem*, which can be solved in $O(n(k-n))$ time [9]. Also, the space needed for this is reduced to $O(k)$, thanks to the simplicity of the modified objective function (*min-max* instead of *min-sum*).

Recall that our aim is to construct an algorithm that fits into $O(n)$ memory. To this end, we need to take care of the space needed to store the learned rank distribution. As a simplification, our objective will be to balance the load on the queues, without caring about the rank distribution inside the rank interval assigned to the queue. In other words, in addition to the queue bounds, we only remember the probability P_i of the next packet being enqueued to queue i . Intuitively, by minimizing the maximum of the P_i values, we even out the load on the queues ($P_i \simeq 1/n$). This will translate to reasonably low inversion rates, measured according to (4), as also reflected in our evaluation results from Sec. 6.

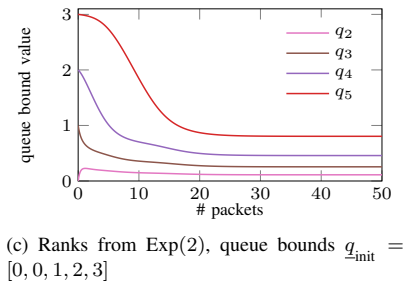
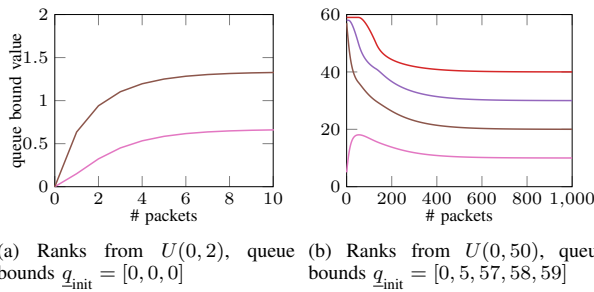


Fig. 3: Queue bounds of the continuous relaxation of the Spring over time in case of continuous rank distributions

A. Continuous relaxation

With the above simplifications, our task is now to learn the per-queue loads P_i and, meanwhile, optimize the integer queue bounds to somewhat heuristically optimize (4). To simplify the development, first, we analyse the continuous relaxation of this problem.

In the continuous model, ranks and queue bounds are real-valued. Packets arrive in infinitesimally small quanta, and so per-queue packet rates P_i can be determined for any queue bound setting. We denote the probability density function for the ranks with $f(x)$. Let the two extreme queue bounds be fixed at $q_1 = 0$ and $q_{n+1} = +\infty$. We can express P_i as $P_i = \int_{q_i}^{q_{i+1}} f(x) dx$, for each $i \in \{1, \dots, n\}$. When the optimization reaches a set of stable queue bounds, the following should hold:

$$P_{i-1} = \int_{q_{i-1}}^{q_i} f(x) dx = \int_{q_i}^{q_{i+1}} f(x) dx = P_i \quad \forall i \in [2, n]$$

For a (small) time quantum Δt , let $\Delta P_i(t)$ denote the number of packets assigned to queue i during time interval $(t - \Delta t, t]$. We define the updated value of of queue bound q_i ($i \in [2, n]$) after this Δt time as:

$$q_i(t) = q_i(t - \Delta t) + \Delta P_i(t) - \Delta P_{i-1}(t). \quad (5)$$

The optimization step, as defined above, just happens to be essentially the same as the Euler method for solving differential equations.

The next Lemma hints that under a variety of circumstances, the optimization as described above converges to the stable equilibrium point, where the loads on the queues are evened out.

Lemma 1: Suppose the following: 1) The probability density function $f(x)$ of the rank distribution is positive on an interval $[0, M]$, for some $M > 0$. 2) On $[0, M]$, the cumulative density function (CDF) of the rank distribution is strictly monotone increasing and continuously differentiable. 3) q_1 and q_{n+1} are fixed at 0 and M , respectively, at all time. 4) At time $t = 0$, we are given some values of the queue bounds $q_1(t) \leq q_2(t) \leq \dots \leq q_{n+1}(t)$. 5) For $i \in \{2, \dots, n\}$, conform with (5), we set $q'_i(t) = (F(q_{i+1}(t)) - 2 \cdot F(q_i(t)) + F(q_{i-1}(t)))$. Then, there exists a globally stable fixpoint of the system, that is $(q_1, \dots, q_{n+1}) = (0, F^{-1}(1/n), F^{-1}(2/n), \dots, F^{-1}(n/n))$. If F is the uniform distribution on $[0, M]$, then $(q_1, \dots, q_{n+1}) = (0, 1/n, \dots, n/n)$ is *exponentially stable* global fixpoint. The proof of Lemma 1 is relegated to the Appendix.

Fig. 3 shows our evaluation results of the continuous model over some famous rank distributions. Since we intend to fix q_1 at 0, it is enough to evaluate the rest of the bounds of a system. Fig. 3a shows the trajectory of the relevant 2 queue bounds for a maximal packet rank of 2 accompanied by a uniform rank distribution on $[0, 50]$. Further, Fig. 3b shows the relevant queue bounds for the maximal packet rank of 50 and a uniform rank distribution on $[0, 50]$. Here, despite the unfortunate initial queue bounds (3 out of the 4 bounds initially exceed the maximum rank), the system quickly converges to the theoretical optimum. Finally, Fig. 3c shows the evolution of queue bounds for an exponential rank distribution.

B. Online learning of the rank distribution

As indicated above, the algorithm, as described so far, should eventually converge to (or slightly oscillate around) a set of stable queue bounds, assuming the packet ranks are i.i.d., but the counters may take on high values as packets are processed. Another problem is, on the other hand, that the number of packets arriving over a short time period Δt is small, yielding imprecise empirical data on the packet rank distribution. Counting the arriving packets with *Exponentially Weighted Moving Averages* (EWMA) solves both problems at once. Let us re-discretize time and the packet arrivals: packets arrive at each positive integer time t (i.e., $\Delta t = 1$), one by one. Let $I_i(t)$ be an indicator (taking on a value of either zero or one) of whether the received packet at time t is assigned to the i -th queue. In addition, let us define a parameter $\alpha \in (0, 1)$ to control how significant new packets should be relative to the packets recorded further in the past (as well as how quickly we forget said packets). We update the EWMA based per-queue packet counters $\mu_i(t)$ on each incoming packet as follows:

$$\mu_i(t) \leftarrow (1 - \alpha) \cdot \mu_i(t - 1) + \alpha \cdot I_i(t),$$

where $\mu_i(0) \in [0, 1]$ can be set arbitrarily, in Bayesian manner. It is easy to see that $\mu_i \in [0, 1]$ holds at any time for each queue. Furthermore, using the moving averages instead of the ΔP_i values makes the random process of the changing of the queue bounds more stable.

Thus, in this more discrete setting, we rewrite (5) as follows: for all $i \in [2, n]$,

$$q_i(t) = q_i(t - \Delta t) + \mu_i(t) - \mu_{i-1}(t). \quad (6)$$

Alg. 2 summarizes our heuristic. Since the mechanics of our algorithm resemble the physical model of serially connected springs, we call our algorithm the *Spring* heuristic. In Alg. 2, we still have to re-discretize the queue bounds. Thus, in the algorithm, we keep track of a real-valued version r_i of each queue bound q_i . More precisely, the often minor adjustments of the optimization are done on the continuous r_i bounds (line 8), while the actual queue bounds q_i are the corresponding integer roundings (line 10). This way, the actual queue bounds are just coarse-grained approximations of the underlying fine-resolution representations.

Another issue is that a careless bound adjustment may result in the bound of a queue falling below that of the lower-ranked neighbour during a transient (recall Fig. 3). To avoid this problem, we have implemented an additional *collision detection* mechanism (lines 6-10), which ensures that q_i is never pushed above $q_{i+1} - 1$, or below $q_{i-1} + 1$. This addition guarantees that queue bounds remain integral and sorted in ascending order during the progression of the algorithm.

C. P4 compatibility and resource usage

The Spring heuristic should be well within the capacity of most P4-compatible devices, thus conducting a thorough P4 compatibility analysis of our heuristic was not in the main scope of this study. Regarding the semantics, we note that there are already examples of fixed point arithmetics implemented in P4 [10], and there are also existing P4 implementations of

Algorithm 2: Spring heuristic

```

// Initialization:
1  $[q_1, \dots, q_n] := [1, \dots, n]$  // queue bounds
2  $[r_1, \dots, r_n] := [1, \dots, n]$  // q. bounds lin. relaxed
3  $[\mu_1, \dots, \mu_n] := [0, \dots, 0]$  // error costs
while Packet arrives with rank  $j$  do
4   Packet enqueued into queue  $Q[j]$  s.t.  $j \geq q_i$ , where  $i$  is
   the greatest queue index satisfying this condition
5    $[\mu_1, \dots, \mu_n] := [\mu_1, \dots, \mu_n] * (1 - \alpha)$ ;  $\mu_i := \mu_i + \alpha$ 
   for  $i = n, \dots, 2$  do
6     lowerBound :=  $r_{i-1} + 1$ ; upperBound :=  $+\infty$ 
7     if  $i < n$  then
8       upperBound :=  $r_{i+1} - 1$ 
9     end
10     $r_i := r_i + \mu_i - \mu_{i-1}$ 
11     $r_i := \min\{\max\{\text{lowerBound}, r_i\}, \text{upperBound}\}$ 
12     $q_i := \text{round}(r_i)$ 
  end
end
    
```

EWMA itself [11]. In terms of memory, the Spring algorithm, as described, should not require considerably more registers than the SP-PIFO itself or the related AIFO [12], which are arguably P4 compatible.

6. EVALUATION

To make our measurements reproducible and comparable to earlier work, we reused an already existing version of NetBench [13], which contained a reference implementation of SP-PIFO. The simulations used the upstream traffic generators from NetBench, labelled ‘Uniform’, ‘Poisson’, ‘Exponential’, ‘Inv. exp.’, ‘Convex’, and ‘Minmax’, respectively, all of them generating i.i.d. integer packet ranks on interval $[0, 100]$. The Exponential and Poisson generate random numbers from distributions $\text{Exp}(1/25)$ and $\text{Pois}(50)$, respectively, and map them to integer values in $[0, 100]$. The Inv.exp. is based on the Exponential distribution, but it subtracts the generated integer from 99. The Convex distribution is based on a random variable $X \sim \text{Pois}(50)$, with the value of Convex being $(X - 1) \bmod 100$ (note that in case of the Convex, rank 99 will be the most probable). Finally, Minmax is based on a variable $Y \sim \text{Convex}$, with a value of $(Y - 10) \bmod 50$.

The configuration of the PUPD matches those used in the inversion-related measurements of [5], with a queue number of $n = 8$. The Spring heuristic uses the same parameters as PUPD, with the addition of a parameter $\alpha = 0.01$ to tune the EWMA component. As a benchmark, we also made measurements with the Greedy heuristic of [5] with basic parameters. In the long turn, with i.i.d. ranks, the bounds of Greedy are considered to converge to the optimum, but its space requirement is infeasibly high [5]. The measurements were configured with a one-second limit on the simulated runtime, resulting in around 10^6 packets.

Fig. 4 summarizes the relative performance of the heuristics in terms of the inversion count built into the NetBench implementation of SP-PIFO. We note that this inversion counter may differ from our simplified metrics, and that the number of inversions consistently exceeded 10^5 . Despite using the NetBench’s own inversion counter, Spring produced a consistently

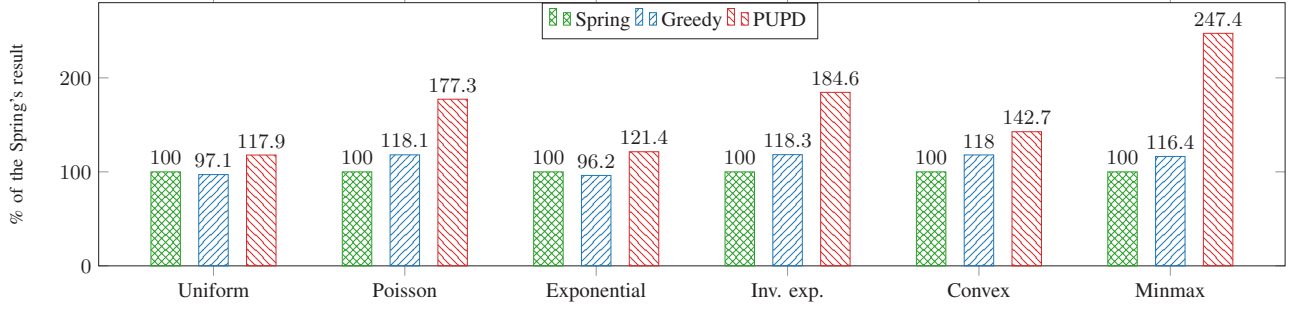


Fig. 4: Number of rank inversions of the different heuristics as percent of the Spring's results

and significantly smaller number of inversions compared to the PUPD, surpassing it on all the distributions studied, with committing only around 40% to 85% of the inversions made by the PUPD. Compared to Spring, PUPD performed the worst at the Minmax, Inv.exp., and Poisson distributions, with roughly 1.7 to 2.5 times more inversions incurred. In the case of the Uniform and Exponential, this ratio was a more moderate 1.2.

Compared to the Greedy, Spring produced a similar number of inversions, beating it on 4 out of the 6 distributions studied. Here, Spring performed the best at the Inv.exp. and Convex distributions committing around 85% of the inversions made by the Greedy, while, at worst, in case of the Uniform and Exponential distributions, this ratio was no more than 1.04.

We have also evaluated the sum of the rank differences of the inverted packet pairs. As Table I demonstrates, for most distributions, the Spring performed slightly better compared to Greedy. The extremes are the Exponential and Inv. exp., where the total inversion size of Greedy is 0.6 and 2.85 times the Spring's, respectively. PUPD is consistently worse, on average making 4.45 times more total inversion size than the Spring.

One intuitive argument on why the PUPD performed consistently the worst among the competing algorithms in our experiment is that its queue bounds are very hectic, driving it to frequent and often unnecessary inversions. See also its behavior on packet sequences built in Sec. 3. Here, a relative advantage of Spring is its relative reluctance to quickly change the queue bounds. To be fair to PUPD, one of its advantages can be that it reacts relatively effortlessly by definition in case of a quickly changing packet rank distribution. In such case, Spring's reaction time depends on parameter α defining its rate of learning.

7. DISCUSSION

Packet scheduling has been extensively studied for decades [14]–[23]. The concept of programmable scheduling

TABLE I
SUM OF THE RANK DIFFERENCE OF THE INVERTED PACKETS IN PERCENT OF THE SPRING'S RESULTS.

	Uniform	Poisson	Exponential	Inv. Exp.	Convex	Minmax
Spring	100	100	100	100	100	100
Greedy	122	158	60	285	101	107
PUPD	290	462	145	904	167	703

was introduced by [4], [24], which proposed the PIFO queue as an enabling abstraction. While promising, implementing PIFO queues in hardware proved challenging. Hence, some works have suggested new hardware designs such as PIEO [25], BMW-Tree [26], BBQ [27], and Sifter [28]. Others have focused on approximating PIFO behaviors on existing programmable data planes (using and efficiently managing a set of FIFO queues): SP-PIFO [5], QCluster [29], PCQ [30], AIFO [12], Gearbox [31], and PACKS [32]. Spring falls into this latter category. Unfortunately, a thorough formal convergence analysis of these approaches, in most cases, is hard to achieve. While the theoretically provable convergence and convergence speed of the Spring to its equilibrium point has no direct implication to its often more complex counterparts, it suggests similar convergence behaviors for similar frameworks.

8. CONCLUSION

Motivated by the industry trends towards rendering the network data plane comprehensively reconfigurable [3], in this study, we revisited the theory and algorithms for programmable packet scheduling.

We presented the first competitive analysis of heuristic PUPD presented in the SP-PIFO framework for approximating theoretical PIFO queues. We encountered a strong negative result: the PUPD may commit up to n times more packet rank inversions than inevitably needed, with both deterministic and stochastic input, where n is the number of SP queues in the system. In other words, the ability of PUPD to take advantage of an additional SP queue largely decreases, and the algorithm gets further from the optimum as the number of SP queues on disposal grows.

Motivated by this finding, we propose an algorithm to compute the optimal static queue bounds minimizing the expected rate of inversions in case of a given rank distribution. The algorithm runs in polynomial time: its complexity is $O(k^2n)$, where k is the maximum rank.

Considering the online setting, a new online bounds adaptation heuristic called *Spring* was also proposed. Crucially, Spring is easy to reason about formally, which is not the case for PUPD, and it provides favorable results during evaluations: in our measurements, it committed up to 2 times fewer inversions than the PUPD.

ACKNOWLEDGEMENTS



Co-funded by the
European Union

The author would like to express his sincere gratitude to Gábor Rétvári and Csaba Sarkadi for their support in this research. This study has received funding from the European Union's Horizon Europe research and innovation programme under the Marie Skłodowska-Curie grant agreement No. 101155116. Research partially supported by the National Research, Development and Innovation Fund of Hungary (grant No. 135606).

REFERENCES

- [1] B. Vass, C. Sarkadi, and G. Rétvári, "Programmable Packet Scheduling With SP-PIFO: Theory, Algorithms and Evaluation," in *IEEE INFOCOM Workshop on Networking Algorithms (WNA)*, London, United Kingdom, May 2022.
- [2] P. Bosshart, D. Daly, G. Gibb, M. Izzard, N. McKeown, J. Rexford, C. Schlesinger, D. Talayco, A. Vahdat, G. Varghese, and D. Walker, "P4: Programming protocol-independent packet processors," *ACM SIGCOMM Computer Communication Review*, vol. 44, no. 3, pp. 87–95, 2014.
- [3] O. Michel, R. Bifulco, G. Rétvári, and S. Schmid, "The programmable data plane: Abstractions, architectures, algorithms, and applications," *ACM Comput. Surv.*, vol. 54, no. 4, 2021.
- [4] A. Sivaraman, S. Subramanian, M. Alizadeh, S. Chole, S.-T. Chuang, A. Agrawal, H. Balakrishnan, T. Edsall, S. Katti, and N. McKeown, "Programmable packet scheduling at line rate," in *Proceedings of the 2016 ACM SIGCOMM Conference*, ser. SIGCOMM '16, New York, NY, USA: Association for Computing Machinery, 2016, p. 44–57. [Online]. Available: [doi: 10.1145/2934872.2934899](https://doi.org/10.1145/2934872.2934899)
- [5] A. G. Alcoz, A. Dietmüller, and L. Vanbever, "SP-PIFO: Approximating Push-In First-Out Behaviors using Strict-Priority Queues," in *17th USENIX Symposium on Networked Systems Design and Implementation (NSDI 20)*, 2020. [Online]. Available: <https://www.usenix.org/conference/nsdi20/presentation/alcoz>
- [6] R. Bellman, "On a routing problem," *Quarterly of applied mathematics*, vol. 16, no. 1, pp. 87–90, 1958.
- [7] A. Agarwal, B. Schieber, and T. Tokuyama, "Finding a minimum weight k-link path in graphs with Monge property and applications," in *Proc. ACM Symposium on Computational Geometry*, 1993, pp. 189–197.
- [8] B. Schieber, "Computing a minimum weight k-link path in graphs with the concave Monge property," *Journal of Algorithms*, vol. 29, no. 2, pp. 204–222, 1998.
- [9] B. Olstad and F. Manne, "Efficient partitioning of sequences," *IEEE Transactions on Computers*, vol. 44, no. 11, pp. 1322–1326, 1995.
- [10] S. Narayana, A. Sivaraman, V. Nathan, P. Goyal, V. Arun, M. Alizadeh, V. Jeyakumar, and C. Kim, "Language-directed hardware design for network performance monitoring," in *Proceedings of the Conference of the ACM Special Interest Group on Data Communication*, ser. SIGCOMM '17, New York, NY, USA: Association for Computing Machinery, 2017. [Online]. Available: [doi: 10.1145/3098822.3098829](https://doi.org/10.1145/3098822.3098829)
- [11] A. Fingerhut, P4, "Floating point operations in <https://github.com/jafingerhut/p4-guide/blob/d03b4d726a75192f8c7cb7e2ee0d4fceda3bacca/docs/floating-point-operations.md>, 2020.
- [12] Z. Yu, C. Hu, J. Wu, X. Sun, V. Braverman, M. Chowdhury, Z. Liu, and X. Jin, "Programmable packet scheduling with a single queue," ser. SIGCOMM '21, New York, NY, USA: Association for Computing Machinery, 2021, p. 179–193. [Online]. Available: [doi: 10.1145/3452296.3472887](https://doi.org/10.1145/3452296.3472887)
- [13] G. Memik, W. Mangione-Smith, and W. Hu, "NetBench: a benchmarking suite for network processors," in *IEEE/ACM International Conference on Computer Aided Design. ICCAD 2001. IEEE/ACM Digest of Technical Papers (Cat. No. 01CH37281)*, 2001, pp. 39–42.
- [14] A. Demers, S. Keshav, and S. Shenker, "Analysis and Simulation of a Fair Queueing Algorithm," in *ACM SIGCOMM*, New York, NY, USA, 1989.
- [15] P. E. McKenney, "Stochastic Fairness Queueing," in *IEEE INFOCOM*, 1990.
- [16] D. D. Clark, S. Shenker, and L. Zhang, "Supporting Real-time Applications in an Integrated Services Packet Network: Architecture and Mechanism," in *ACM SIGCOMM*, Baltimore, MD, USA, 1992.
- [17] M. Shreedhar and G. Varghese, "Efficient Fair Queueing Using Deficit Round Robin," in *ACM SIGCOMM*, Cambridge, Massachusetts, USA, 1995.
- [18] P. Goyal, H. M. Vin, and H. Chen, "Start-time Fair Queueing: A Scheduling Algorithm for Integrated Services Packet Switching Networks," in *ACM SIGCOMM*, Palo Alto, CA, USA, 1996.
- [19] M. Alizadeh, S. Yang, M. Sharif, S. Katti, N. McKeown, B. Prabhakar, and S. Shenker, "Pfabric: Minimal near-optimal datacenter transport," *SIGCOMM Comput. Commun. Rev.*, vol. 43, no. 4, p. 435–446, Aug. 2013. [Online]. Available: [doi: 10.1145/2534169.2486031](https://doi.org/10.1145/2534169.2486031)
- [20] L. E. Schrage and L. W. Miller, "The Queue M/G/1 with the Shortest Remaining Processing Time Discipline," 1966.
- [21] W. Bai, L. Chen, K. Chen, D. Han, C. Tian, and H. Wang, "Information-Agnostic Flow Scheduling for Commodity Data Centers," in *USENIX NSDI*, Oakland, CA, USA, 2015.
- [22] N. K. Sharma, M. Liu, K. Atreya, and A. Krishnamurthy, "Approximating Fair Queueing on Reconfigurable Switches," in *USENIX NSDI*, Renton, WA, USA, 2018.
- [23] R. Mittal, R. Agarwal, S. Ratnasamy, and S. Shenker, "Universal Packet Scheduling," in *USENIX NSDI*, Santa Clara, CA, USA, 2016.
- [24] A. Sivaraman, S. Subramanian, A. Agrawal, S. Chole, S.-T. Chuang, T. Edsall, M. Alizadeh, S. Katti, N. McKeown, and H. Balakrishnan, "Towards Programmable Packet Scheduling," in *ACM HotNets*, Philadelphia, PA, USA, 2015.
- [25] V. Shrivastav, "Fast, Scalable, and Programmable Packet Scheduler in Hardware," in *SIGCOMM '19*, Beijing, China, 2019.
- [26] R. Yao, Z. Zhang, G. Fang, P. Gao, S. Liu, Y. Fan, Y. Xu, and H. J. Chao, "BMW Tree: Large-scale, High-throughput and Modular PIFO Implementation using Balanced Multi-Way Sorting Tree," in *ACM SIGCOMM*, New York, NY, USA, 2023.
- [27] N. Atre, H. Sadok, and J. Sherry, "BBQ: A Fast and Scalable Integer Priority Queue for Hardware Packet Scheduling," in *USENIX NSDI*, Santa Clara, CA, USA, 2024.
- [28] P. Gao, A. Dalleggio, J. Liu, C. Peng, Y. Xu, and H. J. Chao, "Sifter: An Inversion-Free and Large-Capacity Programmable Packet Scheduler," in *USENIX NSDI*, Santa Clara, CA, USA, Apr. 2024.
- [29] T. Yang, J. Li, Y. Zhao, K. Yang, H. Wang, J. Jiang, Y. Zhang, and N. Zhang, "QCluster: Clustering Packets for Flow Scheduling," 2020.
- [30] N. K. Sharma, C. Zhao, M. Liu, P. G. Kannan, C. Kim, A. Krishnamurthy, and A. Sivaraman, "Programmable Calendar Queues for High-speed Packet Scheduling," in *USENIX NSDI*, Santa Clara, CA, USA, 2020.
- [31] P. Gao, A. Dalleggio, Y. Xu, and H. J. Chao, "Gearbox: A Hierarchical Packet Scheduler for Approximate Weighted Fair Queueing," in *USENIX NSDI*, Renton, WA, USA, 2022.
- [32] A. Gran Alcoz, B. Vass, P. Namyar, B. Arzani, G. Rétvári, and L. Vanbever, "Everything matters in programmable packet scheduling," in *22st USENIX Symposium on Networked Systems Design and Implementation (NSDI 2025)*, 2025.
- [33] G. Teschl, *Ordinary differential equations and dynamical systems*. American Mathematical Soc., 2012, vol. 140.
- [34] "Über die abgrenzung der eigenwerte einer matrix," *Bulletin de l'Académie des Sciences de l'URSS. Classe des sciences mathématiques*, no. 6, pp. 749–754, 1931.
- [35] P. A. F. Parsiad Azimzadeh. Weakly chained matrices, policy iteration, and impulse control. Accessed: 2023. [Online]. Available: [arXiv:1510.03928](https://arxiv.org/abs/1510.03928)
- [36] S. Noschese, L. Pasquini, and L. Reichel, "Tridiagonal Toeplitz matrices: properties and novel applications," *Numerical linear algebra with applications*, vol. 20, no. 2, pp. 302–326, 2013.



Balázs Vass received his MSc degree in applied mathematics at ELTE, Budapest in 2016. He finished his PhD in informatics in 2022 on the Budapest University of Technology and Economics (BME). His research interests include networking, survivability, combinatorial optimization, and graph theory. He was an invited speaker of COST RECODIS Training School on Design of Disaster-resilient Communication Networks '19. He has been a TPC member of IEEE INFOCOM since 2023.

APPENDIX

Proof of Lemma 1: For $i \in \{0, \dots, n\}$, let $r_i = q_{i+1} - F^{-1}(i/n)$. Then, $r_0(t) = r_n(t) = 0$, for every $t \geq 0$, thus they do not make any change from the stability perspective, and can be omitted from further investigation. For $i \in \{1, \dots, n-1\}$, we have:

$$\begin{aligned} r'_i &= (q_{i+1} - F^{-1}(i/n))' = q'_{i+1} - 0 = \\ &= F(q_{i+2}(t)) - 2 \cdot F(q_{i+1}(t)) + F(q_i(t)). \end{aligned}$$

For shorthand notation, we will use $F_i(t) := F(q_{i+1}(t))$. With this, we define the following matrix in the function of t :

$$A_F(t) = \begin{pmatrix} -2F_1(t) & F_2(t) & 0 & \dots & 0 \\ F_1(t) & -2F_2(t) & F_3(t) & 0 & \vdots \\ 0 & F_2(t) & -2F_3(t) & F_4(t) & 0 \\ \vdots & \ddots & \ddots & \ddots & \ddots \\ 0 & \dots & 0 & F_{n-2}(t) & -2F_{n-1}(t) \end{pmatrix} \quad (7)$$

Using $A_F(t)$, we can write:

$$\dot{r}(t) = A_F(t) \cdot r(t). \quad (8)$$

To prove that $r = 0$ is globally stable, we will use the following Proposition.

Proposition 1 (Corollary 3.6 of [33]): The linear system $\dot{r}(t) = A \cdot r(t)$, $r(0) = r_0$, $A \in \mathbb{R}^{m \times m}$ is asymptotically stable if and only if all eigenvalues α_j of A satisfy $\text{Re}(\alpha_j) < 0$. Moreover, in this case there is a constant $C = C(\alpha)$ for every $\alpha < \min\{-\text{Re}(\alpha_j)\}_{j=1}^m$ such that

$$\|e^{At}\| \leq Ce^{-t\alpha}, \quad t \geq 0. \quad (9)$$

By Prop. 1, for showing the global stability, it is enough to show that $\text{Re}(\lambda) < 0$ for all eigenvalues of A_F . Gershgorin's circle theorem [34] applied here tells that every eigenvalue of A_F lies in at least one of the disks $D(a_{ii}, \sum_{j \neq i} a_{ji})$. Since, in A_F , $a_{ii} \geq \sum_{j \neq i} a_{ji}$ (i.e., A_F is weakly diagonally dominant, WDD), for each eigenvalue λ we have that either $\text{Re}(\lambda) < 0$, or $\lambda = 0$. This latter case, however, cannot happen since A_F is a *weakly chained diagonally dominant* (WCDD) matrix (defined in a moment), and WCDD matrices are non-singular by [35, Lemma 3.2.]. Here, according to [35, Def. 3.1.], a matrix B is WCDD, if it is WDD, and for each column r , there is a directed path from r to a strictly diagonally dominant column p in the digraph of $D_B = (V_B, A_B)$, where $(i, j) \in A_B$ exactly if $B_{ij} \neq 0$. Clearly, A_F is WCDD.

The global stability of system (8) at $(q_1, \dots, q_{n+1}) = (0, F^{-1}(1/n), F^{-1}(2/n), \dots, F^{-1}(n/n))$ follows from the fact that any system yielded by fixing matrix $A_F(t)$ at time point t is globally (exponentially) stable at $(0, F^{-1}(1/n), F^{-1}(2/n), \dots, F^{-1}(n/n))$.

Finally, the exponential stability of (8) if F is uniform on $[0, M]$ is imminent from Prop. 1 combined with the observation that, in that case, $A_F(t)$ is just a constant tridiagonal matrix A with its elements on the main diagonal equalling -2 , while on the upper and lower diagonal, equalling 1 . By [36, Eq. (4)], we know that its eigenvalues are $2 \cdot (-1 + \cos \frac{k\pi}{n})$, $k \in \{1, \dots, n-1\}$. Further investigating the eigenvalues, we see that the largest among them

equals $2(\cos \pi/n - 1)$, since the cosine function is monotone decreasing on interval $[0, \pi]$. Thus, in Prop. 1, we can use e.g., $\alpha = 1 - \cos \pi/n$. With this α , combining Eq. (8) and Eq. (9), we have:

$$\begin{aligned} \|r(t)\| &= e^{At} \cdot r(0) \leq \|e^{At}\| \cdot \|r(0)\| \leq \\ &\leq Ce^{-t\alpha} \|r(0)\| = Ce^{t(\cos \pi/n - 1)} \|r(0)\|, \end{aligned} \quad (10)$$

for some appropriate constant C . This gives a hint of how fast the queue bounds are converging to the fix point. ■

Challenges of DNS in the Post-Quantum Era: Improving Security with Post-Quantum TLS

Abdullah Aydeger[†], Sanzida Hoque[†], and Engin Zeydan^{*}

Abstract—The Domain Name System (DNS), an important component of the Internet infrastructure, is vulnerable to various attacks that can jeopardize the security and privacy of Internet communications. While DNS over TLS (DoT) is widely used to improve DNS security, the advent of quantum computing poses a significant threat to the underlying cryptographic algorithms used in TLS. In this paper, we propose a comprehensive framework for DNS over Post-Quantum TLS (DoPQT) to address this challenge. Our framework integrates post-quantum cryptographic algorithms into DoT, ensuring robust security against both classical and quantum attacks. We introduce a hybrid key exchange mechanism and post-quantum authentication procedures to protect the confidentiality, integrity, and authenticity of DNS traffic. DoPQT has the potential to offer comparable performance to existing solutions while demonstrating superior quantum resistance. This research contributes to the development of a secure and resilient DNS infrastructure in the post-quantum era. It has been observed that the handshake process is most affected by increased DNS queries and is the main source of the bottleneck. On the other hand, the percentage loss in throughput when using the PQC algorithm (i.e., ML-KEM) is about 33-40% for different DNS queries.

Index Terms—DNS Security, TLS, PQC

I. INTRODUCTION

The Domain Name System (DNS) is a fundamental component of the Internet. It enables the conversion of human-readable domain names into machine-readable IP addresses. This central function makes the DNS a prime target for various cyber threats, including cache poisoning, man-in-the-middle attacks, and DNS hijacking [1], [2]. Such attacks can have serious consequences, such as data exfiltration, unauthorized censorship, and redirection to malicious websites [3]. To mitigate these risks, DNS over TLS (DoT) was introduced, which encrypts DNS requests and responses to prevent eavesdropping and manipulation attempts [4]. However, the emergence of quantum computing poses a significant risk to the security of existing cryptographic protocols, including those used in TLS [5]. The extraordinary computational capabilities of quantum computers could potentially crack the public-key cryptographic algorithms that currently secure DoT, making DNS traffic vulnerable to sophisticated quantum-enabled attacks. This potential vulnerability underscores the urgent need to transition to Post Quantum Cryptography (PQC), which is designed to withstand the challenges posed by both classical and quantum computing [6].

[†] Dept. of Electrical Engineering and Computer Science, Florida Institute of Technology, Melbourne, FL, USA.

^{*} Centre Tecnològic de Telecomunicacions de Catalunya (CTTC/CERCA), Barcelona, Spain.

(E-mail: aaydeger@fit.edu, shoque2023@my.fit.edu, engin.zeydan@cttc.cat)

Recent research has explored various post-quantum cryptographic schemes that could be integrated into DNS protocols to enhance security against quantum threats [7]. For instance, implementations of lattice-based cryptography and hash-based signatures have been suggested as viable options for securing DNS traffic in a post-quantum world [8]. Furthermore, the importance of transitioning DNS infrastructure to support quantum-resistant algorithms is gaining traction, with several proposals advocating for hybrid approaches that combine classical and post-quantum cryptographic techniques [9].

To address these concerns, we propose the integration of PQC into DoT, creating what we term DNS over Post-Quantum TLS (DoPQT). This approach is critical for ensuring the long-term security and resilience of DNS in the face of advancing quantum technologies. The main contributions of this paper are as follows:

- The paper introduces a novel integration of DNS with Post-Quantum TLS, aiming to protect DNS queries and responses against quantum computing attacks. This ensures that DNS communications remain secure in the post-quantum era.
- The paper provides a comprehensive performance evaluation of DNS over Post-Quantum TLS, demonstrating that the approach can be practically implemented with minimal impact on DNS query latency and overall system performance.
- A detailed security analysis is conducted, showing that the proposed method provides strong resistance to both classical and quantum attacks, highlighting the robustness of Post-Quantum TLS for DNS protection.
- By proposing this integration, the paper contributes to the future-proofing of DNS infrastructure against the growing threat of quantum computing, paving the way for broader adoption of post-quantum security measures in internet protocols.

The subsequent sections of the paper are structured in the following manner: Section II presents a concise overview of the existing literature on strategies for secure DNS solutions. Subsequently, the proposed approach is introduced in Section III. Section IV provides theoretical comparisons in terms of security and other metrics. Section V depicts the simulation setup and presents the findings of our evaluation. Section VI discusses the comparative analysis of PQC candidates and outlines future directions and efforts. Finally, Section VII concludes the paper via a thorough discussion and consideration of future research endeavors.

Challenges of DNS in the Post-Quantum Era: Improving Security with Post-Quantum TLS

II. RELATED WORK

Research into quantum-resistant solutions for DNS security has gained momentum. Prior efforts on post-quantum secure DNS primarily center on securing DNSSEC, the application-layer authentication mechanism used for validating DNS records. In contrast, relatively fewer studies focus on the encryption and confidentiality of DNS traffic, which is the core problem addressed by DNS-over-TLS (DoT).

Research by Muller et al. (2020) investigated the potential of lattice-based cryptography, a promising area of PQC, in providing quantum-resistant signatures for DNSSEC [10]. Pan et al. [11] and Raavi et al. [12] enhance DNSSEC using hybrid signatures and fragmentation mitigation, respectively, but do not address transport-layer encryption. Similarly, Goertzen et al. [13], McGowan et al. [14], and Schutijser et al. [15] advance PQ DNSSEC reliability through fragmentation-aware mechanisms and testbeds, yet leave transport-layer encryption untouched. SL-DNSSEC and TurboDNS [16], [17] propose MAC- and KEM-based protocols to reduce resolution cost, diverging from traditional DNSSEC validation and lacking transport integration. These lightweight designs optimize DNSSEC but are not applicable to TLS-based secure DNS. Jafarli, Beernink, and Goertzen [18]–[20] offer implementation insights for DNSSEC using FN-DSA, ML-DSA, and XMSS. Their evaluations highlight signature size trade-offs, but none address encryption-layer issues or compatibility with DoT. TITAN-DNS over HTTPS (DoH), proposed by Ali and Chen [21], integrates trust mechanisms and FrodoKEM-based TLS handshakes to secure DoH. Their adaptive design combines Bayesian inference and verifiable delay functions to detect malicious queries. Their scope is limited to DoH and omits DoT and DNSSEC alignment. Additionally, recent efforts have focused on evaluating the feasibility of post-quantum cryptographic (PQC) algorithms across diverse platforms, particularly constrained environments. Kannwischer et al. [22] benchmarked NIST PQC candidates on the Arm Cortex-M4 using the pqm4 framework, highlighting performance bottlenecks due to large key sizes, dynamic memory usage, and unoptimized code. Fitzgibbon et al. [23] evaluated ML-KEM (aka Kyber), ML-DSA (aka CRYSTALS-Dilithium), and FN-DSA (aka Falcon) on Raspberry Pi 4 devices, identifying ML-KEM and ML-DSA as the most efficient in key encapsulation and signatures, with FN-DSA providing the smallest handshake size. Abbasi et al. [24] conducted extensive cross-platform benchmarks, showing that PQC adoption in high-performance systems incurs minimal overhead, while constrained devices can face up to 12× slowdown depending on algorithm choice. Lonc et al. [25] assessed PQC feasibility in V2X systems, while Kumar [26] surveyed global PQC efforts, emphasizing implementation challenges in real-world settings. Despite these valuable contributions, performance evaluations of PQC integration within DNS security protocols remain relatively limited. While prior studies have explored PQC in TLS and constrained devices more broadly, there is still a lack of detailed analysis targeting DNS-specific use cases, particularly DNS-over-TLS (DoT). Our work contributes to this emerging area by examining the practical implications of post-quantum

transitions in DoT, considering protocol behavior, resource constraints, and algorithm suitability.

Recent advancements have also seen the integration of hybrid cryptographic algorithms, which combine both classical and post-quantum techniques, to gradually transition DNSSEC to a quantum-resistant framework while maintaining compatibility with existing systems [19], [27], [28]. The integration of hybrid cryptography into DNSSEC can, therefore, be seen as a strategic approach to future-proofing DNS against the anticipated quantum threats [27]. The authors in [28] have shown how hybrid key exchange mechanisms can be integrated into TLS 1.3 and SSHv2 so that these protocols can use both classical and post-quantum keys simultaneously. Their work highlights the challenges of negotiating multiple algorithms and effectively combining keys without significantly compromising performance. In these hybrid approaches, conventional cryptographic algorithms (such as RSA or ECDSA) and post-quantum algorithms (such as lattice or hash-based signatures) are used simultaneously. This dual-use strategy ensures that even if one set of algorithms is compromised by quantum computing, the other provides another layer of security, maintaining the integrity of DNSSEC. Studies by Microsoft and others have demonstrated that hybrid TLS implementations, using algorithms like FrodoKEM and classical ECDH, can achieve this balance, making them a viable option for upgrading DNSSEC.

Intensive work has also been done on the development of quantum-resistant key exchange mechanisms specifically for the TLS handshake process [28], [29]. Paquin, Stebila, and Tamvada (2020) compared different cryptographic post-quantum primitives in the context of TLS and highlighted the trade-offs associated with their implementation, especially in environments with high packet loss rates [28]. In addition, significant progress has been made in the development of quantum-resistant key exchange mechanisms for the TLS handshake process, together with industry partners. Microsoft's collaboration with the Open Quantum Safe project [30] has led to the development of a PQC fork of OpenSSL that integrates quantum-resistant key exchange and signature algorithms such as FrodoKEM and qTESLA. These efforts are critical to protecting TLS from potential quantum attacks while enabling hybrid modes of operation that provide security against both current and future threats. Similarly, Meta has addressed the challenges of using post-quantum key exchange mechanisms in large environments, addressing issues such as the larger key sizes that complicate TLS session resumption [31]. Despite these advances, seamlessly integrating PQC into the DoT framework while maintaining backward compatibility and minimizing performance degradation remains a major challenge. The complexity of integrating PQC into existing protocols without significant performance degradation, as highlighted in the Microsoft and Meta studies, underscores the need for further research to optimize these algorithms for practical use in DNS and other Internet security protocols.

Overall, a systematic and deployable integration of PQC into DoT, which secures DNS queries in transit is missing. While prior work makes significant progress in DNSSEC and DoH contexts, a unified and post-quantum-resilient design for

DNS-over-TLS remains underexplored. Our proposed framework, DoPQT, fills this gap by conceptually integrating post-quantum key exchange (ML-KEM) and authentication (FNDSA) into DoT, along with a simulation-based performance analysis to assess real-world feasibility.

III. PROPOSED FRAMEWORK

We propose a comprehensive framework for DNS over Post-Quantum TLS (DoPQT) that addresses the multiple challenges associated with securing DNS in a post-quantum world. Fig. 1 shows a post-quantum secured DNS system. Our framework is designed to strike a balance between improved security, performance efficiency, and compatibility with existing DNS infrastructure. The key components of our proposed framework include:

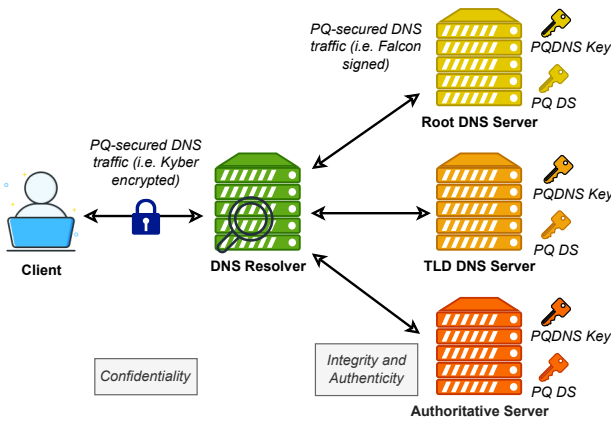


Fig. 1. PQ secured DNS system

- **Hybrid Key Exchange Mechanism:** Our framework uses a hybrid approach to key exchange that combines traditional cryptographic algorithms with post-quantum algorithms. This ensures that communication with both legacy systems and quantum-resistant clients remains secure and provides a smooth transition path when PQC technologies are introduced.
- **Post-Quantum Authentication:** To strengthen the authentication and integrity of DNS responses, we integrate post-quantum digital signature methods. These techniques provide robust protection against spoofing and data tampering and address vulnerabilities that could be exploited by attackers with quantum capabilities.
- **Lightweight Post-Quantum Cipher Suites:** Considering the potential computational overhead associated with PQC, our framework emphasizes the use of lightweight cipher suites. These suites are carefully selected to strike a balance between high security and acceptable performance and to minimize the impact on DNS query and response times.
- **Backward Compatibility:** A key aspect of our framework is its ability to integrate seamlessly with existing DoT implementations. This backward compatibility ensures that the transition to DoPQT can be gradual, allowing

current systems to remain operational while adopting quantum-resistant features as they become available.

A. Steps for DoPQT

The packet exchange with DNS via Post-Quantum TLS (DoPQT), depicted in Fig. 2, proceeds as follows. Note that these steps are conceptually similar to DoT, but have been extended by the integration of post-quantum cryptographic algorithms:

- Step 1: Client-Server Handshake (TLS 1.3)**
- (i) **ClientHello:** The client initiates the connection by sending a ClientHello message containing: a. Supported TLS versions (likely TLS 1.3 for efficiency), b. Supported cipher suites (including post-quantum key encapsulation mechanisms (KEMs) and digital signature algorithms), c. Client random nonce, supported extensions (e.g., Server Name Indication (SNI) to specify the desired domain).
 - (ii) **ServerHello:** The server responds with a ServerHello message: a. Selected TLS version, b. Selected cipher suite (including the chosen post-quantum KEM and signature algorithm), c. Server random nonce, d. Server's certificate (signed with a traditional or post-quantum digital signature).
 - (iii) **EncryptedExtensions:** (Optional) The server may send additional extensions if needed.
 - (iv) **CertificateRequest:** (Optional) If client authentication is required, the server requests the client's certificate.
 - (v) **Certificate:** (Optional) The client sends its certificate if requested.
 - (vi) **CertificateVerify:** (Optional) The client provides a digital signature to prove ownership of its certificate.
 - (vii) **Finished:** Both client and server send Finished messages, verifying the handshake and establishing encrypted communication using the shared secret derived from the post-quantum KEM.

- Step 2: DNS Query and Response (over Encrypted TLS Channel)**
- (i) **DNS Query:** The client sends a DNS query message (e.g., an A or AAAA query to resolve a domain name to an IP address) over the encrypted TLS channel.
 - (ii) **DNS Response:** The server processes the query and sends back a DNS response message containing the requested information (e.g., the IP address associated with the domain name).

Note that the post-quantum Key Encapsulation Mechanism (KEM) and digital signature algorithms used may depend on the chosen cipher suite and implementation. The exact format of the DNS query and response messages must comply with the standard DNS protocol specifications. Finally, the TLS handshake ensures the confidentiality, integrity, and authenticity of DNS traffic, protecting it from eavesdropping and tampering. The use of PQC provides an additional layer of security against possible quantum attacks in the future.

B. Example: Hybrid Key Exchange

In a hybrid key exchange scenario, the ClientHello could contain both traditional (e.g., X25519) and post-quantum (e.g., ML-KEM) KEMs. The server would then select one and include the corresponding public key in its ServerHello. The

Challenges of DNS in the Post-Quantum Era: Improving Security with Post-Quantum TLS

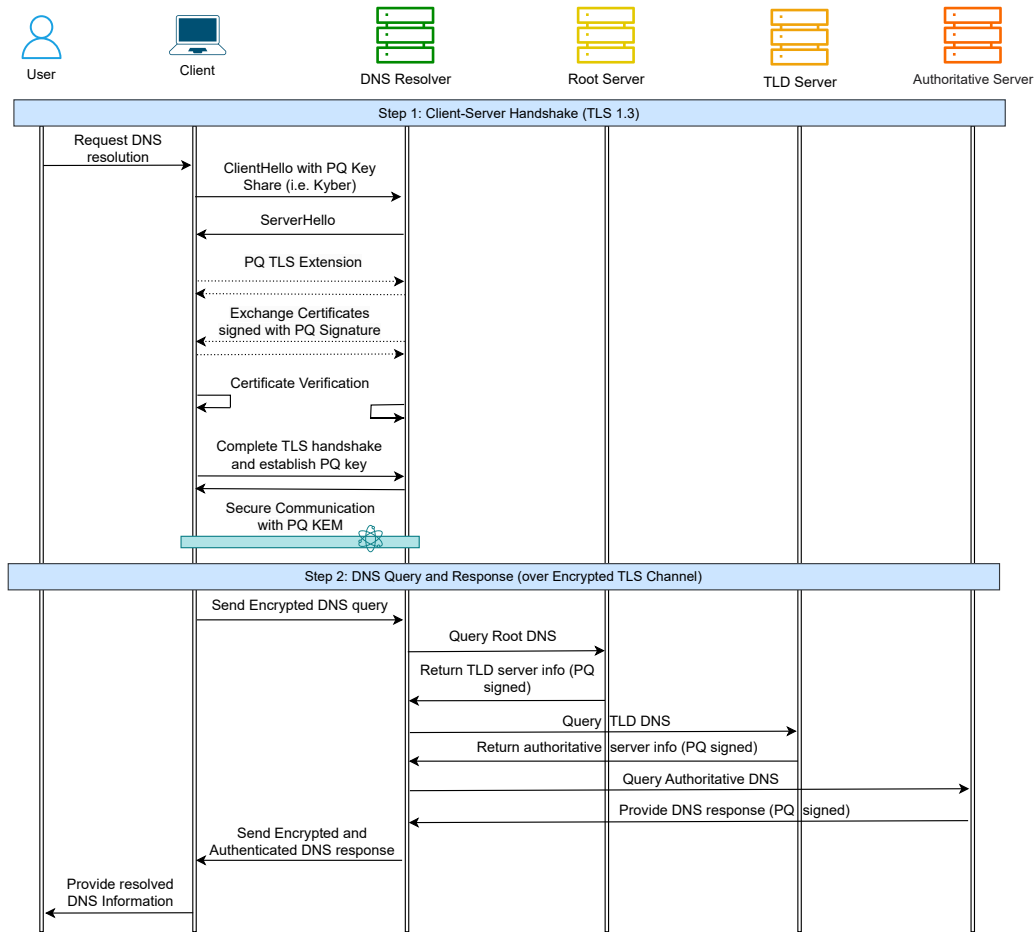


Fig. 2. Packet Exchange of DoPQT

client and server would then use their chosen KEMs to derive a shared secret for encryption.

IV. DoPQT: THEORETICAL AND QUANTIFIABLE COMPARISONS

In the context of DoPQT, performance metrics are critical for assessing whether the integration of post-quantum cryptographic algorithms maintains, improves or reduces the efficiency and practicality of DNS over TLS (DoT). Key performance metrics for DoPQT compared to DoT with potential impacts include:

- **Latency:** The time it takes to complete a DNS query and response cycle. Post-quantum algorithms usually require more computing resources, which can increase latency. This can result in slightly slower DNS query resolution times compared to DoT. However, if optimized post-quantum algorithms are used, the increase in latency can be minimized, although DoPQT may still have a higher latency than conventional DoT due to the added complexity of quantum-resistant cryptography.
- **Computational Overhead:** The amount of computing power required to perform cryptographic operations during the TLS handshake. PQC typically involves higher computational

overhead because it requires more complex mathematical operations, larger key sizes, and more intensive signature verification processes. This increased overhead could lead to slower performance during the initial handshake phase, resulting in higher CPU utilization and potentially longer connection setup times compared to DoT.

- **Bandwidth Consumption:** The amount of data transmitted during the cryptographic handshake and subsequent DNS query and response exchanges. Post-quantum algorithms often use larger key sizes and signatures, which can increase the overall bandwidth required during the handshake phase. This might lead to slightly higher data transfer requirements compared to DoT, particularly during the initial connection setup.
- **Memory Usage:** The amount of memory required to store cryptographic keys, certificates, and other relevant data structures. The larger key sizes and signatures associated with PQC may demand more memory resources. This increase in memory usage could affect systems with limited resources, making DoPQT more resource-intensive than DoT.
- **Security Strength:** The resilience of the cryptographic protocol against attacks, especially those that could be launched

by quantum computers. While DoPQT might introduce performance trade-offs in terms of latency, computational overhead, and bandwidth, it significantly enhances security. The use of post-quantum cryptographic algorithms provides robust protection against quantum computing threats, making DoPQT far more secure in the long term compared to DoT.

- **Scalability:** The ability of the protocol to handle a large number of concurrent DNS queries and responses. The increased computational and memory requirements of DoPQT may affect its scalability, especially in high-traffic environments. Systems may require optimization or more powerful hardware to maintain the same level of scalability as DoT.

A. DoPQT vs. DoT

Table I provides comparisons of DNS, DoT, and DoPQT in terms of various features. DoPQT will have a slightly higher latency and a slightly higher computational overhead due to the higher complexity of post-quantum operations. DoPQT may require more bandwidth and memory due to larger key sizes and cryptographic data, especially during the TLS handshake. DoPQT offers significantly more security compared to DoT, especially when it comes to defending against potential attacks on quantum computers, which outweigh the performance losses. The scalability of DoPQT could pose a problem due to the additional computing and resource requirements, but this can be mitigated by optimized algorithms and a more powerful infrastructure.

TABLE I
COMPARISON OF DNS, DoT, AND DoPQT

Feature	DNS	DoT	DoPQT
Encryption	No	Yes	Yes
Authentication	No	Yes	Yes
Integrity	No	Yes	Yes
Confidentiality	No	Yes	Yes
Quantum Resistance	No	No	Yes
Performance	Fastest	Slower than DNS	Slower than DoT
Deployment	Widely deployed	Increasing adoption	Future deployment
Complexity	Low	Moderate	High

B. Security Comparisons

The conventional DNS protocol works without encryption and is therefore vulnerable to eavesdropping and manipulation. DoT provides confidentiality, integrity, and authentication by encrypting DNS traffic with TLS. However, it is vulnerable to attacks from quantum computers. DoPQT offers the same security benefits as DoT, but with the added benefit of being resistant to quantum attacks. This is achieved by incorporating post-quantum cryptographic algorithms. The disadvantage is increased complexity and potentially lower performance compared to DoT. On the other hand, while DoPQT is the most secure option, its widespread adoption will depend on the development and standardization of PQC. The choice

between DNS, DoT, and DoPQT depends on specific security requirements and available resources. Organizations should consider transitioning to DoT as an interim solution while preparing for the later adoption of DoPQT.

C. Discussion on security performance trade-offs

1) **Security:** DoPQT can effectively prevent passive eavesdropping on DNS traffic due to the strong encryption provided by the post-quantum TLS channel. The use of post-quantum authentication mechanisms prevents an attacker from impersonating the DNS server or client, thwarting man-in-the-middle attacks. DoPQT's integrity protection ensures that DNS responses cannot be tampered with, preventing cache poisoning attacks.

2) **Key performance metrics:** The computational overhead and latency caused by the post-quantum cryptographic algorithms in DoPQT can be higher. The handshake time is the time required to establish a secure TLS connection. DoPQT may have a slightly longer handshake time compared to DoT due to the additional computations involved in hybrid key exchange and post-quantum authentication. Query/Response Latency is the time it takes for a DNS query to be sent, processed, and answered. DoPQT can have a slightly higher latency compared to DoT, which is primarily due to the encryption and decryption of messages using post-quantum cipher suites.

3) **Comparison with Existing Solutions:** The performance of DoPQT with traditional DoT and other relevant solutions, such as DNS over HTTPS (DoH) are compared in Table I. DoPQT offers comparable performance to DoH, while providing superior quantum resistance. While DoT remains the fastest option, its lack of quantum resistance makes it a less desirable choice for long-term security.

Overall, DoPQT can provide a viable solution for securing DNS in the post-quantum era. While there is a performance trade-off compared to traditional DoT, the improved security and resistance to quantum attacks outweigh the slight increase in latency. It is expected that further optimizations and advancements of PQC will reduce the performance overhead in the future, making DoPQT an even more attractive solution for securing the DNS infrastructure.

V. SIMULATION RESULTS

A. Simulation Parameters, Setup and Metrics

We used a discrete-event simulator in Python for event-driven updates to model the DNS resolution workflow, including cryptographic operations, network latency, and system load. The simulation parameters were calibrated using realistic latency distributions and published cryptographic benchmarks. While direct network and system measurements were outside the scope of this study, the log-normal distribution for network latency was chosen to reflect skewed delay distributions seen in Internet measurements. The log-normal distribution is often used to model network latencies because it naturally deals

Challenges of DNS in the Post-Quantum Era: Improving Security with Post-Quantum TLS

with the skewness observed in real latency data, where most latencies are small but there are occasionally higher values [32], [33]. The gamma distribution is often used to model the times required for operations or processes that are additive in nature, such as the time required for multiple steps in the TLS handshake or PQC operations. For instance, the authors in [34] discuss how end-to-end delay distribution (which is additive), when individual router delays are assumed to be exponential, leads to a gamma distribution. The exponential distribution is useful for modeling the time between events, e.g., the response time of a server, where the likelihood of longer times decreases exponentially. Table II contains the assumed parameters for the DNS query scenario for the above distributions. For ML-KEM-512, we use the key generation of 19.8 microseconds, 22.4 microseconds for key encapsulation, and 14.8 microseconds for key decapsulation¹. Server load is modeled using a simple, linear model that assumes processing time increases gradually as more queries are handled, reflecting a straightforward degradation in performance due to resource contention. Although simplified, this approach captures the essential trend of load-dependent processing delays observed in large-scale systems [35]. DNS resolution processes are simulated by breaking down the total latency into several components:

- **Network Latency:** Time taken for a DNS query to reach the root server. It is modelled with a log-normal distribution. This choice is selected based on empirical findings from latency measurements in wide area networks, where latency distributions are known to be positively skewed and have been successfully modelled with lognormal distributions in the past [32], [33].
- **TLS Handshake Latency:** Time required to establish a secure connection using TLS. This includes cryptographic operations such as key exchange. When PQC is enabled, this includes post-quantum primitives like ML-KEM-512. TLS handshake latency without PQC is modeled using a Gamma distribution to reflect multi-step operations in classical TLS. Parameters were selected to approximate handshake latencies observed in production TLS deployments. The Gamma distribution was chosen due to its suitability for modeling the total duration of sequential, variable-length processes, which aligns with the structure of TLS handshakes [34]. TLS handshake latency with PQC is the same distribution family (Gamma), but with a larger scale to account for the added cost of post-quantum cryptographic primitives.
- **PQC Operations Latency:** Additional time introduced by PQC operations (key generation, encapsulation, decapsulation). This is a sub-component of the handshake when PQC is used.
- **Response Latency:** Time taken for the DNS server to process the query and send a response. This is modeled as exponential, representing time to process and generate a DNS response on the server side. This reflects the memoryless nature of lightweight tasks like DNS query parsing and cache lookups.

¹Online: <https://medium.com/asecuritysite-when-bob-met-alice/bursting-the-myth-on-post-quantum-cryptography-pqc-being-slow-8d98d30b9a69>, Available: September 2024.

- **Overall DNS Query Latency:** The sum of all above components (network + handshake [+ PQC] + response) for a complete resolution round.

Throughout the paper, we refer to individual latency components using the terms above and reserve overall latency to mean the total DNS resolution time observed by the client.

TABLE II
ASSUMED PARAMETERS FOR DNS QUERY SCENARIO

Component	Distribution	Parameters
Network Latency	Log-Normal	Mean: 3, Sigma: 0.5
TLS Handshake Without PQC	Gamma	Shape: 2, Scale: 5
TLS Handshake With PQC	Gamma	Shape: 3, Scale: 20

The simulation was performed with multiple DNS queries and the average latencies for each of these components were calculated, both with and without PQC. We simulated DNS queries for different batch sizes (50 to 1000 queries) and ran each simulation 100 times to obtain average values. DNS resolution for a given number of queries with and without using PQC uses a priority queue (*heapq*) to manage events, ensuring they are processed in the correct order based on the event time. It involves the following steps: (i) DNS Query that simulates the network latency. (ii) TLS Handshake that simulates the time required to establish a secure connection. (iii) DNS Response that simulates the time taken by the server to process and respond to the query. Additionally, with PQC, we have (iv) PQC Operations that simulate ML-KEM-512 key generation, encapsulation, and decapsulation.

B. Numerical Results

1) **Latency:** Fig. 3 shows a detailed breakdown of how different latency components contribute to the total time required to respond to DNS queries, both with and without the use of PQC (ML-KEM-512). The impact of PQC operations on DNS resolution performance is visualized by running the simulation for different query counts and averaging the results. As the number of DNS queries increases, handshake latency becomes the dominant component, far exceeding network, PQC and response latency. The handshake latency increases by approximately 250% for 400 queries and by 500% for 800 queries. This indicates that the process of establishing secure connections significantly impacts performance under high load. On the other hand, both network and PQC latencies remain relatively stable under varying workloads, which means that they are not significantly affected by the number of DNS queries. The network latency is generally stable due to the constant transmission overhead, while the PQC latency is low and efficient in this context. The response latency increases steadily and is about 71.4% to 114.3% at 400 and 800 queries, respectively. This indicates that the server's ability to process and respond to requests decreases with the load, although not as much as with the handshake process.

2) **Throughput:** Throughput as the number of DNS requests processed per second. This can be done by measuring the

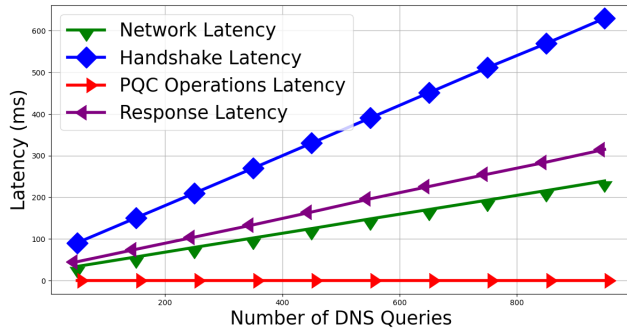


Fig. 3. Latency breakdown compared to the number of DNS queries (averaged over 100 runs) with PQC.

total time taken to process a batch of DNS requests and then dividing the number of requests by this total time. In simulations, throughput is calculated for each stage by dividing the number of DNS queries by the total time spent in that stage (converted from milliseconds to seconds). Total throughput is the total throughput taking into account the entire DNS resolution process. Network throughput is the throughput considering only the network latency. Handshake throughput is the throughput taking into account the TLS handshake latency. PQC throughput is the throughput taking into account the PQC operations (ML-KEM-512). Response throughput is the throughput taking into account the response time of the server.

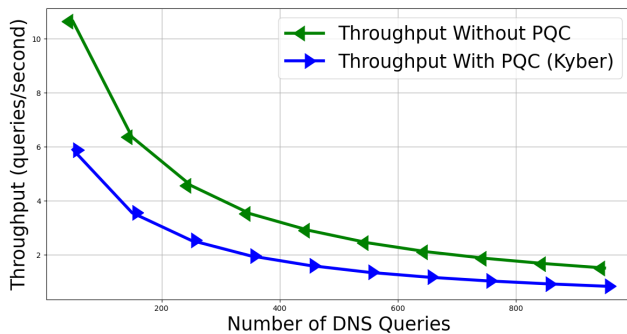


Fig. 4. Throughput vs. number of DNS queries (average over 100 runs) with and without PQC

Fig. 4 shows the impact of PQC on the throughput of DNS query processing. The overall throughput decreases significantly when PQC is enabled, especially at higher load (i.e. with a large number of DNS queries). It stabilizes around 1.5-2 queries/second for higher query counts with PQC. There is performance cost of PQC. As the number of DNS queries increases, the computational load of PQC becomes more significant, resulting in lower throughput compared to the scenario without PQC. Throughput decreases by about 40% and 33% when PQC is used for 100 and 800 DNS queries respectively. The results in Fig. 4 also show that while PQC (especially ML-KEM, also known as Kyber) ensures post-quantum security, it incurs additional overhead that can affect the server's ability to process DNS queries efficiently, especially when the load increases.

Fig. 5 shows all throughput (queries per second) across

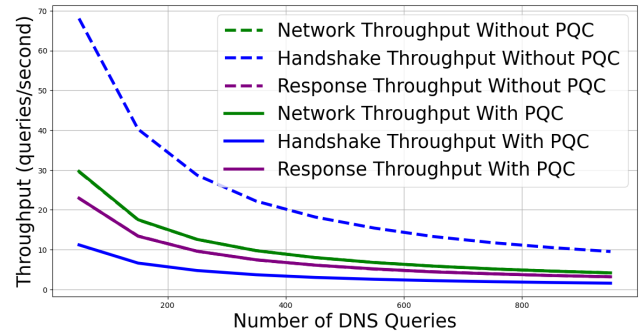


Fig. 5. Throughput vs. number of DNS queries for each phase of the workflow (averaged over 100 runs) with and without PQC.

different stages of DNS resolution (network, handshake, and response) both with and without PQC. The most important effect of PQC is the handshake throughput. The handshake process is much slower with PQC, with a decrease of 80-82% due to the computational overhead introduced by PQC. This is due to the additional cryptographic operations (e.g., ML-KEM-512 operations such as key generation, encapsulation, and decapsulation). The decrease in throughput with increasing number of queries is more pronounced in the handshake phase, especially when PQC is enabled. It decreases by 80-82% due to the computational overhead caused by PQC. The network throughput is largely unaffected by PQC, as PQC mainly influences the cryptographic processes and not the network transmission itself. With PQC, the throughput only decreases by 28.6-33.3% as the number of queries increases. Response throughput is affected to some extent, but remains stable as the server processes an increasing number of queries. It decreases by 25-33.3%, which shows that PQC moderately affects the server's ability to respond to DNS queries.

VI. DISCUSSION

A. Comparative View of PQC Candidates

While our performance evaluation focused on ML-KEM (KEM) due to the strong NIST standing and compatibility with TLS handshake constraints [24], other PQC candidates offer different trade-offs. Table III [36]–[38] compares major classes of PQC primitives, emphasizing their relevance to DNS security protocols, particularly DNS-over-TLS (DoT). Among the NIST-standardized algorithms, lattice-based schemes (ML-KEM for KEM and ML-DSA/FN-DSA for DS) are best suited for DNS environments due to their balance of small key/signature sizes, high performance, and stateless operation. Notably, FN-DSA offers the smallest signature sizes (e.g., 666 bytes for FN-DSA-512), but its sensitivity to side-channel attacks and reliance on constant-time floating-point arithmetic complicate secure deployment. SPHINCS+, a hash-based and conservative digital signature algorithm, provides strong security guarantees without relying on algebraic structures, but its large signatures [23] (8–17KB) and statefulness introduce challenges in bandwidth-limited or high-throughput DNS scenarios. Classic McEliece remains the only code-based KEM accepted in NIST's standardization process. It is known for its mature security foundations and fast encryption/decryption.

Challenges of DNS in the Post-Quantum Era: Improving Security with Post-Quantum TLS

TABLE III
COMPARISON OF POST-QUANTUM CRYPTOGRAPHIC PRIMITIVES FOR DNS SUITABILITY

Parameter	Isogeny-based (SIKE, SQISign)	Lattice-based (ML-KEM, ML-DSA)	Code-based (McEliece)	Multivariate-based (Rainbow)	Hash-based (SPHINCS)
Platform Suitability	IoT-friendly (but slow)	Suitable for general and embedded platforms	Not ideal for constrained devices	Not practical	Good for high-assurance systems
Computational Efficiency	Slow	Fast	Moderate	Fast	Fast
Hardware Acceleration	Limited availability	Strong (widely implemented)	Moderate	Limited	Moderate
Public Key Size	Small (around hundreds of bytes)	Moderate (a few KB)	Large (tens to hundreds of KB)	Large (tens to hundreds of KB)	Small (a few bytes to KB)
Private Key Size	Small	Moderate	Large	Large	Small
Signature Size	Small (around hundreds of bytes)	Moderate to large (KB)	Large	Moderate to large (KB)	Large (tens to hundreds of KB)
NIST Status	SIKE: Rejected (Broken), others: experimental	Standardized (ML-KEM, ML-DSA)	Alternate finalist	Rainbow: Rejected (Broken)	Standardized (SPHINCS+)
Side-Channel Resistance	Poor (needs careful masking)	Moderate to good (with constant-time impl.)	Good (structure-dependent)	Weak	Good
Stateful / Stateless	Stateless	Stateless	Stateless	Stateless	Stateful
DNS Suitability	Low	High	Low	Low	Medium

However, its massive public key sizes (over 1MB) make it entirely unsuitable for DNS or TLS scenarios, especially those relying on UDP transmission. Due to this key size constraint [23], McEliece is rarely considered in PQC DNS discussions and is not a viable candidate for DNS-over-TLS deployments. Isogeny-based schemes such as SIKE previously attracted attention due to their remarkably small key sizes, which seemed ideal for constrained protocols like DNS, especially in IoT environment [39]. However, SIKE was broken in 2022 by Castryck and Decru [40], removing it from NIST consideration. SQISign [41], a newer isogeny-based signature scheme, remains a research candidate but is not yet mature enough for deployment. While isogeny cryptography holds theoretical appeal for bandwidth-constrained environments, no current scheme from this class is suitable or secure enough for PQC DNS integration. Multivariate schemes like Rainbow have been cryptanalytically broken and excluded from NIST's standardization.

Among all classes, lattice-based cryptography, specifically ML-KEM for key encapsulation and ML-DSA or FN-DSA for digital signatures, offers the best balance of security, performance, and compatibility for post-quantum DNS security, especially DNS-over-TLS. Our work reflects this by choosing ML-KEM for integration and evaluation within the DoPQT framework, supported by simulation-based measurements and compatibility analysis.

B. Future Directions and Efforts

The need to transition to PQC to secure the Internet infrastructure goes beyond DoT and also includes other important DNS protocols. DoH and DNS over QUIC (DoQ), both of which are growing in popularity due to their enhanced privacy and performance benefits, must also be adapted to withstand potential quantum attacks. We provide the implementation strategies for the DoPQT in Table IV.

- *Post-Quantum DoH (DoH-PQC)*: The integration of PQC in DoH is crucial for the protection of DNS requests transmitted over HTTPS. This requires careful selection of post-quantum algorithms for the TLS layers, striking a balance between security, performance, and backward compatibility. Research efforts must focus on optimizing the implementation to minimize the latency and computational overhead incurred by PQC within DoH. The project also aims to investigate how PQC can further improve privacy protection in DoH to ensure that quantum attackers cannot compromise user data. Finally, prototyping and testing under real-world conditions are necessary to evaluate the feasibility and performance of DoH PQC under different network conditions.
- *Post-Quantum DoQ (DoQ-PQC)*: The development of a quantum-resistant DoQ requires the identification and implementation of post-quantum key exchange mechanisms that are tuned to the fast handshake process of QUIC. The challenge is to preserve the low-latency advantages of DoQ even when integrating PQC. In addition, since QUIC supports multiplexed connections, efficient integration of PQC into this model is critical to avoid performance bottlenecks. Researchers must ensure that DoQ-PQC remains interoperable with existing QUIC implementations and is resistant to both quantum and conventional network attacks.
- *Unified Framework for PQC in DNS Protocols*: A standardized framework for the seamless integration of PQC across different DNS transport protocols (DoT, DoH, DoQ) is an important long-term goal. This includes standardizing PQC implementations and ensuring a consistent transition to post-quantum security across DNS-over-encryption protocols. Collaboration with standardization bodies such as the IETF will be critical to promote broad adoption and interoperability. The framework should also consider phased transition strategies that take into account the evolving net-

TABLE IV
IMPLEMENTATION STAGES FOR DoPQT

Stage	Objective	Key Activities	Expected Outcome	Reference Performance Metric [42]
Research & Development	Identify PQC algorithms suitable for DoPQT	Evaluate NIST PQC finalists (e.g., ML-KEM, FN-DSA); prototype DoPQT resolver	Ranked list by key metrics; Identified PQC candidates	$Latency \leq 36$ ms median (Baseline: Google DoT)
Pilot Testing	Test DoPQT in controlled environments	Deploy in test networks, evaluate DoPQT performance	Performance data, feedback	$Query\ success\ rate \geq 99\%$ within 150 ms (Baseline: Cloudflare DoT)
Optimization	Enhance efficiency and compatibility with existing DNS infrastructure	Algorithm tuning, hardware acceleration strategies	Optimized DoPQT implementation	$Throughput \geq 500$ queries/sec (Baseline: DoT throughput trends)
Deployment	Roll out DoPQT in production networks	Phased deployment, monitoring	Secure DNS with minimal disruption	$Failure\ rate \leq 0.1\%$ (Baseline: DoT failure rates)

work infrastructure and ensure uninterrupted DNS services during the transition.

- *Prototype Validation with OQS and OpenSSL*: While our evaluation provides insight into the performance implications of post-quantum cryptography in DNS resolution workflows, we acknowledge that our current evaluation is limited to simulation-based analysis. Although the parameters were chosen based on realistic distributions and empirical cryptographic benchmarks, actual deployment in live networks may introduce additional overheads or optimizations not captured here. To further strengthen the credibility and practical relevance of our work, we plan to implement the DoPQT framework using DoPQT framework using DNS server software such as BIND, Unbound, or Knot DNS, built with libraries such as Open Quantum Safe (OQS) and patched versions of OpenSSL PQC support or alternative TLS libraries such as BoringSSL and wolfSSL [43]. This will allow us to assess real-world performance, interoperability with existing TLS implementations, and protocol behavior in dynamic network environments. Prototype-based testing will also provide insights into handshake latency, cryptographic overhead, and compatibility across different resolver configurations. Validating our findings on an experimental testbed with real DNS server implementations and hybrid TLS configurations will be an important direction for future work, and we plan to explore this as PQC support matures in production systems.
- *Real-World Integration Challenges*: While the proposed DoPQT framework has been evaluated in a simulated environment, its integration into real-world DNS infrastructure introduces additional challenges. The hierarchical structure of DNS, including recursive resolvers, forwarders, and caching mechanisms, can impact performance, especially when incorporating post-quantum TLS handshakes with larger key exchanges and signature payloads. Moreover, CDN infrastructures often optimize DNS queries using proprietary techniques. This ensures that the privacy, efficiency, and reliability that users expect from a modern Internet infrastructure are maintained even in the face of threats from quantum computing. Cross-layer optimizations across these layers are essential for practical deployment and will be explored in future work.

The simulation-based evaluation provided indicative trends to understand the impact of PQC on DNS resolution perfor-

mance. Future work will also include measurements from a deployed PQC-enabled DNS prototype to validate and refine these results.

VII. CONCLUSION

The impending threat of quantum computing necessitates immediate action to fortify DNS security through post-quantum cryptographic solutions. Our proposed DNS over Post-Quantum TLS (DoPQT) framework represents a proactive and comprehensive approach to addressing these emerging challenges. By combining traditional and post-quantum cryptographic algorithms, we offer a robust defense against both classical and quantum-enabled attacks, ensuring the continued integrity and security of DNS. While DoPQT introduces some performance trade-offs, particularly in terms of latency, computational overhead, bandwidth, and memory usage, these are balanced by the substantial increase in security against quantum threats. Optimizations and future advancements in PQC are expected to reduce these impacts, potentially bringing DoPQT performance closer to that of traditional DoT over time. The hybrid key exchange mechanism, emphasis on lightweight cipher suites, and focus on backward compatibility further enhance the practicality and effectiveness of our solution. As PQC continues to evolve, ongoing research and development will be essential to optimizing performance, identifying and mitigating vulnerabilities, and promoting widespread adoption. The successful implementation of DoPQT will play a pivotal role in securing the DNS infrastructure and safeguarding the future of internet communications in a post-quantum world.

REFERENCES

- [1] M. Taleby Ahvanooy, W. Mazurezyk, J. Zhao, *et al.*, "Future of cyberspace: A critical review of standard security protocols in the post-quantum era," *Computer Science Review*, vol. 57, p. 100 738, 2025. doi: 10.1016/j.cosrev.2025.100738.
- [2] A. Aydeger, P. Zhou, S. Hoque, M. Carvalho, and E. Zeydan, "Mtdns: Moving target defense for resilient dns infrastructure," in *2025 IEEE 22nd Consumer Communications & Networking Conference (CCNC)*, IEEE, 2025, pp. 1–6. doi: 10.1109/CCNC54725.2025.10975971.
- [3] M. M. Nesary and A. Aydeger, "Vdns: Securing dns from amplification attacks," in *2022 IEEE International Black Sea Conference on Communications and Networking (BlackSeaCom)*, IEEE, 2022, pp. 102–106. doi: 10.1109/BlackSeaCom54372.2022.9858278.
- [4] C. Deccio and J. Davis, "Dns privacy in practice and preparation," in *Proceedings of the 15th International Conference on Emerging Networking Experiments And Technologies*, 2019, pp. 138–143. doi: 10.1145/3359989.3365435.

Challenges of DNS in the Post-Quantum Era: Improving Security with Post-Quantum TLS

- [5] N. Alnahawi, J. Müller, J. Oupický, and A. Wiesmaier, "A comprehensive survey on post-quantum tls," *IACR Communications in Cryptology*, 2024. **doi:** 10.62056/ahee0iuc.
- [6] D. Joseph, R. Misoczki, M. Manzano, *et al.*, "Transitioning organizations to post-quantum cryptography," *Nature*, vol. 605, no. 7909, pp. 237–243, 2022. **doi:** 10.1038/s41586-022-04623-2.
- [7] Y. Hanna, D. Pineda, K. Akkaya, A. Aydeger, R. Harrilal-Parchment, and H. Albalawi, "Performance evaluation of secure and privacy-preserving dns at the 5g edge," in *2023 IEEE 20th International Conference on Mobile Ad Hoc and Smart Systems (MASS)*, IEEE, 2023, pp. 89–97. **doi:** 10.1109/MASS58611.2023.00019.
- [8] L. Malina, P. Dobias, J. Hajny, and K.-K. R. Choo, "On deploying quantum-resistant cybersecurity in intelligent infrastructures," in *Proceedings of the 18th International Conference on Availability, Reliability and Security*, 2023, pp. 1–10. **doi:** 10.1145/3600160.3605038.
- [9] A. Aydeger, E. Zeydan, A. K. Yadav, K. T. Hemachandra, and M. Liyanage, "Towards a quantum-resilient future: Strategies for transitioning to post-quantum cryptography," in *2024 15th International Conference on Network of the Future (NoF)*, IEEE, 2024, pp. 195–203. <https://doi.org/10.1109/NoF62948.2024.10741441>.
- [10] M. Müller, J. de Jong, M. van Heesch, B. Overeinder, and R. van Rijswijk-Deij, "Retrofitting post-quantum cryptography in internet protocols: A case study of dnssec," *ACM SIGCOMM Computer Communication Review*, vol. 50, no. 4, pp. 49–57, 2020. **doi:** 10.1145/3431832.3431838.
- [11] S. W. S. Pan, D. D. N. Nguyen, R. Doss, W. Armstrong, P. Gauravaram, *et al.*, "Double-signed fragmented dnssec for countering quantum threat," *arXiv preprint arXiv:2411.07535*, 2024. **doi:** 10.48550/arXiv.2411.07535.
- [12] M. Raavi, S. Wuthier, and S.-Y. Chang, "Securing post-quantum dnssec against fragmentation misassociation threat," in *ICC 2024-IEEE International Conference on Communications*, IEEE, 2024, pp. 97–102. **doi:** 10.1109/icc51166.2024.10622607.
- [13] J. Goertzen and D. Stebila, "Post-quantum signatures in dnssec via request-based fragmentation," in *International Conference on Post-Quantum Cryptography*, Springer, 2023, pp. 535–564. **doi:** 10.1007/978-3-031-40003-2_20.
- [14] C. McGowan, J. Liu, and S. Ruj, "Security considerations for post-quantum signatures in dnssec via request-based fragmentation," in *Companion Proceedings of the ACM on Web Conference 2025*, 2025, pp. 1189–1193. **doi:** 10.1145/3701716.3715585.
- [15] C. Schutijser, E. E. Lastdrager, R. Koning, and C. E. Hesselman, "A testbed to evaluate quantum-safe cryptography in dnssec," in *DNS and Internet Naming Research Directions, DINR 2024*, 2024.
- [16] A. S. Rawat and M. P. Jhanwar, "Quantum-safe signatureless dnssec," *Cryptology ePrint Archive*, 2024. **doi:** 10.1145/3708821.3710837.
- [17] A. S. Rawat and M. P. Jhanwar, "Post-quantum dnssec with faster tcp fallbacks," in *International Conference on Cryptology in India*, Springer, 2024, pp. 212–236. **doi:** 10.1007/978-3-031-80311-6_11.
- [18] S. Jafarli, "Providing dns security in post-quantum era with hash-based signatures," M.S. thesis, University of Twente, 2022.
- [19] G. Beernink, "Taking the quantum leap: Preparing dnssec for post quantum cryptography," M.S. thesis, University of Twente, 2022.
- [20] J. Goertzen, "Enabling post-quantum signatures in dnssec: One arrf at a time," M.S. thesis, University of Waterloo, 2022.
- [21] B. Ali and G. Chen, "Titan-doh: Trust-integrated threat adaptive network for post-quantum secure dns over https," *Available at SSRN 5230452*, **doi:** 10.2139/ssrn.5230452.
- [22] M. J. Kannwischer, M. Krausz, R. Petri, and S.-Y. Yang, *Pqm4: Benchmarking NIST additional post-quantum signature schemes on microcontrollers*, Cryptology ePrint Archive, Paper 2024/112, 2024.
- [23] G. Fitzgibbon and C. Ottaviani, "Constrained device performance benchmarking with the implementation of post-quantum cryptography," *Cryptography*, vol. 8, no. 2, p. 21, 2024. **doi:** 10.3390/cryptography8020021.
- [24] M. Abbasi, F. Cardoso, P. Váz, J. Silva, and P. Martins, "A practical performance benchmark of post-quantum cryptography across heterogeneous computing environments," *Cryptography*, vol. 9, no. 2, p. 32, 2025. **doi:** 10.3390/cryptography9020032.
- [25] B. Lonc, A. Aubry, H. Bakhti, M. Christofi, and H. A. Mehrez, "Feasibility and benchmarking of post-quantum cryptography in the cooperative its ecosystem," in *2023 IEEE Vehicular Networking Conference (VNC)*, IEEE, 2023, pp. 215–222. **doi:** 10.1109/vnc57357.2023.10136335.
- [26] M. Kumar, "Post-quantum cryptography algorithm's standardization and performance analysis," *Array*, vol. 15, p. 100 242, 2022. **doi:** 10.1016/j.array.2022.100242.
- [27] A. A. Giron, J. P. A. do Nascimento, R. Custódio, L. P. Perin, and V. Mateu, "Post-quantum hybrid kemtls performance in simulated and real network environments," in *International Conference on Cryptology and Information Security in Latin America*, Springer, 2023, pp. 293–312. **doi:** 10.1007/978-3-031-44469-2_15.
- [28] C. Paquin, D. Stebila, and G. Tamvada, "Benchmarking post-quantum cryptography in tls," in *Post-Quantum Cryptography: 11th International Conference, PQCrypto 2020*, Paris, France, April 15–17, 2020, Proceedings 11, Springer, 2020, pp. 72–91. **doi:** 10.1007/978-3-030-44223-1_5.
- [29] P. Schwabe, D. Stebila, and T. Wiggers, "Post-quantum tls without handshake signatures," in *Proceedings of the 2020 ACM SIGSAC Conference on Computer and Communications Security*, 2020, pp. 1461–1480. **doi:** 10.1145/3372297.3423350.
- [30] D. Stebila and M. Mosca, "Post-quantum key exchange for the internet and the open quantum safe project," in *International Conference on Selected Areas in Cryptography*, Springer, 2016, pp. 14–37. **doi:** 10.1007/978-3-319-69453-5_2.
- [31] Meta Engineering, *Post-Quantum Readiness: TLS PQR at Meta*, [Online]. Available: <https://engineering.fb.com/2024/05/22/security/post-quantum-readiness-tls-pqr-meta/>, Accessed: August 2024, 2024.
- [32] A. B. Downey, "Lognormal and pareto distributions in the internet," *Computer Communications*, vol. 28, no. 7, pp. 790–801, 2005. **doi:** 10.1016/j.comcom.2004.11.001.
- [33] D. Mödinger, J.-H. Lorenz, R. W. van der Heijden, and F. J. Hauck, "Unobtrusive monitoring: Statistical dissemination latency estimation in bitcoin's peer-to-peer network," *Plos one*, vol. 15, no. 12, e0243475, 2020. **doi:** 10.1371/journal.pone.0243475.
- [34] R. Wallace, X. G. Andrade, P. Kayser, *et al.*, "Models of network delay," in *International Workshop on Statistical Modelling*, Springer, 2024, pp. 231–238. **doi:** 10.1007/978-3-031-65723-8_36.
- [35] R. Jain, *The art of computer systems performance analysis*. john wiley & sons, 1990.
- [36] M. Raavi, S. Wuthier, P. Chandramouli, Y. Balytskyi, X. Zhou, and S.-Y. Chang, "Security comparisons and performance analyses of post-quantum signature algorithms," in *International Conference on Applied Cryptography and Network Security*, Springer, 2021, pp. 424–447. **doi:** 10.1007/978-3-030-78375-4_17.
- [37] J. Henrich, A. Heinemann, A. Wiesmaier, and N. Schmitt, "Performance impact of pqc kems on tls 1.3 under varying network characteristics," in *International Conference on Information Security*, Springer, 2023, pp. 267–287. **doi:** 10.1007/978-3-031-49187-0_14.
- [38] S. Hoque, A. Aydeger, and E. Zeydan, "Exploring post quantum cryptography with quantum key distribution for sustainable mobile network architecture design," in *Proceedings of the 4th Workshop on Performance and Energy Efficiency in Concurrent and Distributed Systems*, 2024, pp. 9–16. **doi:** 10.1145/3659997.3660033.
- [39] A. Aydeger, S. Hoque, E. Zeydan, and K. Dev, "Analysis of robust and secure dns protocols for iot devices," *arXiv preprint arXiv:2502.09726*, 2025. **doi:** 10.48550/arXiv.2502.09726.
- [40] W. Castryck and T. Decru, "An efficient key recovery attack on sidh," in *Annual international conference on the theory and applications of cryptographic techniques*, Springer, 2023, pp. 423–447. **doi:** 10.1007/978-3-031-30589-4_15.

- [41] L. De Feo, D. Kohel, A. Leroux, C. Petit, and B. Wesolowski, "Sqsign: Compact post-quantum signatures from quaternions and isogenies," in *International conference on the theory and application of cryptology and information security*, Springer, 2020, pp. 64–93.
doi: 10.1007/978-3-030-64837-4_3.
- [42] A. Hounsel, K. Borgolte, P. Schmitt, J. Holland, and N. Feamster, "Comparing the effects of dns, dot, and doh on web performance," in *Proceedings of The Web Conference 2020*, 2020, pp. 562–572.
doi: 10.1145/3366423.3380139.
- [43] S. Hoque, A. Aydeger, and E. Zeydan, "Post-quantum secure ue-to-ue communications," in *2024 15th International Conference on Network of the Future (NoF)*, IEEE, 2024, pp. 28–30.
doi: 10.1109/nof62948.2024.10741456.



Abdullah Aydeger is currently an assistant professor at the Electrical Engineering and Computer Science Department at Florida Institute of Technology. Prior to joining Florida Tech in August 2022, he was an assistant professor at the School of Computing at Southern Illinois University, Carbondale, since 2020. Dr. Aydeger obtained a Ph.D. Degree in Electrical and Computer Engineering from Florida International University in 2020. His research interests are network security, network virtualization, and post-quantum cryptography.



Sanzida Hoque is a Ph.D. student in Computer Science at the Florida Institute of Technology, currently working as a Graduate Assistant. She completed her MS in Computer Science from the University of Missouri-St. Louis in August 2023 and her B.Sc. in Computer Science and Engineering from the Military Institute of Science and Technology (MIST) in March 2019. She worked as a Graduate Research Assistant from the Fall of 2021 to the Summer of 2023. Her research interests span various areas, including Network Systems, Post-Quantum Cryptography, DNS Security, Next-Generation Internet, Wireless Networks, Mobile/Edge Computing, IoT, SDN, and related fields. Website: sites.google.com/view/sanzidahoque



Engin Zeydan received his Ph.D. degree in February 2011 from the Department of Electrical and Computer Engineering at Stevens Institute of Technology, Hoboken, NJ, USA. He is currently a Senior Researcher at Services as Networks (SaS) Research Unit in CTTC, Barcelona, Spain and Project Coordinator of the Horizon Europe Unity6G European Project (January 2025-present). He was the Project Coordinator of the European H2020 MonB5G Project [2021-2023]. His research interests are in the areas of telecommunications, data engineering and network security.

OSINT Based Recognized Cyber Picture

Gergo Gyebnar

Abstract—The Recognized Cyber Picture (RCP) is a critical NATO initiative aimed at enhancing situational awareness in the cyber domain by consolidating intelligence on cyber threats, vulnerabilities, and adversarial tactics. This study explores the feasibility of developing an RCP based on Open-Source Intelligence (OSINT), addressing the absence of publicly available implementations. Leveraging frameworks such as MITRE ATT&CK, methodologies like threat intelligence, weighted scoring model and detection as code concept. The research highlights how OSINT can complement classified data by identifying and prioritizing threats. Through targeted intelligence the study maps adversarial Tactics, Techniques, and Procedures (TTPs) to critical military operations. The findings underscore the importance of a scalable, resource-efficient RCP to counter increasingly sophisticated hybrid warfare threats effectively. The study also suggests onboarding and maintenance methodologies via cutting edge technology called Detection-as- Code.

Index Terms—Cybersecurity, Detection-as-Code, MITRE ATT&CK, OSINT, Recognized Cyber Picture, Situational Awareness

I. INTRODUCTION

Military operations are increasingly characterized using hybrid warfare tactics, where cyberspace plays a crucial role as a battlefield. Hybrid warfare reflects the complexities of contemporary international relations, which have evolved into a polyarchic structure. In this system, state actors often seek to advance their interests not through conventional warfare but via hybrid methods, including cyber capabilities. [1].

The polyarchic nature of today's global order emphasizes the interconnectedness and interdependence of states, which complicates traditional notions of power and conflict. Cyberattacks exemplify this dynamic, offering states the ability to influence, disrupt, or coerce adversaries without engaging in direct military confrontation. Unlike conventional military engagements, cyber operations allow actors to obscure their involvement, leveraging state-supported hacker groups to carry out attacks. These operations often target critical infrastructure, communications, or strategic information, creating significant disruption while maintaining plausible deniability.

Such tactics blur the lines between state and non-state action, and between peace and war. The lack of clear attribution and accountability complicates international responses and raises the stakes for nations to develop robust cyber defense mechanisms. Within this context, NATO's Recognized Cyber Picture (RCP) concept gains heightened importance. As cyberspace becomes a central element of hybrid conflicts, RCP serves as a critical framework for situational awareness and coordination among member states, ensuring collective security in the face of evolving threats.

Faculty of Military Sciences and Officer Training, Ludovika University of Public Service Budapest, Hungary (E-mail: gergo.gyebnar@blackcell.io)

DOI: 10.36244/ICJ.2025.3.3

Threat actors employ increasingly sophisticated techniques, and potential attacks may target both traditional IT systems and complex military assets, including Operational Technology (OT) or Industrial Control Systems (ICS). Additionally, the RCP must be flexible enough to address the unique security and privacy needs of each member state, while still aligning with NATO's overarching cybersecurity strategy.

The RCP leverages various frameworks, such as the MITRE ATT&CK framework, to map and prioritize known threats and adversarial techniques, while also providing a structure for intelligence sharing and coordination. Using a framework approach allows NATO to classify, score, and visualize the risk landscape through methods like heatmaps, which help in setting prioritized defense actions across different operational levels.

Regarding RCPs, which primarily rely on classified data, my work is constrained to open-source intelligence (OSINT), which inherently has limitations. To enhance its effectiveness, OSINT should ideally be supplemented with signals intelligence (SIGINT), human intelligence (HUMINT), and social media intelligence (SOCMINT). Nevertheless, OSINT has demonstrated its effectiveness in this domain.

This paper aims to compile, assess, and illustrate the most used attack techniques targeting military targets in cyberspace. The resulting heatmaps presented here offer a unique visualization, though they require ongoing updates to remain effective. **This paper aims to explore this gap by examining the potential of an OSINT-based RCP and to provide military targets with a threat-informed defensive strategy.**

II. THEORETICAL BACKGROUND

In the evolving landscape of modern warfare, hybrid tactics—which blend conventional military strategies with unconventional and asymmetric approaches—have become a predominant mode of operation. Among these tactics, cyberspace has emerged as a critical domain for both offensive and defensive operations. Cyberattacks are increasingly leveraged to disrupt, degrade, and manipulate adversaries' critical infrastructure, command systems, and information networks, often blurring the line between peacetime cyber activities and acts of war.

A. Recognized Cyber Picture (RCP)

The RCP is NATO's efforts to address the challenges posed by cyberspace as a battlefield. RCP seeks to create a comprehensive and near real-time view of the cyber threat environment. Its primary objective is to enhance situational awareness by consolidating intelligence on cyber activities, vulnerabilities, and adversarial actions that may impact NATO's operational capabilities. Through the RCP, NATO

aims to provide military commanders and decision-makers with actionable insights to enable swift and informed responses to cyber incidents.

Despite its importance, there is currently not publicly available Open-Source Intelligence (OSINT)-based implementation of the RCP. These are confidential. This study seeks to provide initial insights and recommendations that could contribute to the development of such a system.

B. Components of the RCP

An effective RCP encompasses several critical elements. Threat intelligence aggregation to integrate intelligence to RCP to ensure a multifaceted understanding of cyber threats. Realtime monitoring and analysis that requires continuous monitoring of cyber activities provides the means to detect emerging threats, analyze attack patterns, and identify vulnerabilities to be able to create an actionable threat intel. The framework utilization where RCP leverages established frameworks, such as the MITRE ATT&CK framework, to classify, map, and prioritize adversarial techniques. Tools such as heatmaps visualize risks and guide decision-making on defense priorities. It facilitates structured intelligence sharing and operational coordination among NATO members, ensuring alignment with NATO's overarching cybersecurity strategy.

C. Challenges in Implementing the RCP

Developing and maintaining an effective RCP is a complex undertaking due to the dynamic nature of cyber threats and the diversity of NATO's member states. Key challenges include the sophistication of threat actors because adversaries employ advanced techniques to target both traditional IT systems and operational technology (OT) assets used in military environments. They may also change their adversarial infrastructure and develop new TTPs. Interoperability is also a kind of challenge since member states' varied cybersecurity policies, capabilities, and legal frameworks necessitate a flexible yet cohesive approach to intelligence sharing and response coordination. The rapid pace of cyber threat evolution demands continuous updates to the RCP, requiring substantial investment in personnel, technology, and training.

D. Hypothesis

Given the increasing sophistication of hybrid warfare and the integral role of cyberspace as a modern battlefield, the development of an OSINT-based RCP can provide a viable, flexible, and scalable solution for enhancing situational awareness and decision-making capabilities. While the RCP concept has been central to NATO's cybersecurity framework, there is currently no publicly available implementation of an OSINT-based RCP. This gap presents a significant opportunity to explore how OSINT, combined with proven frameworks like the MITRE ATT&CK, can be leveraged to aggregate, analyze, and visualize cyber threats relevant to military operations.

The hypothesis is grounded in prior studies and frameworks utilized in cybersecurity, particularly in the electric sector.

The Pyramid of Pain (D. Bianco, 2013) [2] emphasizes the varying levels of effort adversaries face in cyber operations, which can inform the prioritization of intelligence in an RCP. MITRE ATT&CK Framework (The MITRE Corporation, 2022) [6] is widely recognized as a structured approach to identifying adversarial techniques, offering a foundation for mapping threats in the electric sector, such as examples of cyberattacks on power grids, including Crashoverride [3] and incidents involving APT groups like Sandworm Team [4], demonstrate the utility of structured intelligence-sharing platforms and visualization techniques for proactive threat management.

Techniques from studies like MISP [5] have demonstrated the efficacy of visualizing gaps in cyber defenses and patterns of adversarial behavior. By integrating OSINT into an RCP model, this study aims to demonstrate how OSINT can fill gaps that arise due to resource constraints or the unavailability of classified data. It seeks to establish methods for visualizing and prioritizing threats using tools like heatmaps and the MITRE ATT&CK Navigator [6].

III. METHODOLOGY

Many cybersecurity assessment frameworks exist, but some focus more on compliance, lack objectivity, or are immature. In contrast, the MITRE ATT&CK Framework is a widely accepted knowledge base detailing adversary tactics, techniques, and procedures (TTPs). TTPs describe adversarial behaviors, processes, actions, and workflows used to infiltrate target infrastructures. This paper will specifically examine the "Enterprise" domain within MITRE ATT&CK v16, with the Enterprise domain covering 14 tactics, 203 techniques, and 453 sub-techniques. This approach is among the most pragmatic for defending critical infrastructure today. [7]

The results highlight the critical techniques that entities should monitor, alert and respond to effectively. To accomplish this, organizations need visibility into detailed datasets. Without relevant data, effective detection is impossible.

Drawing on historical data and analysis of cyberattacks, this research aggregates actionable methods for detecting adversary Tactics, Techniques, and Procedures (TTPs).

The methodology followed a structured approach to identify and analyze incidents. Initially, incidents of interest were identified through targeted queries, followed by a deeper investigation phase that reduced irrelevant data and focused on relevant materials. Each incident was analyzed to identify specific malware, tools, and techniques used, and to extract patterns and profiles linked to adversary groups.

Using the MITRE ATT&CK framework, specific TTPs were precisely identified and mapped to tactics. MITRE ATT&CK Navigator was then used to visualize threat actor tactics and techniques across various operational domains. This tool allows for creating customized views, scoring, color-coding, and filtering to highlight critical areas for analysis and defense.

To assign priority levels, each layer in the matrix was evaluated and assigned weighted scores based on several

factors. The **Impact Score**, ranging from 1 to 5, measures the severity of a technique, from minor disruptions to life-threatening consequences. The **Evasion Score**, also on a scale of 1 to 5, assesses how effectively a technique can avoid detection mechanisms. The **Complexity Score** reflects the skill level and resources required to execute a given technique, again ranging from 1 to 5. Historical data informed the **Historical Successfulness Score**, which indicates the past effectiveness of a technique, scored on the same scale. Additionally, the **Data Accuracy Score**, a multiplier ranging from 0.5 to 1.5, adjusts the weighting to reflect the reliability of the data. The scores based on the identified threats are shown in *Table 1. Scoreboard*

For visualization, a Red-Amber-Green (RAG) scoring model was employed to convey priority levels in an intuitive and easily interpretable format.

IV. ANALYZED INCIDENTS, THREAT ACTORS AND IDENTIFIED TOOLS OR MALWARE

A. Andariel

Andariel is a North Korean state-sponsored cyber threat group active since at least 2009. This group has primarily targeted South Korean government agencies, military entities, and various domestic companies, often employing destructive tactics. Notable operations attributed to Andariel include Operation Black Mine, Operation GoldenAxe, and Campaign Rifle. Andariel is regarded as a subset of the broader Lazarus Group and is linked to North Korea's Reconnaissance General Bureau. [8].

B. APT-C-23

APT-C-23 is a cyber threat group active since at least 2014, primarily targeting entities in the Middle East, including Israeli military assets. Since 2017, APT-C-23 has developed and deployed mobile spyware aimed at both Android and iOS devices, enabling extensive surveillance and data collection capabilities [8].

C. APT1

APT1 is a Chinese cyber threat group attributed to Unit 61398, a division of the 2nd Bureau within the 3rd Department of the People's Liberation Army (PLA) General Staff Department (GSD) [8].

D. APT28

APT28, a cyber threat group linked to Russia's General Staff Main Intelligence Directorate (GRU), specifically the 85th Main Special Service Center (GTSS), military unit 26165, has been active since at least 2004. This group reportedly targeted the Hillary Clinton campaign, the Democratic National Committee (DNC), and the Democratic Congressional Campaign Committee (DCCC) in 2016 to interfere with the U.S. presidential election. Some of these activities were conducted in coordination with GRU Unit 74455, also known as the Sandworm Team [8].

E. Confucius

Confucius is a cyber espionage group active since at least 2013, primarily targeting military personnel, high-profile individuals, business figures, and government organizations across South Asia. Security researchers have observed

similarities between Confucius and the threat group Patchwork, particularly in their use of custom malware code and their choice of targets [8].

F. Gallmaker

Gallmaker is a cyberespionage group active since at least December 2017, primarily targeting organizations in the Middle East. The group has focused its efforts on the defense, military, and government sectors, aiming to infiltrate and gather intelligence from these high-value targets [8].

G. Gamaredon Group

Gamaredon Group is a suspected Russian cyber espionage threat group that has been active since at least 2013, primarily targeting military, NGO, judiciary, law enforcement, and non-profit organizations in Ukraine. The group's name originates from an early campaign where the word "Armageddon" was misspelled as "Gamaredon." In November 2021, the Ukrainian government officially attributed Gamaredon Group to Center 18 of Russia's Federal Security Service (FSB) [8].

H. Ke3chang

Ke3chang is a Chinese threat group active since at least 2010, targeting entities in the oil sector, government, diplomatic, military, and NGOs. Its operations have spanned Central and South America, the Caribbean, Europe, and North America, focusing on intelligence gathering from high-value organizations across these regions [8].

I. Machete

Machete is a suspected Spanish-speaking cyber espionage group active since at least 2010. The group has primarily focused its operations on Latin America, with a particular emphasis on Venezuela, while also targeting entities in the United States, Europe, Russia, and parts of Asia. Machete typically targets high-profile organizations, including government institutions, intelligence agencies, military units, as well as telecommunications and power companies, seeking to gather sensitive information [8].

J. Magic Hound

Magic Hound is an Iranian-sponsored cyber threat group that conducts long-term, resource-intensive cyber espionage operations, likely on behalf of the Islamic Revolutionary Guard Corps (IRGC). Active since at least 2014, the group has targeted government and military personnel, academics, journalists, and organizations such as the World Health Organization (WHO) across Europe, the United States, and the Middle East. Their operations typically involve complex social engineering campaigns aimed at gaining access to sensitive information [8].

K. Mofang

Mofang is a likely China-based cyber espionage group known for its tactic of imitating a victim's infrastructure to carry out attacks. Active since at least May 2012, Mofang has conducted targeted operations against government entities and critical infrastructure in Myanmar, as well as other sectors and countries, including military, automobile, and weapons industries. The group's activities are primarily focused on intelligence gathering through sophisticated and stealthy cyber intrusions [8].

L. Naikon

Naikon is a state-sponsored cyber espionage group attributed to the Chinese People's Liberation Army's (PLA) Chengdu Military Region Second Technical Reconnaissance Bureau (Military Unit Cover Designator 78020). Active since at least 2010, Naikon has primarily targeted government, military, and civil organizations in Southeast Asia, as well as international entities such as the United Nations Development Programme (UNDP) and the Association of Southeast Asian Nations (ASEAN). While Naikon shares some similarities with APT30, the two groups are not considered identical in their tactics or operations [8].

M. Sandworm Team

Sandworm Team is a destructive cyber threat group attributed to Russia's General Staff Main Intelligence Directorate (GRU), specifically to the Main Center for Special Technologies (GTsST), military unit 74455. Active since at least 2009, Sandworm Team has carried out a series of high-profile cyber operations. In October 2020, the United States indicted six officers from GRU Unit 74455 in connection with several major attacks, including the 2015 and 2016 cyberattacks on Ukrainian electrical companies and government organizations, the global NotPetya attack in 2017, interference with the 2017 French presidential campaign, the 2018 Olympic Destroyer attack on the Winter Olympic Games, an operation against the Organisation for the Prohibition of Chemical Weapons in 2018, and attacks on Georgia in 2018 and 2019. Some of these operations were carried out in collaboration with GRU Unit 26165, also known as APT28 [13].

N. Sidewinder

Sidewinder is a suspected Indian threat actor group active since at least 2012. The group has primarily targeted government, military, and business entities across Asia, with a particular focus on Pakistan, China, Nepal, and Afghanistan. Sidewinder's operations are believed to involve cyber espionage, aimed at gathering sensitive information from these high-value targets [8].

O. The White Company

The White Company is a likely state-sponsored threat actor with advanced cyber capabilities. Between 2017 and 2018, the group conducted an espionage campaign known as Operation Shaheen, primarily targeting government and military organizations in Pakistan. The operation involved sophisticated techniques aimed at gathering intelligence from these high-value sectors [8].

P. Tonto Team

Tonto Team is a suspected Chinese state-sponsored cyber espionage group that has been active since at least 2009. Initially focused on South Korea, Japan, Taiwan, and the United States, the group expanded its operations by 2020 to include other Asian and Eastern European countries. Tonto Team has targeted a wide range of sectors, including government, military, energy, mining, financial, education, healthcare, and technology organizations. Notable operations linked to the group include the Heartbeat Campaign (2009-2012) and Operation Bitter Biscuit (2017), which involved sophisticated cyber espionage tactics [8].

Q. WIRTE

WIRTE is a threat group active since at least August 2018, known for targeting government, diplomatic, financial, military, legal, and technology organizations across the Middle East and Europe. The group is believed to conduct cyber espionage operations, focusing on high-value sectors to gather sensitive information [8].

R. Lotus Blossom

Lotus Blossom is a cyber threat group known for targeting government and military organizations across Southeast Asia. The group's operations focus on intelligence gathering, often aimed at supporting geopolitical interests in the region [9].

S. Turla

Turla is a cyber espionage threat group attributed to Russia's Federal Security Service (FSB). Active since at least 2004, Turla has compromised victims in over 50 countries, targeting a wide range of industries, including government, embassies, military, education, research, and pharmaceutical sectors. The group is known for conducting watering hole and spearphishing campaigns, and for using custom-developed tools and malware, such as Uroburos, to carry out sophisticated intelligence-gathering operations [10].

T. ToddyCat

ToddyCat is a sophisticated threat group active since at least 2020. The group employs custom loaders and malware in multi-stage infection chains, primarily targeting government and military organizations across Europe and Asia. ToddyCat is known for its advanced tactics and persistence in infiltrating high-value targets to conduct cyber espionage operations [11].

U. Agent.btz

Agent.btz is a worm that spreads primarily through removable devices, such as USB drives. It is known for having infected U.S. military networks in 2008, causing significant concerns due to its ability to propagate through unprotected systems and compromise sensitive networks [12].

V. COATHANGER

COATHANGER is a remote access tool (RAT) designed to target FortiGate networking appliances. First identified in 2023, it was used in targeted intrusions against military and government entities in the Netherlands, among other victims. The malware was publicly disclosed in early 2024, with high confidence linking it to a state-sponsored actor from the People's Republic of China. COATHANGER is typically delivered after compromising a FortiGate device, with exploitation of CVE-2022-42475 being associated with its in-the-wild use. The name "COATHANGER" comes from a unique string in the malware that is used to encrypt configuration files: "She took his coat and hung it up." [12].

W. DOGCALL

DOGCALL is a backdoor malware used by the APT37 threat group, primarily targeting South Korean government and military organizations. First observed in 2017, DOGCALL is typically deployed via an exploit in Hangul Word Processor (HWP) documents, enabling the attackers to maintain persistent access to compromised systems for cyber espionage activities [12].

X. Gootloader

Gootloader is a JavaScript-based infection framework active since at least 2020, primarily used as a delivery method for various malicious payloads, including the Gootkit banking trojan, Cobalt Strike, REvil ransomware, and others. Operating on an "Initial Access as a Service" model, Gootloader leverages SEO poisoning techniques to gain access to a wide range of entities across multiple sectors globally, including financial, military, automotive, pharmaceutical, and energy industries. This approach allows attackers to compromise organizations through malicious websites that appear in search engine results [12].

Y. Ninja

Ninja is a C++-developed malware used by the ToddyCat threat group to penetrate networks and remotely control compromised systems since at least 2020. It is believed to be part of a post-exploitation toolkit exclusively used by ToddyCat, enabling multiple operators to work simultaneously on the same machine. Ninjas have been deployed in attacks against government and military entities in Europe and Asia and have been observed in specific infection chains deployed by another group, Samurai [12].

Z. ShimRat

ShimRat is a malware used by the suspected China-based threat group Mofang in cyber campaigns targeting various sectors, including government, military, critical infrastructure, automobile, and weapons development. The malware's name, "ShimRat," comes from its extensive use of Windows Application Shimming techniques to maintain persistence on compromised systems. ShimRat is part of a broader espionage effort aimed at gathering sensitive information from high-value targets across multiple countries [12].

AA. ShimRatReporter

ShimRatReporter is a tool used by the suspected Chinese threat group Mofang to automate the initial discovery phase of cyber intrusions. The information gathered during this discovery is used to customize follow-up payloads, such as ShimRat, and to set up deceptive infrastructure that mimics the adversary's targets. ShimRatReporter has been deployed in campaigns targeting multiple countries and sectors, including government, military, critical infrastructure, automobile, and weapons development, facilitating sophisticated espionage operations [13].

BB. USBferry

USBferry is an information-stealing malware used by the threat group Tropic Trooper in targeted attacks against air-gapped military environments in Taiwan and the Philippines. While USBferry shares an overlapping codebase with another malware, YAHOOYAH, it includes several distinct features that differentiate it from its counterpart. USBferry is primarily used to steal sensitive information from compromised systems, with a particular focus on military targets in highly secure, isolated networks [13].

CC. C0011

C0011 was a suspected cyber espionage campaign conducted by the threat group Transparent Tribe, which primarily targeted students at universities and colleges in India. This campaign marked a significant shift from Transparent Tribe's usual focus on Indian government,

military, and think tank personnel. Security researchers noted that the C0011 campaign appeared to be ongoing as of July 2022, with the group using new tactics to compromise student networks and potentially gather intelligence [13].

DD. Operation Ghost

Operation Ghost was a cyber espionage campaign conducted by APT29, also known as the "Cozy Bear" group, beginning in 2013. The operation targeted ministries of foreign affairs in Europe as well as the Washington, D.C. embassy of a European Union country. During this campaign, APT29 employed new families of malware and utilized sophisticated techniques such as web services, steganography, and unique command-and-control (C2) infrastructure for each victim.[13].

V. SCOREBOARD

TABLE I.
SCOREBOARD

Layer	Score					
	Impact	Evasion	Complexity	Successfulness	Accuracy	SUM
Andariel	2	2	2	2	1	8
APT-C-23	2	3	3	3	1	11
APT 1	4	4	4	4	1	16
APT 28	4	5	5	4	1	18
Confucius	4	4	3	4	1	15
Gallmaker	3	3	3	2	1	11
Gamaredon Group	3	4	4	3	1	14
Kc3chang	4	3	4	3	1	14
Lotus Blossom	2	3	2	3	1.2	12
Machete	3	3	2	3	1	11
Magic Hound	3	4	4	4	1	15
Mofang	2	3	2	3	1	10
Naikon	4	3	3	3	1	13
Sandworm Team	5	5	5	4	1.2	22.8
Sidewinder	3	3	3	3	1	12
The White Company	3	4	2	4	1	12
ToddyCat	4	4	3	4	1.3	19.5
Tonto Team	4	4	3	3	1	14
Turla	4	5	5	5	1.4	26.6
WIRTE	3	3	2	3	1	11
Agent.btz	2	2	2	2	1	8
COATHANGER	3	2	3	4	1	12
DOGCALL	3	2	2	3	1	10
ShimRat	3	2	4	4	1	13
ShimRatReporter	3	2	3	3	1	11
USBferry	3	5	3	4	1	10
C0011	3	4	2	4	1	14
Operation Ghost	3	4	2	4	1	13

VI. HEATMAP

The analysis of military RCP has revealed a set of emerging Tactics, Techniques, and Procedures (TTPs) that adversaries increasingly leverage to compromise systems, evade detection, and maintain persistence. These TTPs span a wide range of tactics, from phishing and infrastructure acquisition to advanced scripting and exploitation techniques. Each TTP is accompanied by distinctive indicators that enable detection and mitigation, offering a roadmap a prioritized action plan for strengthening defenses against evolving cyber threats. The following breakdown highlights the most critical and actionable TTPs identified. The heatmap is shown below.

A. Emerging Threat TTP Analysis Based on Military RCP heatmap

1) *T1566.001 – Phishing: Spearphishing Attachment:* This TTP remains a cornerstone for adversarial initial access strategies. Key indicators include suspicious execution of HH.EXE, Microsoft OneNote spawning unexpected child processes, and file creations in Outlook Temp Folders. Additionally, activity involving Office macros generating files and ISO file operations in temporary directories signals potential spearphishing campaigns. Monitoring these patterns can disrupt early-stage compromises effectively.

2) *T1566.002 - Phishing: Spearphishing Link:* Like T1566.001, this method emphasizes adversaries leveraging links to execute payloads. Indicators overlap with suspicious HH.EXE executions but extend to HTML Help child processes, highlighting the necessity of scrutinizing link-based attachments and their follow-on behaviors.

3) *T1583.001 - Acquire Infrastructure: Domains:* This TTP is characterized by adversaries setting up malicious infrastructure. Observations primarily derive from Cyber Threat Intelligence (CTI) operations, as detection is complex without additional telemetry or proactive monitoring of domain registration trends.

4) *T1588.002 - Obtain Capabilities: Tool:* Adversaries demonstrate increasing sophistication using renamed tools, such as Sysinternals DebugView and RegistrySet utilities. These tactics allow malicious activities to blend with legitimate processes, demanding nuanced detection strategies to differentiate genuine tool usage from adversarial actions.

5) *T1189 - Drive-by Compromise:* Drive-by compromises often exploit user interactions with vulnerable web services. DNS rebinding attacks stand out as the primary indicator, requiring proactive monitoring of web server behaviors and DNS activity.

6) *T1047 - Windows Management Instrumentation (WMI):* Adversaries increasingly leverage WMI for lateral movement and persistence. Anomalies include WMIC executions, encoded scripts within WMI consumers, and unquoted service paths for reconnaissance, which warrant close inspection of WMI-related command-line activity and associated processes.

7) *T1204.001/002 - User Execution (Malicious Link or File):* These tactics involve the execution of malicious files or links delivered through spearphishing or drive-by techniques. Indicators include VBA DLL loading, creation of

binaries or files through Office applications, and abnormal activity from user directories or temporary locations. Detecting these patterns is crucial for mitigating user-targeted exploits.

8) *T1053.005 - Scheduled Task/Job: Scheduled Task:* Adversaries exploit scheduled tasks for persistence or payload execution. Abnormal task creation, modification, or disabling is a key indicator, particularly when tasks are linked to PowerShell-encoded payloads or registry-based persistence attempts.

9) *T1106 - Native API:* Monitoring adversaries' exploitation of WinAPI calls via PowerShell or command-line interfaces is vital. This includes observing deviations in expected script behavior or unusual API access attempts.

10) *T1203 - Exploitation for Client Execution:* This TTP highlights client-side vulnerabilities exploited via malicious documents or scripts. Detection focuses on sub-process activities and network connections initiated by Excel or similar applications, signaling potential exploitation.

11) *T1059.003 - Command and Scripting Interpreter: Windows Command Shell:* Adversaries use the Windows command shell for malicious execution. Notable patterns include path traversal anomalies and commands embedding suspicious URLs, often targeting AppData or temporary directories.

12) *T1059.005 - Command and Scripting Interpreter: Visual Basic:* Indicators focus on WScript or CScript droppers and unexpected parent processes for Csc.exe, emphasizing the importance of monitoring Visual Basic-based execution paths.

13) *T1059.001 - Command and Scripting Interpreter: PowerShell:* PowerShell remains a favored tool for adversaries due to its flexibility and obfuscation capabilities. Detection includes encoded commands, reverse shell activity, and obfuscated script patterns designed to evade traditional defenses.

14) *T1505.003 - Server Software Component: Web Shell:* Web shells are deployed to maintain persistent access. Suspicious ASPX file drops, SQL server anomalies, and webshell reconnaissance activities highlight the need for vigilant monitoring of server and database operations.

15) *T1547.001 - Boot or Logon Autostart Execution: Registry Run Keys / Startup Folder:* Registry and startup folder abuse enable adversaries to maintain persistence. Indicators include modifications to user shell folders, registry run key entries, and PowerShell-based startup shortcuts.

16) *T1027 - Obfuscated Files or Information:* Adversaries commonly obfuscate payloads to avoid detection. This includes Base64 encoding, renamed executables, and misuse of utilities like Certutil to download and encode files.

17) *T1036.005 - Masquerading: Match Legitimate Name or Location:* Masquerading involves using legitimate-looking names or locations for malicious files or processes. Indicators include system process names in unusual locations and misused execution paths.

18) *T1070.004 - Indicator Removal: File Deletion*: File deletion patterns associated with malicious activity include Prefetch file deletion and combinations of ping/del commands, suggesting attempts to erase evidence post-operation.

19) *T1140 - Deobfuscate/Decode Files or Information*: Decoding tools such as Certutil and MSHTA are leveraged for obfuscation reversal. These activities often precede more overt malicious execution.

20) *T1003.001 - OS Credential Dumping: LSASS Memory*: Credential dumping from LSASS memory is a persistent threat, with indicators such as Dump64.exe execution and memory extraction via Comsvcs.DLL.

21) *T1056.001 - Input Capture: Keylogging*: PowerShell-based keylogging tactics target sensitive input data. Monitoring unusual scripting behaviors can aid in early detection [17].

VII. DETECTION AS CODE

The Detection as Code (DaC) fundamentally transforms the processes on the defensive side. It offers numerous advantages over traditional, outdated rule-management methods, the most important of which are detailed in the following subsections. DaC is a transformative approach that enables automated, version-controlled detection rules tailored specifically for Military RCP. The primary goal is to facilitate quick adaptation to new threats and seamless integration with security tooling. This approach leverages version control (e.g., using Git) to manage detection rules and configurations, allowing teams to track changes and collaborate on detection development.

Automated processes accelerate the creation, testing, and deployment of rules, enabling faster responses to new threats. Equipping defense systems in a timely manner to counter new attacks is essential and can be the difference between a thwarted and a successful attack. Storing rules in a code-like format significantly enhances their clarity, interoperability, and ease of management. Rule sets stored in this manner represent a substantial advancement over the traditional, now outdated solutions used by defenders. F

Version control also simplifies teamwork and enables tracking of changes. Detection rules have designated owners who are responsible for creating and monitoring the lifecycle of the rules. Any fine-tuning or modifications made are automatically logged, making them fully traceable and auditable. Different versions resulting from specific client requirements or unique environmental characteristics can be managed more easily in a structured and organized manner. The use of version control systems opens new possibilities for collaboration, allowing detection rules to reflect the combined expertise.

A key advantage of code-like rule sets is that they can be annotated with MITRE ATT&CK mappings, enabling the creation of an automatically updated matrix that clearly reflects the current state of defense.

The built-in and extensible features of IDEs (such as syntax highlighting and code completion), multi-level automated testing and validation, and mandatory, traceable

peer reviews integrated into the process minimize risks arising from human error. [14].

By adopting automation principles, detection rules are automatically deployed to environments, ensuring consistency and repeatability across all devices and systems. Continuous Integration and Continuous Deployment (CI/CD) pipelines further enhance this by automating the testing, validation, and deployment of detection rules, ensuring that updates are made swiftly in response to new attack techniques and vulnerabilities. This approach, when combined with continuous asset discovery, vulnerability management, MITRE ATT&CK mapping creates a comprehensive and adaptive security framework.

VIII. DISCUSSION AND CONCLUSION

Compared to earlier findings, the current study provides advancements in conceptualizing an Open-Source Intelligence (OSINT)-based Recognized Cyber Picture (RCP). By contrast, this study leverages an OSINT-centric approach, highlighting its potential to address situational awareness gaps without reliance on classified data. The absence of an OSINT-based implementation of RCP at the time underscores the novelty of this research [15]. The results indicate that OSINT can indeed serve as a foundation for developing an RCP. The dynamic nature of hybrid warfare necessitates adaptive and proactive measures. OSINT provides a viable alternative to traditional intelligence by aggregating publicly available information, which is often overlooked in classified frameworks. This enhances the feasibility and operational value of OSINT in supporting threat mapping and situational awareness. Tools such as heatmaps and detection coverage analysis and Detection-as-code are borrowed from other sectors like the electric, prove effective in identifying and addressing gaps in cyber defense.

These methods enable actionable insights. For networks to become more resilient against cyberattacks, especially in critical sectors, it is essential to keep these repositories up to date with Detection as code concept. Utilizing frameworks like MITRE ATT&CK, integrating advanced detection technologies, continuously managing risks, and updating threat intelligence systems are key to strengthening defense capabilities.

1) Limitations

While this paper presents avenues for the development of an Open-Source Intelligence (OSINT)-based Recognized Cyber Picture (RCP), several limitations must be acknowledged. One of the most significant limitations is the absence of classified data in this research. While OSINT provides valuable insights, classified intelligence often contains critical details about advanced persistent threats (APTs), attack methods, and vulnerabilities. Without access to this information, the RCP may lack the depth and precision required for high-stakes military applications. The retrospective nature of OSINT poses another challenge. Threat intelligence derived from publicly available sources may lag the rapidly evolving cyber threat landscape. Adversaries frequently update their tactics, techniques, and procedures (TTPs), rendering historical data potentially outdated or less effective for real-time decision-making, thereby introducing latency. Another significant concern is

data quality, which I have addressed through the implementation of a data accuracy multiplier within my methodology. The current RCP methodology lacks a mechanism for incorporating detailed vulnerability information. Effective cyber situational awareness depends not only on understanding adversarial behavior but also on identifying and addressing vulnerabilities within critical infrastructure and systems. This gap could lead to blind spots in the RCP, undermining its utility in proactive defense strategies. Developing an OSINT-based RCP requires substantial technical expertise and resources. Many NATO member states may face challenges in implementing and maintaining the necessary infrastructure, further exacerbating disparities in cybersecurity capabilities across the alliance.



Gergő Gyebnár is the CEO of Black Cell. With over 15 years of experience in IT security, he brings valuable expertise to the role. As CEO, Gergő is responsible for leading the expansion of Black Cell's OT Security and Detection Engineering services worldwide. He is committed to delivering high-quality services and staying ahead of the curve in terms of technological advancements in the field of IT security.

REFERENCES

- [1] M. Boda, "Hybrid War: Theory and Ethics", AARMS – Academic and Applied Research in Military and Public Management Science. Budapest, vol. 23, no. 1., pp. 5–17, 2024,
doi: 10.32565/aarms.2024.1.1.
- [2] D. Bianco, "The Pyramid of Pain", Mar. 1, 2013. [Online]. Available: <https://detect-respond.blogspot.com/2013/03/the-pyramid-of-pain.html>
- [3] R. M. Lee, "Crashoverride – Analysis of the Threat to Electric Grid Operation", Dragos Inc., Hanover, MD, USA, Jun. 13, 2017. [Online]. Available: <https://www.dragos.com/wp-content/uploads/CrashOverride-01.pdf>
- [4] The MITRE Corporation, (Sep. 22, 2022) Sandworm Team. [Online]. Available: <https://attack.mitre.org/groups/G0034/>
- [5] MISP project, (2022, October 5) MISP Threat Sharing. [Online]. Available: <https://www.misp-project.org/>
- [6] The MITRE Corporation, (Sep. 19, 2022) ATT&CK Matrix for Enterprise. [Online]. Available: <https://attack.mitre.org/>
- [7] J. L. Kurtz, L. E. Damianos, R. Kozierok, and L. Hirschman, "The MITRE Map Navigation Experiment", MITRE Corporation, Bedford, MA, USA, Jun. 1., 1999. [Online]. Available: <https://dl.acm.org/doi/pdf/10.1145/323216.323365>
- [8] The MITRE Corporation, (2024) Groups. [Online]. Available: <https://attack.mitre.org/groups/>
- [9] R. Falcone, J. Grunzweig, J. Miller-Osborn, and R. Olson, "Operation Lotus Blossom", Palo Alto Networks, Santa Clara, CA, USA, 2015. [Online]. Available: https://www.paloaltonetworks.com/apps/pan/public/downloadResource?pagePath=/content/pan/en_US/resources/research/unit42-operation-lotus-blossom
- [10] Unit 42, "Threat Group Assessment: Turla (aka Pensive Ursa)", Palo Alto Networks, Santa Clara, CA, USA, Sep. 15, 2023. [Online]. Available: <https://unit42.paloaltonetworks.com/turla-pensive-ursa-threat-assessment/>
- [11] G. Dedola, D. Caldarella, A. Fedotov, and A. Gunkin, "ToddyCat: Keep calm and check logs", Kaspersky Lab, Moscow, Russia, Oct. 12, 2023. [Online]. Available: <https://securelist.com/toddycat-keep-calm-and-check-logs/110696/>
- [12] The MITRE Corporation, (2024) Software. [Online]. Available: <https://attack.mitre.org/software/>
- [13] The MITRE Corporation, (2024) Campaigns. [Online]. Available: <https://attack.mitre.org/campaigns/>
- [14] M. Roddie, G. J. Katz, and J. Deyalsingh, Practical Threat Detection Engineering: A hands-on guide to planning, developing, and validating detection capabilities. Birmingham, UK: Packt Publishing Ltd, 2023.
- [15] Kumar, Shekhar, Ummineni, Srikar Lyu, Ran Mondal, Sourav, Muthe, Laxman, Advanced Open-Source Intelligence (OSINT) Platform to Combat Misinformation on Social Media. [Online]. Available: <https://vtechworks.lib.vt.edu/items/fbd991c7-a005-4562-8670-053d686b5c68>

Cost Analysis of Mode Division Multiplexing Bidirectional Transmission System

Ahmed S. Mohamed¹, and Eszter Udvary^{1,2}

Abstract—This paper presents a cost analysis and performance evaluation of a Mode Division Multiplexing (MDM) system with Power Over Fiber (PWoF) for bidirectional transmission, to support high-capacity data and centralized power distribution in mobile networks. Two configurations are proposed: an unsymmetrical Multicore Fiber (MCF) and a Double Clad Fiber (DCF). The MCF setup features separate data and power transmission cores, reducing inter-core crosstalk, while the DCF setup combines data and power within a single fiber to simplify deployment. A cost model is developed to compare the capital expenditures (CapEx) of each configuration, revealing an estimated cost of 322,000 Euros for the MCF system and 212,000 Euros for the DCF system. Given that MDM technology is still in the research phase, commercially available equipment is limited, contributing to high initial costs. However, similar to Dense Wavelength Division Multiplexing (DWDM), the price of MDM is expected to decline over time as the technology matures. Additionally, both configurations leverage PWoF to centralize power generation at the Central Office (CO), enabling the use of renewable energy sources and supporting sustainable network infrastructure. This study highlights the potential of MDM with PWoF as a cost-effective, environmentally friendly solution for future high-capacity mobile networks.

Index Terms—Cost Model, Few-Mode Fiber, In-Building Solutions, Double Clad Fiber, Mode Division Multiplexing.

I. INTRODUCTION

The rise of LTE-based radio access networks has put the spotlight on Cloud Radio Access Network (CRAN) due to its potential for reducing both operating and capital costs for 5G and 6G networks. [1]. In CRAN, the Remote Radio Unit (RRU) and Baseband Unit (BBU) are separated, with the centralized BBU pool handling the processing, control, and management of a large number of RRUs. The link between the RRUs and BBUs, known as fronthaul, uses the Common Public Radio Interface (CPRI) for digital data transmission[2]. To keep up with the high demands of 4G/5G mobile communication services, especially those using Multiple Input Multiple Output (MIMO) antennas, WDM has been proposed as a solution for transporting CPRI data over the fronthaul. WDM is a good fit here because it provides high spectral efficiency, service transparency, and energy savings.

When it comes to deploying and maintaining systems on the RRU side, cost-effectiveness and colorless operation are essential for WDM-based fronthaul. However, the proposed laser sources for RRUs are expensive, which raises concerns

about their commercial viability [3]. MDM has attracted considerable interest as a promising solution to boost transmission capacity. This approach has been widely researched for high-speed optical transmission and access networks, aiming to support a larger number of users [4]. Unlike WDM, systems based on MDM are colorless by its nature, allowing them to be well-suited for network applications that need colorless operations. This feature can potentially lead to cost savings for the RRU [3]. Moreover, compared with the other emerging techniques like Space Division Multiplexing (SDM), MDM ensures higher spectral efficiency by utilizing multiple modes in a single fiber or core which will reduce the fiber footprint, thus reduces the cabling complexity compared to the parallel single mode or multimode fiber in SDM [5]. Additionally, powering the RRUs poses a challenge that can add to their complexity and, in turn, increase manufacturing costs [6]. To address the challenge of power, this research incorporates the PWoF concept. PWoF is a simple solution for transmitting both optical data and power over a single optical fiber. In mobile networks, PWoF offers a dual function: it supplies power to RRUs from a CO [6] while consolidating the power source within the CO. This approach is compared to other wireless domain solutions, such as the energy transfer methods in Radio Frequency ID (RFID), Near Field Communications (NFC) applications [7], or the use of several Power Stations (PSs) for energy transfer in Wireless Powered Communication Networks (WPCN) [8]. Adopting the PWoF approach helps reduce overall system complexity and lowers operational costs in mobile networks.

II. MOTIVATION

The rapid expansion of mobile networks, driven by 5G and beyond, demands high-capacity, cost-effective, and energy-efficient solutions. While WDM has been widely adopted for optical fronthaul, its reliance on wavelength-specific laser sources increases system costs and complexity. MDM presents a compelling alternative by enabling multiple data streams to share a single fiber, reducing the dependency on costly wavelength-tuned components. Another challenge in mobile fronthaul is power distribution to RRUs. Traditional methods require distributed power stations, adding operational overhead. PWoF with MDM centralizes power generation at the CO, simplifying deployment and lowering infrastructure costs. Despite its advantages, MDM adoption faces technical hurdles, such as modal dispersion and crosstalk, which require advanced Digital Signal Processing (DSP) techniques. This study evaluates the cost-effectiveness of MDM with PWoF, comparing different fiber configurations to assess its feasibility for mobile networks.

¹ Department of Networked Systems and Services, Faculty of Electrical Engineering and Informatics, Budapest University of Technology and Economics, Budapest, Hungary

² ELKH-BME Information Systems Research Group, Budapest, Hungary
(E-mail: asayed@hit.bme.hu, udvary.eszter@vik.bme.hu)

III. RELATED WORK

MDM has been widely investigated as an alternative to WDM for increasing transmission capacity. Chen et al. [3] demonstrated bidirectional mobile fronthaul using MDM, showcasing its ability to eliminate wavelength-specific dependencies. Mercy Kingsta et al. [9] explored multi-core and few-mode fiber implementations, proving their potential for high-capacity optical communication.

One of the main challenges of MDM is modal dispersion and crosstalk, which degrade signal integrity. Alon et al. [10] proposed equalization techniques, while Yaman et al. [11] demonstrated the advantages of Few-Mode Fibers (FMF) in reducing interference. Despite these advances, MDM is still in the research phase, with limited commercial deployment due to high initial costs and the need for complex DSP techniques. Li et al. [12] examined the CapEx advantages of multi-core fiber networks, concluding that as the technology matures, costs will decrease. This study extends existing research by providing a detailed cost model for MDM-based mobile fronthaul with PWO, evaluating its feasibility against traditional solutions.

IV. SYSTEM CONCEPT

MDM has recently attracted significant attention as an alternative technology to increase transmission capacity in high-speed optical transmission and access networks, enabling support for more users [3]. In MDM, light modes are orthogonal, allowing each mode to act as a distinct data carrier channel. This contrasts with conventional approaches in Single Mode Fibers (SMF) and Multimode Fibers (MMF), where the light itself serves as the data channel without multiplexing, or techniques like Coarse WDM (CWDM) and DWDM, where channels are separated by wavelength. MDM utilizes the multimode properties of MMFs to expand the number of data channels, though employing MMFs for long-distance transmission can result in energy transfer between modes, causing crosstalk and modal dispersion [10]. These effects degrade the quality of received data, leading to increased Bit Error Rate (BER) and limiting the achievable bandwidth-distance product [10]. To mitigate crosstalk and modal dispersion, FMF have been developed. FMFs have a core size that is larger than SMFs but smaller than MMFs, supporting only a limited number of modes, which helps reduce crosstalk and improve transmission quality [11].

Despite the advantages of MDM, several challenges remain that prevent its widespread commercial adoption. The processes for multiplexing and demultiplexing light modes are still evolving and require further research. Furthermore, sophisticated DSP techniques are needed to manage crosstalk and control modal dispersion, adding complexity to the system [13]. Figure 1 shows a mobile fronthaul system utilizing MDM, where each RRU is assigned a specific mode over a few-mode fiber in a shared session. This MDM-based system is inherently colorless, unlike WDM, making it well-suited for network applications that require colorless operation [3]. There remains potential to further improve MDM-based mobile systems.

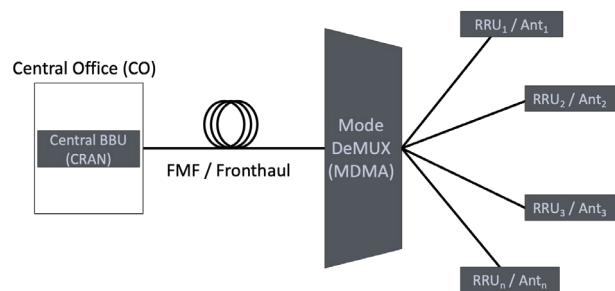


Fig. 1. Mobile System

We present an MDM system designed for bidirectional and symmetrical high-speed mobile fronthaul applications. In this configuration, one spatial mode transmits the signal carrier from the BBU pool, while another mode is utilized for downstream signal channels.

Unlike conventional wavelength-reused techniques, this system achieves symmetrical bidirectional transmission without introducing re-modulation crosstalk. Figure 2 depicts the proposed setup, showcasing the carrier along with RF/DL and RF/UL branches as part of the design.

Two optical modes are generated and delivered to the RRU through a FMF. The second mode, responsible for carrying downlink information, is detected at the RRU, amplified, and filtered before being sent as a radio frequency signal to the antenna. For the uplink, the signal collected from the antenna is filtered, amplified, and then used to modulate the first optical mode originating from the central unit. This modulated first mode is then directed back into the FMF using an optical circulator.

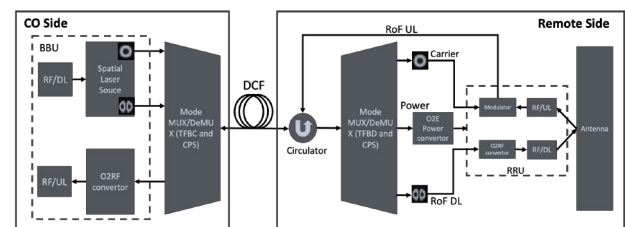


Fig. 2. Conceptual Model

To further simplify the RRU system, the proposed setup integrates the PWO concept, where the RRU is powered by optical energy generated at the CO, as shown in the power branch of Figure 2. This centralizes the power source at the CO, allowing mobile operators to manage power resources independently without depending on third parties. This approach is anticipated to reduce Operating Expenditures (OpEx), as discussed in later sections of this paper.

V. IMPLEMENTATION CHALLENGES

The suggested system comes with several technical challenges, which must be investigated and solved.

A. Double-clad and MultiCore few-mode fiber

The cost model under discussion in this study suggests using two types of waveguides. Dual Clad Fiber (DCF) where this concept provides the transmission of power required to drive the RRUs. Recent literature on this context shows that to

Cost Analysis of Mode Division Multiplexing Bidirectional Transmission System

obtain the high power required to drive the RRUs; larger fiber cores are needed, as the upper limit of the ultimate fiber damage power density is around 5 MW/cm² [14][15]. Moreover, the second approach suggests the use of Multicore Fiber MCF with two unsymmetrical cores, one ordinary FMF core for data transmission and another larger core for power transmission due to the same reasons mentioned earlier [14] [15].

A related concept is discussed in [15], where MCF with symmetrical core diameters is proposed instead of using DCF or asymmetrical MCFs. The approach explores two scenarios: in the first, one or more cores are dedicated to delivering the power needed to operate the RRU, while the remaining cores handle data transmission. In the second scenario, the same cores are shared for both power and data transmission simultaneously. Optical fibers with larger cores offer an advantage, as they can safely carry more power without the risk of damage. However, when high-power light is injected into small-core optical fibers, nonlinear effects like Stimulated Brillouin Scattering (SBS) and Stimulated Raman Scattering (SRS) may occur, significantly degrading the quality of concurrent data transmissions [6].

B. Required RRU Power

The PWOF link is constrained by the efficiency of both the optical source and the power converter at the RRU. Table 1 presents the power requirements of commercially available RRUs from various vendors.

TABLE I.
NEEDED POWER FOR THE COMMERCIAL RRUS

Model of RRU	Required Power	Manufacturer
Flexi MultiRadio 10	Up to 1KW	Nokia [16]
5G+4G Dot 4475	55W	Ericsson [17]
Street Macro 6701	Up to 400W	
AIR 1281	Up to 125W	
AIR 5322	Up to 220W	

Currently, available High-Power Laser Diodes (HPLDs) and Photovoltaic Converters (PPCs) require further research and testing to improve their conversion efficiencies. According to [18], efforts are being made to increase High-Density Laser Power (HDLP) efficiency to 74.7%, as current commercial laser sources reach only about 40%. For optimal Power Transmission Efficiency (PTE), the transmission wavelength should fall within the 8XX nm or 9XX nm range [19]. Additionally, the operating wavelength of the Photovoltaic Power Converter (PPC) must match the transmitted wavelength [19]. In [14], an optical feed of 150W is considered where several HDLPs used by coupling them into the fiber using Tapered Fiber Bundle Combiner (TFBC) and Divider (TFBD); Cladding Power Stripper (CPS) is used to separate the power feed from the data signal. Despite of the injected 150W, only 29.1W is transmitted to the DCF, yielding a PTE of 19.46%. On the receiving end, this results in an electrical power output of just 7.08W, with a PTE of 4.84%. Table 2 below summarizes the parameters of these power components.

TABLE II.
POWER COMPONENTS PARAMETERS

Component	Conversion Efficiency	Manufacturer
Commercial PPC	50% O/E	[20]
Commercial HPLD	40% E/O	[21]
Under research HPLD	74.7% E/O	[18]

Implementing this concept enables a compact RRU design, reducing overall system costs by moving many RRU functions to the CO, which further decreases RRU complexity and cost. Additionally, using the Radio over Fiber (RoF) approach helps minimize latency and increase capacity [22]. However, transmitting power and data simultaneously over the same fiber can lead to energy transfer from the inner, smaller core of the DCF to the outer, larger core, resulting in crosstalk, which can severely affect the quality of the received data. Recent research on PWOF using DCF is discussed in [23].

This research focuses on optimizing energy transfer from the "inner cladding" to the fiber core by identifying the ideal fiber design that maximizes energy absorption in the core. According to the study, certain fiber designs allow "skew-rays" to exit, which negatively affects the system by reducing energy absorption. These optical rays travel in a helical path around the fiber axis without entering the core, diminishing overall efficiency [23]. Moreover, in the case of using unsymmetrical MCF, less crosstalk is expected.

C. Modulation Bandwidth of few-mode transmission

As outlined in Section II, two light modes are necessary, one is designated for the uplink and the other for the downlink. The RF signals from both ends are transmitted via these optical signals using RoF technology. Each mode is modulated according to its respective RF signal, with the downlink signal from the BBU, and the uplink signal from the antenna. The required bandwidth varies depending on the specific mobile communication technology (4G, 5G, or 6G). Table 3 below lists the bandwidth requirements for each technology [24].

TABLE III.
POWER COMPONENTS PARAMETERS[24]

Technology	RF Frequency	Data Capacity
3G	1.6-2.5GHz	385Kbps-30Mbps
4G	2-8GHz	200Mbps-1Gbps
5G	>6GHz	>Gbps

D. Intra-Core and Inter-Core Crosstalk

Although crosstalk is reduced with FMF, energy still couples between modes during propagation [25], impacting the received signals for both downlink and uplink [26]. Proper DSP is needed to mitigate the effects of this crosstalk [27]. The intra-core crosstalk of a similar system has been extensively studied in a previous work where systems with FMF supporting two and four modes were considered [28]. Results show that the higher the order of the mode, the more crosstalk is observed, resulting in one of the main constraints of MDM systems. Moreover, in a system where MCF is adopted, inter-core crosstalk between the adjacent cores is expected. Fortunately, many literatures studied this issue proposing solutions to mitigate it by optimizing core and spectrum assignment [12] or by using MIMO algorithms [29].

VI. APPLIED OPTICAL WAVEGUIDES

We suggest using a specialized DCF with an FMF in the inner core, supporting two modes—one for downlink and one for uplink transmission—while the larger outer core is reserved for power signals. As outlined, the proposed system requires the excitation of two modes in the FMF, factoring in the Numerical Apertures (NAs) for commercially available DCFs of similar specifications [30]. Using the normalized frequency (V) equation, the required inner core diameter is calculated to be 10 μm . Alternatively, calculations can yield the same result by using core and cladding refractive indices from [31], where $n_1 = 1.45$ and $n_2 = 1.44$ at 1550 nm. For power transmission, the outer core of the DCF, with a diameter of 200 μm , is utilized. Additionally, an asymmetrical MCF will be considered in the second scenario, with cores similar to those in the DCF. In the DCF scenario, multiple HPLDs will be injected into the DCF in parallel. The design leverages commercially available HPLDs and PPCs; PPCs operate at 8XX nm with 50% optical-to-electrical (O/E) conversion efficiency [20], while HPLDs operate at 808 nm, handling a maximum input power of 100 W with 40% electrical-to-optical (E/O) conversion efficiency [21].

VII. SYSTEM COST MODEL

In this section, a cost comparison between the current DWDM commercial systems, DWDM system with PwOF on a dedicated fiber, and MDM with PwOF on a single fiber are shown. Regarding the DWDM system, a 40 channel DWDM system capable of handling a traffic of 400Gbps (40X10Gbps) is considered. On the other hand, an MDM system with the same capacity has been considered as well.

A. DWDM System

As previously mentioned, DWDM is the multiplexing technique of choice that being used for several years to handle the high-capacity needs of mobile networks. Having an exact cost for the DWDM system can vary depending on several business wise factors, vendors, and the location where these systems are implemented. However, and easier approach of having this calculation can be done by dividing these systems to the main parts that have the significant impact on the total cost. As compared with any other data transmission system.

A DWDM system consists of a transmitter, Medium, and a receiver. For the transmitter part, it is usually consisting of a laser source, multiplexer, and an optical amplifier. On the other hand, a receiver is consisting of a demultiplexer and photo detection device plus an amplifier. The component of a typical bidirectional DWDM system is showing in Figure 3 below.

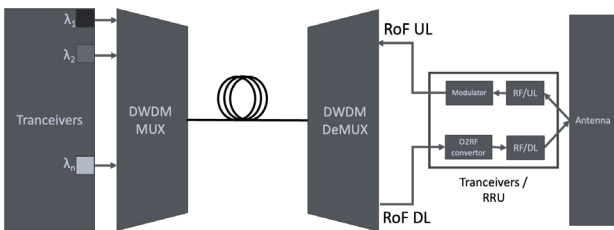


Fig. 3. DWDM System

Based on the above, the below formula can be used for the cost model:

Total System Cost

$$= (N_{Tx} \times C_{Tx}) + (N_{Rx} \times C_{Rx}) \\ + (N_{Mux/Demux} \times C_{Mux/Demux}) \\ + (N_{Amplifiers} \times C_{Amplifiers}) + \\ + (L_{Fiber} \times C_{Fiber})$$

Where:

N_{Tx} = Number of Transmitters

C_{Tx} = Cost per Transmitter

N_{Rx} = Number of Receivers

C_{Rx} = Cost per Receiver

$N_{Mux/Demux}$ = Number of Mux/Demux units

$C_{Mux/Demux}$ = Cost per Mux/Demux unit

$N_{Amplifiers}$ = Number of Amplifiers

$C_{Amplifiers}$ = Cost per Amplifier

L_{Fiber} = Total length of fiber cable (in meters)

C_{Fiber} = Cost per meter of fiber cable

In 5G networks, download speeds are anticipated to reach up to 20 Gbps, with upload speeds achieving 10 Gbps on average[32]. For a Four-floors building with approximately 40 active users simultaneously online, this would necessitate a data throughput of around 400 Gbps. Assuming a distance of 15 km between the building and the CRAN pool at the mobile operator's premises, the cost model in this study has been calculated accordingly to account for the infrastructure and operational demands over this span and as shown in Table 4 below.

TABLE IV.
DWDM COST MODEL

System Part	Price	Quantity	Manufacturer
Transceiver	308€	40	FS.com [33]
40CH MUX	2420€	1	FS.com [34]
Amplifier-Near-end	2808€	1	FS.com [35]
Fiber / m	1.5€	15000	ThorLabs.com [36]
Amplifier – Far-end	2808€	1	FS.com [35]
40CH DeMUX	2420€	1	FS.com [35]
Transceiver	308€	40	FS.com [34]

Based on the previous information, a DWDM system costs around 57 thousand Euros of CapEx excluding the labor cost of the fiber implementation as they vary a lot depending on the country where the system is being implemented considering no impact on the system cost itself and as follow:

Total Cost = Total system cost

+ (labor cost per Km $\times$ Fiber Length

Moreover, according to Eurostat website [37], the electricity average cost in Europe in 2023 is about 0.2€ per KWh. For Ericsson AIR 5322 [17] shown in table 1 above, the total power OpEx needed to drive this device is about 520€ yearly with a total of 2080€ yearly assuming installing one device per floor.

Cost Analysis of Mode Division Multiplexing Bidirectional Transmission System

B. DWDM system with PwoF on a dedicated fiber

This system is similar to the previous one, however, it has the capability of handling PWO_F on a dedicated optical fiber to drive the RRUs.

Table 5 below shows the cost model of the additional equipment needed for the power side.

TABLE V.
PWO_F COST MODEL

System Part	Price	Quantity	Manufacturer
HPLD	3122€	1	laserdiodesource.com [21]
Large Core Optical Fiber	0.95€	15000	Thorlabs.com [38]
PPC	10000€	1	MHGOPOWER.com [39]

The cost of the system is 57,000 Euros for the DWDM system plus 27,000 Euros for the power system, with a total of about 85,000 Euros excluding the labor cost of the fiber implementation.

C. Bidirectional MDM system with PWO_F over Unsymmetrical two cores MCF

This section shows the cost model of the first proposed system of this study where MDM system is being adopted with a two-core unsymmetrical MCF. Smaller two-modes FMF core for the data transmission and larger core for the power transmission, and as shown in Table 6 below.

TABLE VI.
MDM WITH UNSYMMETRICAL TWO-CORES MCF

System Part	Price	Quantity	Manufacturer
MDM Spatial Laser Source	12000€	2 (1/mode)	ThorLabs [40]
Photo Detector	350€	2 (1/mode)	ThorLabs [41]
MCF Fan-in	3500€	2	FiberCore [42]
MUX	7000€	1	CabiLabs [43], [44]
MultiCore FMF / m	18€	15000	FiberCore [42]
DeMUX	7000€	1	CabiLabs [43], [44]
HPLD	3122€	1	laserdiodesource.com [21]
PPC	10000€	1	MHGOPOWER.com [39]

As MDM is an emerging transmission multiplexing technique, optical amplifiers were not considered to reduce the system's cost for the relatively short fronthaul transmission, however, they can be used for long-haul ones to achieve higher distances. Moreover, laser sources (SLMs) prices are quite high due to the fact that they are currently available for lab usage only. Based on the above, the MDM system over 15 Km costs around 328,000 Euros of CapEx excluding the labor cost of the fiber implementation.

D. Bidirectional MDM system with PWO_F over DCF

This system is similar to the previous one with the difference of using different types of fiber where DCF is being utilized. As mentioned earlier, DCF can reduce the complexity even more as both power and data will be carried on the same fiber simultaneously. The cost model is shown in Table 7 below.

TABLE VII.
MDM WITH DCF

System Part	Price	Quantity	Manufacturer
MDM Spatial Laser Source / Mode	12000€	2 (1/mode)	ThorLabs [40]
Photo Detector /Mode	350€	2 (1/mode)	ThorLabs [41]
MDM MUX	7000€	1	CaiLabs [43], [44]
DCF / m	11.6€	15000	ThorLabs [30]
MDM DeMUX	7000€	1	CabiLabs [43], [44]

Based on the above, this MDM system cost around 212 thousand Euros of CapEx excluding the labor cost of the fiber implementation.

VIII. RESULTS AND DISCUSSION

This study evaluates two configurations for MDM with PWO_F: one based on DCF and the other on an unsymmetrical MCF. Each configuration offers distinct cost, performance, and sustainability benefits that make them promising solutions for high-capacity, centralized power transmission.

The DCF configuration, with an estimated CapEx of 212,000 Euros, consolidates both data and power transmission within a single fiber, simplifying infrastructure and reducing deployment complexity. In contrast, the MCF setup, with a projected cost of 322,000 Euros, separates data and power channels into distinct cores, which enhances resilience to interference and minimizes inter-core crosstalk. This unsymmetrical design is particularly advantageous for high-data-rate applications due to its added layer of signal isolation. A critical consideration for both configurations is crosstalk.

In the DCF system, intra-core crosstalk occurs when power and data signals interfere within the same core, which can degrade signal quality and increase the BER. Additionally, inter-core crosstalk arises when signals in adjacent cores interact, a risk present in both DCF and MCF setups. However, the unsymmetrical core design in the MCF configuration inherently reduces inter-core crosstalk by dedicating separate cores for data and power, though not eliminating it entirely. Both systems rely on advanced DSP like MIMO techniques to manage and mitigate these crosstalk effects, ensuring stable transmission quality.

Given that MDM technology is still under research, selecting standard commercial lab equipment for deployment is currently not feasible, as many components remain in research-grade forms. This research-stage limitation contributes to the high initial costs of MDM-based systems. However, it is anticipated that MDM technology will follow a cost reduction trend similar to that of DWDM as it matures, making it more economically accessible over time by reducing the price per Mbps, as per Jefferies.com, the price of the Mbps has been reduced from 6.5 Euro (7USD) in 2004 to less than 0.9 Euro(1USD) in 2015 [45]. This decrement occurred due to the decrement in the DWDM prices, similar trend in expected with MDM systems, allowing to provide economic solution to the extensive needs of capacities which the current DWDM system will not be able to handle in the future.

Figure 4 shows a side-by-side comparison of the scenarios discussed in this study, highlighting the impact of the costliest active and passive parts for each system.

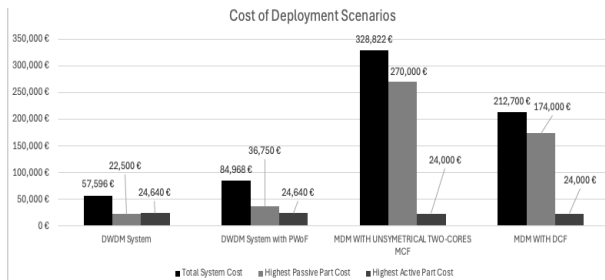


Fig. 4. Cost Comparison Per System

In terms of sustainability, both configurations benefit from the PwOF concept by centralizing power generation at the CO. This centralized setup enables the integration of renewable energy sources, reducing the reliance on distributed power infrastructure and lowering the overall carbon footprint. Adopting centralized, green power generation aligns with modern sustainability goals, providing an eco-friendly solution for powering RRUs across the network.

IX. CONCLUSION

This study assessed the cost-effectiveness and feasibility of deploying MDM with PwOF for bidirectional fronthaul, comparing DCF and MCF configurations. The DCF-based system, with a CapEx of 212,000 Euros, simplifies deployment by carrying both data and power within a single fiber but faces intra-core crosstalk issues. The MCF-based system, costing 322,000 Euros, improves signal isolation by using dedicated cores for data and power transmission, making it more suitable for high-data-rate applications. Although MDM systems are currently expensive, historical trends from DWDM suggest that costs will decline over time, making MDM a viable long-term solution. Moreover, PwOF enhances network sustainability by centralizing power generation, reducing energy consumption and carbon footprint. Moving forward, future research should focus on Improving DSP algorithms to mitigate crosstalk and enhance signal stability, scaling MDM transmission beyond fronthaul applications, and optimizing PwOF efficiency for better power transmission. With its potential for cost savings, scalability, and environmental sustainability, MDM with PwOF offers a promising path for next-generation mobile networks. As technology advances, it could become a cornerstone of future high-capacity optical infrastructures.

I. ACKNOWLEDGEMENT

The work has been supported by the National Research Development and Innovation Office (NKFIH) through the OTKA Grant K 142845.

REFERENCES

- [1] A. Fayad, T. Cinkler, and J. Rak, "5G/6G optical fronthaul modeling: cost and energy consumption assessment," *Journal of Optical Communications and Networking*, vol. 15, no. 9, p. D33, Sep. 2023, doi: 10.1364/JOCN.486547.
- [2] A. Fayad, T. Cinkler, and J. Rak, "Toward 6G Optical Fronthaul: A Survey on Enabling Technologies and Research Perspectives," *IEEE Communications Surveys & Tutorials*, vol. 27, no. 1, pp. 629–666, Feb. 2025, doi: 10.1109/COMST.2024.3408090.
- [3] Y. Chen et al., "Bidirectional mobile fronthaul based on wavelength reused MDM," in *ICOON 2016 - 2016 15th International Conference on Optical Communications and Networks*, Institute of Electrical and Electronics Engineers Inc., Mar. 2017, doi: 10.1109/ICOON.2016.7875835.
- [4] R. Mercy Kingsta and R. Shantha Selvakumari, "A review on coupled and uncoupled multicore fibers for future ultra-high capacity optical communication," *Optik (Stuttg.)*, vol. 199, Dec. 2019, doi: 10.1016/j.jleo.2019.163341.
- [5] Y. Su, Y. He, H. Chen, X. Li, and G. Li, "Perspective on mode-division multiplexing," *Appl Phys Lett*, vol. 118, no. 20, May 2021, doi: 10.1063/5.0046071.
- [6] M. Matsuura, "Power-Over-Fiber Using Double-Clad Fibers," *Journal of Lightwave Technology*, vol. 40, no. 10, pp. 3187–3196, May 2022, doi: 10.1109/JLT.2022.3164566.
- [7] A. Lazaro, R. Villarino, and D. Girbau, "A Survey of NFC Sensors Based on Energy Harvesting for IoT Applications," *Sensors*, vol. 18, no. 11, p. 3746, Nov. 2018, doi: 10.3390/s18113746.
- [8] P. Vamvakas, E. E. Tsiropoulou, M. Vomvas, Papavassiliou, "Adaptive power management in wireless powered communication networks: a user-centric approach," in *2017 IEEE 38th Sarnoff Symposium*, IEEE, Sep. 2017, pp. 1–6, doi: 10.1109/SARNOF.2017.8080386.
- [9] R. Mercy Kingsta and R. Shantha Selvakumari, "A review on coupled and uncoupled multicore fibers for future ultra-high capacity optical communication," *Optik (Stuttg.)*, vol. 199, Dec. 2019, doi: 10.1016/j.jleo.2019.163341.
- [10] E. Alon, V. Stojanovic, J. M. Kahn, S. Boyd, and M. Horowitz, "Equalization of modal dispersion in multimode fiber using spatial light modulators," in *IEEE Global Telecommunications Conference, 2004. GLOBECOM'04.*, IEEE, 2004, pp. 1023–1029, doi: 10.1109/glocom.2004.1378113.
- [11] F. Yaman, N. Bai, B. Zhu, T. Wang, and G. Li, "Long distance transmission in few-mode fibers," *Opt Express*, vol. 18, no. 12, p. 13 250, Jun. 2010, doi: 10.1364/OE.18.013250.
- [12] Y. Li, N. Hua, and X. Zheng, "CapEx advantages of multi-core fiber networks," *Photonic Network Communications*, vol. 31, no. 2, pp. 228–238, Apr. 2016, doi: 10.1007/s11107-015-0536-9.
- [13] A. S. MOHAMED, "Mode Division Multiplexing Zero Forcing Equalization Scheme Using LU Factorization," Thesis (Masters), Universiti Utara Malaysia (UUM), 2016. Accessed: Feb. 14, 2025. [Online]. Available: <https://etd.uum.edu.my/id/eprint/6565>
- [14] M. Matsuura, N. Tajima, H. Nomoto, and D. Kamiyama, "150- W Power-Over-Fiber Using Double-Clad Fibers," *Journal of Lightwave Technology*, vol. 38, no. 2, pp. 401–408, Jan. 2020, doi: 10.1109/JLT.2019.2948777.
- [15] D. S. Montero, J. D. Lopez-Cardona, F. M. A. Al-Zubaidi, I. Perez, P. C. Lallana, and C. Vazquez, "The Role of Power-over-Fiber in C-RAN Fronthauling Towards 5G," in *2020 22nd International Conference on Transparent Optical Networks (ICTON)*, IEEE, Jul. 2020, pp. 1–4, doi: 10.1109/ICTON51198.2020.9203531.
- [16] Nokia, "Nokia Flexi Multiradio 10 Base Station," 2015.
- [17] Ericsson, "Radio Portfolio."
- [18] S. Fafard and D. P. Masson, "74.7% Efficient GaAs-Based Laser Power Converters at 808 nm at 150 K," *Photonics*, vol. 9, no. 8, Aug. 2022, doi: 10.3390/photonics9080579.

Cost Analysis of Mode Division Multiplexing Bidirectional Transmission System

- [19] M. Perales et al., "Characterization of high performance silicon-based VMJ PV cells for laser power transmission applications," in *High-Power Diode Laser Technology and Applications XIV*, SPIE, Mar. 2016, p. 97 330U. **doi:** 10.1117/12.2213886.
- [20] LUMENTUM, "Photonic Power Module (21135472)."
- [21] Aerodiode, "Data Sheet of 808NM MULTI-MODE LASER DIODE."
- [22] R. Singh, M. Kumar, and D. Sharma, "A Review on Radio over Fiber communication System An overview of Wireless Sensor Networks View project An Overview of various Digital Halftone Processing Technique View project," 2017. [Online]. Available: <https://www.researchgate.net/publication/319738806>
- [23] M. Grabner, K. Nithyanandan, P. Peterka, P. Koska, A. A. Jasim, and P. Honzatko, "Simulations of Pump Absorption in Tandem-Pumped Octagon Double-Clad Fibers," *IEEE Photonics J*, vol. 13, no. 2, Apr. 2021, **doi:** 10.1109/JPHOT.2021.3060857.
- [24] B. B.S and S. Azeem, "A survey on increasing the capacity of 5G Fronthaul systems using RoF," *Optical Fiber Technology*, vol. 74, Dec. 2022, **doi:** 10.1016/j.yofte.2022.103078.
- [25] E. Udvary, "Integration of QKD Channels to Classical High-speed Optical Communication Networks," *Infocommunications Journal*, vol. 15, no. 4, pp. 2–9, 2023, **doi:** 10.36244/ICJ.2023.4.1.
- [26] E. Alon, V. Stojanović, J. M. Kahn, S. Boyd, and M. Horowitz, "Equalization of Modal Dispersion in Multimode Fiber using Spatial Light Modulators," 2004. **doi:** 10.1109/glocom.2004.1378113
- [27] Y. Su, Y. He, H. Chen, X. Li, and G. Li, "Perspective on mode-division multiplexing," May 17, 2021, *American Institute of Physics Inc.* **doi:** 10.1063/5.0046071.
- [28] A. S. Mohamed and E. Udvary, "Enabling Enhanced In- Building Solutions Fronthaul Connectivity: The Role of Mode Division Multiple Access in Mobile Networks," in *2024 24th International Conference on Transparent Optical Networks (ICTON)*, IEEE, Jul. 2024, pp. 1–5. **doi:** 10.1109/ICTON62926.2024.10647457.
- [29] C. Rottondi, P. Martelli, P. Boffi, L. Barletta, and M. Tornatore, "Crosstalk-Aware Core and Spectrum Assignment in a Multicore Optical Link With Flexible Grid," *IEEE Transactions on Communications*, vol. 67, no. 3, pp. 2144–2156, Mar. 2019, **doi:** 10.1109/TCOMM.2018.2881697.
- [30] ThorLabs, "ThorLabs DCF Cable Data Sheet".
- [31] F. A. Shnain and A. R. Salih, "Design and Study of Few-Mode Fibers at 1550 nm," *Journal of Educational and Scientific Studies*, vol. 18, no. 1, pp. 83–95, 2021.
- [32] A. I. Zreikat and S. Mathew, "Performance Evaluation and Analysis of Urban-Suburban 5G Cellular Networks," *Computers*, vol. 13, no. 4, p. 108, Apr. 2024, **doi:** 10.3390/computers13040108.
- [33] FS.com, "Cisco C19 DWDM-SFP10G-62.23 Compatible SFP+10G DWDM 1562.23nm 100GHz 40km DOM Duplex LC/UPC SMF Optical Transceiver Module for Transmission," FS.com. Accessed: Oct. 21, 2024. [Online]. Available: [Cisco C19 DWDM-SFP10G-62.23 Compatible SFP+10G DWDM 1562.23nm 100GHz 40km DOM Duplex LC/UPC SMF Optical Transceiver Module for Transmission](https://www.fs.com/de-en/products/182058.html?attribute=70022&id=3461426)
- [34] FS.com, "40 Channels 100GHz C21-C60 Active, with 1310nm and Monitor Port, 4.5dB Typical IL, LC/UPC, Dual Fiber DWDM Mux Demux, 1U Rack Mount," FS.com. Accessed: Oct. 21, 2024. [Online]. Available: <https://www.fs.com/de-en/products/182058.html?attribute=70022&id=3461426>
- [35] FS.com, "Customized In-Line EDFA for DWDM Solution," FS.com. Accessed: Oct. 21, 2024. [Online]. Available: <https://www.fs.com/de-en/products/35924.html>
- [36] ThorLabs, "SMF-28-J9-SpecSheet."
- [37] Eurostat, "Electricity Price Statistics," Eurostat. Accessed: Oct. 27, 2024. [Online]. Available: https://ec.europa.eu/eurostat/statistics-explained/index.php?title=Electricity_price_statistics
- [38] ThorLabs, "Typical Multimode Fiber Attenuation 0.39 NA, High-OH, Hard Clad Silica, TECS-Clad Attenuation (dB/km)," 2024. [Online]. Available: www.thorlabs.com/contact
- [39] MHGOPOWER.COM, "YCH Series MIH® Photovoltaic Power Converter," http://www.mhgopower.com/laser_pof_YCHPPC.html.
- [40] ThorLabs, "Thorlabs Spatial Laser," https://www.thorlabs.com/newgrouppage9.cfm?objectgroup_id=10378. Accessed: Oct. 28, 2024. [Online]. Available: https://www.thorlabs.com/newgrouppage9.cfm?objectgroup_id=10378
- [41] ThorLabs.com, "MDM Photo Detector Cost," https://www.thorlabs.com/newgrouppage9.cfm?objectgroup_id=1297. Accessed: Oct. 28, 2024. [Online]. Available: https://www.thorlabs.com/newgrouppage9.cfm?objectgroup_id=1297
- [42] FiberCore, "FAN OUTS," 2020. Accessed: Nov. 15, 2024. [Online]. Available: <https://fibercore.humaneticsgroup.com/products/multicore-fiber/fan-outs/fan-4c>
- [43] Cailabs, "Cailabs MDM MUX." Accessed: Nov. 06, 2024. [Online]. Available: <https://www.cailabs.com/fiber-networks/optical-networks-of-the-future/proteus/>
- [44] Cailabs, "PROTEUS-S6-SI."
- [45] R. Wu, * Jefferies, and H. Kong, "China (PRC) I Technology Technology Optical Transceiver: How It Differs in 5G and Cloud EQUITY RESEARCH CHINA," 2017.



Ahmed S. Mohamed received his B.Sc. from the university of Technology/ Iraq in 2007 and M.Sc. from the Universiti Utara Malaysia (UUM) in 2016. He is currently pursuing his PhD in Budapest University of Technology and Economics/Hungary in optical networks for mobile networks applications.



Eszter Udvary received her Ph.D. degree in electrical engineering from the Budapest University of Technology and Economics (BME), Budapest, Hungary, in 2009. She is currently an Associate Professor at BME, Mobile Communication and Quantum Technologies Laboratory. Dr. Udvary's research interests are in the broad area of optical communications, including microwave photonics, optical access network, visible light communication, and quantum communication.

Improved Outage Probability using Multi-IRS-Assisted MIMO Wireless Communications in Cellular Blockage Scenarios

Suresh Penchala, Shravan Kumar Bandari, and V.V. Mani

Abstract—Future wireless communication technologies must be upgraded to serve the upcoming seamless data-intensive applications and support a minimum information rate. To this end, an intelligent reflecting surface (IRS) technology was recently proposed as a viable solution to cover the uncovered regions with enhanced performance by means of a controlled wireless environment using multiple reflecting elements. In this paper, to elevate the end-user experience, we intend to improve the outage probability (OP) performance utilizing multiple IRS for multiple-input-multiple-output (MIMO) communication systems under a generalized η - μ fading channel, which is suitable for non-line-of-sight (NLOS) scenarios. We derived a closed-form expression for OP and validated the same using a rigorous Monte-Carlo (MC) simulation setup under the considered system model. A comprehensive analysis of each system parameter impacting the likelihood that the communication channel supports the information rate has been detailed based on the number of IRSs, the number of reflecting elements, antenna count at transmitter and receiver, fading parameters, and the placement of IRSs. Results suggest that employing multiple IRSs reduces the outage scenarios in blockage zones, and further improvement can be observed with multiple antennas positioning the IRSs closer to either the transmitter or the receiver. Furthermore, we present an energy efficiency (EE) assessment for the multi-IRS system.

Index Terms—Intelligent reflecting surface, outage probability, multi-IRS, multi-antenna, η - μ fading channel.

I. INTRODUCTION

The intelligent reflective surface (IRS) (*a.k.a.* reconfigurable intelligent surface) is emerging as a transformative technology with significant potential to influence the development of next-generation wireless communication systems, including 6G cellular networks [1]. In a nutshell, an IRS panel comprises multiple low-cost passive meta-material surface elements steered by a controller to guide the incident electromagnetic (EM) signals toward the intended direction by adjusting their reflecting angles to enhance the end-user experience by maximizing the received signal-to-noise ratio (SNR) [2]. The configuration thus established will now be able to control the uncontrollable wireless radio environment with ease, which has attracted significant attention from academia and industry [3], [4]. IRS technology tunes the beam (active/passive/jointly) to direct

electromagnetic waves toward a specified target [5]. This functionality enables IRS-assisted systems to mitigate deep-fading conditions by dynamically reconfiguring the wireless environment.

IRS technology is adaptable and suitable for various applications, covering indoor and outdoor wireless systems, 6G mobile networks, and satellite communications [6]. Various technological advancements, including multiple antenna systems [7], generalized frequency division multiplexing (GFDM) [8], [9], orthogonal time frequency space modulation (OTFS) [10], and network optimization [11], have been developed to enhance the performance of wireless communication. The utilization of the non-orthogonal multiple access (NOMA) scheme in the IRS application was examined in [12]. The authors in [13] employed an IRS to improve the performance of a dual-hop system that combines free-space optical and radio frequency (FSO-RF) technologies. Subsequently, the authors in [14] expanded their investigation on the IRS to include unmanned aerial vehicle (UAV) networks. Moreover, the concept of employing IRS-assisted transmission for efficient communication has been thoroughly examined in relation to co-operative relay [15], space shift keying [16], and backscatter communication technologies [17]. Recent studies have explored practical IRS deployment aspects such as indoor panel positioning and field trials for single and multi reflection architectures [18], as well as AI-driven control approaches like deep reinforcement learning applied to UAV-mounted IRS [19].

On the other hand, wireless communication channels are subject to fading due to interactions like reflection, scattering, and diffraction as signals propagate. These effects cause fluctuations in signal strength and quality, which traditional models such as Rayleigh and Nakagami- m distributions only partially capture [20]. To address the limitations of these models, the η - μ fading channel was introduced as a more versatile framework [21]. It can describe a wide range of fading scenarios by accounting for diverse scattering environments and providing a better fit to experimental data. It is customary to analyze the performance of any new technology in an uncontrollable, random wireless fading environment. The deployment of IRS technology allows for improved control over channel conditions by intelligently positioning the IRS panel at desired locations. This helps mitigate channel randomness and enhances overall system performance in terms of bit/symbol error rates (BER/SER), Ergodic capacity, outage probability (OP), and energy efficiency (EE).

Suresh Penchala and Shravan Kumar Bandari are with the Department of Electronics and Communication Engineering, National Institute of Technology Meghalaya, India (e-mail: penchalasuresh@gmail.com, shravnbandari@nitm.ac.in).

V.V. Mani is with the Department of Electronics and Communication Engineering, National Institute of Technology Warangal, India (e-mail: vvmani@nitw.ac.in).

DOI: 10.36244/ICJ.2025.3.5

Improved Outage Probability using Multi-IRS-Assisted MIMO Wireless Communications in Cellular Blockage Scenarios

In [22], an IRS-assisted network in Nakagami- m fading was proposed, and closed-form expressions for the outage probability, average symbol error probability, and channel capacity were derived. Results suggest the superior performance of the IRS-aided communication system compared to the conventional one. Study in [23] examines the IRS technology in enhancing wireless communication under a Weibull fading channel. It highlights the improvement in spectral efficiency, outage probability, and SER through increased IRS elements and better phase accuracy. A dual IRS-assisted system with separate IRS placements near the source and destination, enabling communication without a direct link, was investigated in [24]. Closed-form OP and SER expressions are derived for Rician and Rayleigh fading, with results showing superior performance and energy efficiency compared to a single IRS setup. The OP of a single-input-single-output (SISO) wireless system using multiple IRS panels in κ - μ fading channels was investigated in [25]. Selecting the best IRS panel ensures optimal quality of service (QoS) for single-node communication. [26] investigates the OP and asymptotic sum rate in multi-IRS-assisted networks operating over Rayleigh fading channels. It was observed that the number of IRSs and the number of reflecting elements play an important role in the capacity scaling law of multiple RIS-aided networks. A system with multiple IRSs has been studied in [27] and analyzed under Rician fading conditions, focusing on the OP, which depends on the phase shifts of all IRSs. Studies demonstrate that utilizing multiple IRSs significantly enhances wireless communication in terms of coverage and service quality. The collaborative utilization of multiple IRSs significantly enhances the received signal power [28]. Furthermore, situating IRSs in diverse locations guarantees that signals can still access the receiver via alternative routes, even if certain paths encounter significant fading.

From these observations, it is crucial to highlight that the outage probability analysis of a multi-IRS-enabled MIMO communication system under η - μ fading channels under the NLOS scenario between the transmitter (Tx) and the receiver (Rx) remains unexplored. Additionally, we adopt a more practical approach by modeling the fading behavior with an η - μ channel, which more accurately represents real-world conditions. To address this gap, in this work, we intend to investigate the outage probability performance analysis of a multi-IRS-assisted communication system in a MIMO wireless fading channel environment. The key contributions of the proposed system model are summarized as follows:

- A versatile channel model for evaluating various IRS-assisted systems is highly significant. In this study, we adopt such a model for the Tx-IRSs and IRSs-Rx links, employing a generalized η - μ fading distribution combined with large-scale path loss, tailored for a multi-antenna MIMO communication system supported by multiple IRSs.
- Using the derived signal-to-noise ratio and the central limit theorem, we determined the mean and variance of the overall cascaded wireless channel links.

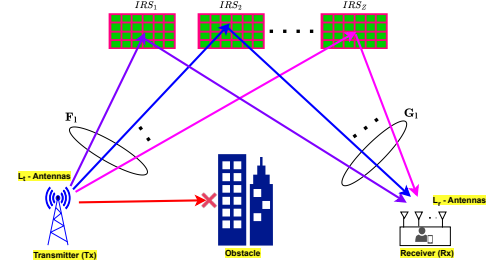


Fig. 1. Multi IRS aided MIMO wireless communication model

- A framework to find the probability that a user is in an outage scenario is also discussed through the obtained theoretical expression. For completeness, the energy efficiency comparing the multi-IRS system to the reference single-IRS system is also discussed.
- We demonstrated the optimal placement of the IRS, emphasizing its critical positioning near both the Tx and Rx for a fixed transmitted signal power.
- The theoretical expressions closely match the Monte Carlo (MC) simulations, validating the proposed multi-antenna, multi-IRS MIMO system in the η - μ channel. We also analyzed and highlighted the impact of key parameters on performance.

The structure of the paper is organized as follows: Section II describes the system model, Section III outlines the theoretical formulations for analyzing OP performance, Section IV outlines the energy efficiency, and Section V examines the results obtained through Monte Carlo simulations and analytical calculations. Finally, Section VI provides a discussion of the conclusions.

II. SYSTEM MODEL

This paper explores the design of a MIMO wireless system with L_t antennas at the transmitter (Tx), L_r antennas at the receiver (Rx), and Z IRSs as shown in Fig. 1. The system operates under η - μ fading channels, and we derive a mathematical expression for the OP in the NLOS scenario. Notably, we consider the direct link between the Tx and Rx to be completely blocked due to tall buildings and is therefore excluded from the analysis. Combining the effects of all Z IRSs, the signal received at the Rx via Z independent paths can be expressed as [29],

$$y = \left[\sum_{l=1}^Z \sum_{i=1}^{N_l} \sum_{p=1}^{L_t} \sum_{q=1}^{L_r} g_{lip} r_{liq} f_{liq} \right] x + n \quad (1)$$

The parameters of equation (1) along with their corresponding notations are listed in Table I. $g_{lip} = d_{Tl}^{-\frac{\alpha}{2}} \delta_{lip} e^{-j\phi_{lip}}$ is the channel between the p -th Tx and the l -th IRS, where δ_{lip} is the channel gain between the p -th Tx antenna and the i -th element of the l -th IRS, d_{Tl} is the distance between the Tx and the l -th IRS, and ϕ_{lip} is the phase of the Tx to the i -th element of the l -th IRS channel. $f_{liq} = d_{lR}^{-\frac{\alpha}{2}} \vartheta_{liq} e^{-j\psi_{liq}}$ is the channel between the l -th IRS and the q -th Rx, where ϑ_{liq} is the channel gain

between the i -th IRS element and the q -th Rx antenna, d_{lIR} is the distance between the l -th IRS and the Rx, and ψ_{liq} is the phase of the IRS-to-Rx channel. $r_{li} = b_{li}e^{j\theta_{li}}$ is the reflection coefficient of the i -th element of the l -th IRS, where θ_{li} is the phase shift applied by the i -th element of the l -th IRS. Additionally, α represents the average path loss exponent. For simplicity, this study assumes the optimal phase shift (OPS) selection as $\theta_{li} = \phi_{lip} + \psi_{liq}$ [30]. Assuming the optimal

TABLE I
PARAMETERS NOTATION SUMMARY

Symbol	Description
L_t	Number of transmit antennas
L_r	Number of receive antennas
Z	Number of IRS panels
N_l	Number of reflecting elements per IRS
$\mathbf{x}$	Transmitted signal from Tx
$\mathbf{y}$	Received signal at Rx
$\mathbf{n}$	Additive white gaussian noise (AWGN)
g_{lip}	The channel between the p -th Tx and the l -th IRS
f_{liq}	The channel between the l -th IRS and the q -th Rx
r_{li}	IRS reflection coefficient
θ_{li}	Phase shift applied by i th element of l th IRS
d_{TlI}, d_{lIR}	Tx-to-IRS and IRS-to-Rx distances
α	Path loss exponent
δ_{lip}	Fading gain from Tx to IRS element
ϑ_{liq}	Fading gain from IRS element to Rx
Υ	Effective cascaded channel gain
ζ	Normalization term for SNR
γ	Instantaneous received SNR
γ_{th}	Outage SNR threshold
Λ_e, Λ_v^2	Mean and variance of Υ
ω_m, ω_v	Mean and variance of η - μ variables
η	Scattering parameter (η - μ model)
μ	Clustering parameter (η - μ model)
$Q(\cdot)$	Gaussian Q-function
$Q_{1/2}(\cdot, \cdot)$	Generalized Marcum Q-function (order 1/2)

phase and perfect reflection ($b_{li} = 1$) at the IRS. In practical implementations, discrete phase shifting (DPS) is employed to restrict phase values to a finite set within the range $[0, 2\pi]$, where the phase shift is given by $\theta_{lip} = \frac{2\pi}{2^b}f$. Here f can be $0, 1, \dots, 2^b - 1$, and b indicates how many digits are used to show each level [31]. As a result, the maximum received SNR at the Rx can be expressed as [32],

$$\gamma = \frac{\left(\sum_{l=1}^Z \sum_{i=1}^{N_l} \sum_{p=1}^{L_t} \sum_{q=1}^{L_r} d_{TlI}^{-\frac{\alpha}{2}} \delta_{lip} d_{lIR}^{-\frac{\alpha}{2}} \vartheta_{liq} \right)^2 E_x}{N_o} = \frac{\Upsilon^2 E_x}{\left(\sum_{l=1}^Z d_{TlI}^{\alpha} d_{lIR}^{\alpha} \right) N_o} = \Upsilon^2 \zeta \quad (2)$$

where $\Upsilon = \left(\sum_{l=1}^Z \sum_{i=1}^{N_l} \sum_{p=1}^{L_t} \sum_{q=1}^{L_r} \delta_{lip} \vartheta_{liq} \right)$, $\zeta = \frac{E_x}{\left(\sum_{l=1}^Z d_{TlI}^{\alpha} d_{lIR}^{\alpha} \right) N_o}$. When N_l is sufficiently large, the central limit theorem (CLT) suggests that Υ is likely to follow a Gaussian distribution [33]. As a result, γ will follow a noncentral chi-squared (NCCS) distribution with one degree of freedom. Furthermore, δ_{lip} and ϑ_{liq} are assumed to be independent and identically distributed (IID) $\eta - \mu$ random variables, characterized by the following probability density function (PDF) [9], [34],

TABLE II
 η - μ SPECIAL CASES

Description	η	μ
Rayleigh	1	0.5
Nakagami- m	1	$m/2$
Nakagami- q (Hoyt)	q^2	0.5

$$f_P(\rho) = \frac{2\sqrt{\pi}(1+\eta)^{\mu+\frac{1}{2}}\rho^{2\mu}}{\sqrt{\eta}(1-\eta)^{\mu-\frac{1}{2}}\Gamma(\mu)} \exp\left[-\frac{\mu(1+\eta)^2\rho^2}{2\eta}\right] I_{\mu-\frac{1}{2}}\left[\frac{\mu(1-\eta^2)\rho^2}{2\eta}\right] \quad (3)$$

where $\Gamma(\cdot)$ refers to the Gamma function, $I_n(\cdot)$ signifies the n -th modified Bessel function of the first kind, η represents the ratio of the non-centrality parameter to the scale parameter, primarily influencing the spread or variability of the fading envelope. The parameter μ represents the shape parameter of the distribution, affecting the asymmetry of the fading envelope. The generalized η - μ fading distribution is a versatile model capable of encompassing several other distributions, including Rayleigh, Nakagami- q , and Nakagami- m fading channels, as special cases shown in Table II. It is especially suitable for NLOS environments due to its capability to model the intricate characteristics of signal propagation, including non-uniform conditions influenced by scattering elements, reflective surfaces, and diffraction effects. The j^{th} moment of $\eta - \mu$ distribution is given by [34],

$$E(P^j) = \frac{2^{(2\mu+j/2)}\Gamma(2\mu+j/2)}{(2+\eta^{-1}+\eta)^{\mu+j/2}\mu^{j/2}\Gamma(2\mu)} {}_2F_1\left[\mu+\frac{j}{4}+\frac{1}{2}, \mu+\frac{j}{4}; \mu+\frac{1}{2}; \left(\frac{1-\eta}{1+\eta}\right)^2\right] \quad (4)$$

where ${}_2F_1(\cdot)$ is the Gauss hypergeometric function. Consequently, by utilizing equation (4), we can determine the mean and variance (at the top of the next page) as follows,

$$\varpi_m \triangleq E(P) = \frac{2^{(2\mu+1/2)}\Gamma(2\mu+1/2)}{(2+\eta^{-1}+\eta)^{\mu+1/2}\mu^{1/2}\Gamma(2\mu)} {}_2F_1\left[\mu+\frac{3}{4}, \mu+\frac{1}{4}; \mu+\frac{1}{2}; \left(\frac{1-\eta}{1+\eta}\right)^2\right] \quad (5)$$

where $E(\cdot)$ and $\text{Var}(\cdot)$ are the expectation and variance operators respectively.

III. OUTAGE PROBABILITY ANALYSIS

In this section, we derive the analytical expression for the OP within the proposed generalized framework. The statistical properties of Υ are necessary to derive the analytical expression of the performance metric under consideration. With this, next we aim to find the mean and variance of Υ . Recalling δ_{kip} and ϑ_{kig} follow IID η - μ distribution functions, we have the following,

$$E[\delta_{lip}\vartheta_{liq}] = E[\delta_{lip}]E[\vartheta_{liq}] = \varpi_m^2 \quad (7)$$

Improved Outage Probability using Multi-IRS-Assisted MIMO
 Wireless Communications in Cellular Blockage Scenarios

$$\varpi_v \triangleq \text{Var}(P) = \left(\frac{4}{2+\eta^{-1}+\eta} \right)^{\mu+1} {}_2F_1 \left[\mu+1, \mu+\frac{1}{2}; \mu+\frac{1}{2}; \left(\frac{1-\eta}{1+\eta} \right)^2 \right] - \left[\frac{2^{(2\mu+1/2)} \Gamma(2\mu+1/2)}{(2+\eta^{-1}+\eta)^{\mu+1/2} \mu^{1/2} \Gamma(2\mu)} {}_2F_1 \left[\mu+\frac{3}{4}, \mu+\frac{1}{4}; \mu+\frac{1}{2}; \left(\frac{1-\eta}{1+\eta} \right)^2 \right] \right]^2 \quad (6)$$

$$\begin{aligned} \text{Var}[\delta_{lip}\vartheta_{liq}] &= \text{Var}[\delta_{lip}] \text{Var}[\vartheta_{liq}] + \text{Var}[\delta_{lip}] (\text{E}[\vartheta_{liq}])^2 \\ &\quad + \text{Var}[\vartheta_{liq}] (\text{E}[\delta_{lip}])^2 \\ &= \varpi_v^2 + 2\varpi_v\varpi_m^2 \end{aligned} \quad (8)$$

Consequently, the statistics of Υ can be derived as,

$$\Lambda_e \triangleq \text{E}[\Upsilon] = \text{E} \left[\sum_{l=1}^Z \sum_{i=1}^{N_l} \sum_{p=1}^{L_t} \sum_{q=1}^{L_r} \delta_{lip}\vartheta_{liq} \right] = ZN_l L_t L_r \varpi_m^2 \quad (9)$$

$$\begin{aligned} \Lambda_v^2 \triangleq \text{Var}[\Upsilon] &= \text{Var} \left[\sum_{l=1}^Z \sum_{i=1}^{N_l} \sum_{p=1}^{L_t} \sum_{q=1}^{L_r} \delta_{lip}\vartheta_{liq} \right] \\ &= ZN_k L_t L_r (\varpi_v^2 + 2\varpi_v\varpi_m^2) \end{aligned} \quad (10)$$

Recalling $\gamma = \Upsilon^2 \zeta$ and Υ^2 is a NCCS distribution with one degree of freedom having parameters Λ_e and Λ_v^2 , the PDF of Υ^2 can be given as [33],

$$f_{\Upsilon^2}(s) = \frac{1}{2\Lambda_v^2} \left(\frac{s}{\Lambda_e} \right)^{-\frac{1}{2}} e^{-\frac{(s+\Lambda_e)}{2\Lambda_v^2}} I_{-\frac{1}{2}} \left[\frac{\sqrt{s\Lambda_e}}{\Lambda_v^2} \right] \quad (11)$$

Accordingly, the PDF of γ can be written as,

$$f_{\gamma}(w) = \frac{1}{2\Lambda_v^2 \zeta} \left(\frac{w}{\zeta \Lambda_e} \right)^{-\frac{1}{2}} e^{-\frac{(w+\zeta \Lambda_e)}{2\zeta \Lambda_v^2}} I_{-\frac{1}{2}} \left[\frac{1}{\Lambda_v^2} \sqrt{\frac{w\Lambda_e}{\zeta}} \right] \quad (12)$$

In order to find out the dead zones caused by blockages, one has to find out the outage probability. To this end, the probability that the received SNR falls below a specified threshold value (γ_{th}) at the Rx can be expressed as [26], [35],

$$\begin{aligned} P_{\text{outage}} &= \Pr(\gamma < \gamma_{th}) = \int_{-\infty}^{\gamma_{th}} f_{\gamma}(w) dw \\ &= 1 - \int_{\gamma_{th}}^{\infty} f_{\gamma}(w) dw \end{aligned} \quad (13)$$

$$\int_{\gamma_{th}}^{\infty} f_{\gamma}(w) dw = \int_{w=\gamma_{th}}^{\infty} \frac{1}{2\Lambda_v^2 \zeta} \left(\frac{w}{\zeta \Lambda_e} \right)^{-\frac{1}{2}} e^{-\frac{(w+\zeta \Lambda_e)}{2\zeta \Lambda_v^2}} I_{-\frac{1}{2}} \left[\frac{1}{\Lambda_v^2} \sqrt{\frac{w\Lambda_e}{\zeta}} \right] dw$$

Substituting $\frac{w}{\zeta \Lambda_v^2} = \tau^2$, the above equation can be modified as,

$$\int_{\gamma_{th}}^{\infty} f_{\gamma}(w) dw = \frac{1}{\left(\frac{\Lambda_e}{\Lambda_v^2} \right)^{-\frac{1}{2}}} \int_{\sqrt{\frac{\gamma_{th}}{\zeta \Lambda_v^2}}}^{\infty} \tau^{\frac{1}{2}} e^{-\left(\frac{\tau^2 + \frac{\Lambda_e}{\Lambda_v^2}}{2} \right)} I_{m-1} \left[\sqrt{\frac{\Lambda_e}{\Lambda_v^2}} \tau \right] d\tau \quad (14)$$

The Marcum-Q function can be given as [35],

$$Q_m(a, b) = \frac{1}{a^{m-1}} \int_b^{\infty} \tau^m e^{-\left(\frac{\tau^2 + a^2}{2} \right)} I_{m-1}(a\tau) d\tau \quad (15)$$

Now comparing equations (14) and (15), we have the following,

$$\int_{\gamma_{th}}^{\infty} f_{\gamma}(w) dw = Q_{\frac{1}{2}} \left(\sqrt{\frac{\Lambda_e}{\Lambda_v^2}}, \sqrt{\frac{\gamma_{th}}{\zeta \Lambda_v^2}} \right) \quad (16)$$

TABLE III
SIMULATION PARAMETERS [28], [37]

Description	Values
Number of reflecting elements, N	36 – 576
Distance between Tx to Rx, d	100 m
Height of IRS (h_I)	15 m
Height of Tx (h_T)	10 m
Height of Rx (h_R)	2 m
Path loss Exponent, α	3
Outage threshold, γ_{th}	10 dB
Carrier frequency (f_c)	3 GHz
Transmit power (P_x)	[0, 30] dBm
Power dissipated per IRS element (P_i) [mW]	7.8
Power conversion efficiency (ξ)	80%
Circuit dissipated power at Tx (P_c^{Tx}) [dBm]	10
Circuit dissipated power at Rx (P_c^{Rx}) [dBm]	10
Hardware static power of phase shift circuit (P_{ps}) [dBm]	10
Bandwidth	10 MHz
Noise Figure	10 dBm
Modulation scheme	M-PSK

Finally, using equation (16) in equation (13), the outage probability of the considered system model can be derived as follows,

$$P_{\text{outage}} = 1 - Q_{\frac{1}{2}} \left(\sqrt{\frac{\Lambda_e}{\Lambda_v^2}}, \sqrt{\frac{\gamma_{th}}{\zeta \Lambda_v^2}} \right) \quad (17)$$

where $Q_{\frac{1}{2}}(A, B)$ is the fractional order 1/2 generalized Marcum Q -function and may be written in terms of Gaussian- Q function as [36],

$$Q_{\frac{1}{2}}(A, B) = Q(B - A) + Q(B + A) \quad (18)$$

Consequently, the OP can be restructured as shown here,

$$P_{\text{outage}} = Q \left(\sqrt{\frac{\gamma_{th}}{\zeta \Lambda_v^2}} - \sqrt{\frac{\Lambda_e}{\Lambda_v^2}} \right) + Q \left(\sqrt{\frac{\gamma_{th}}{\zeta \Lambda_v^2}} + \sqrt{\frac{\Lambda_e}{\Lambda_v^2}} \right) \quad (19)$$

IV. ENERGY EFFICIENCY

The EE of an IRS-assisted system can be described as [28],

$$\text{EE} = \text{BW} \times \left(\frac{R}{P_{\text{total}}} \right) \quad (20)$$

where BW is the bandwidth of the system, $R = b \times \text{BW}$ represents the rate at which bits are transmitted, where b signifies the bits per symbol (for BPSK, $b = 1$). P_{total} denotes the comprehensive power utilized by the system to attain a specified BER, which can be calculated as,

$$P_{\text{total}} = P_x + P_x^{HPA} + P_c^{Tx} + NP_i + NP_{ps} + P_c^{Rx} \quad (21)$$

where P_x is power used for transmitting the information, $P_x^{HPA} = \frac{P_x}{\xi}$ is the power consumed by the high power amplifier (HPA) with ξ being the power conversion efficiency (80%), P_i is power at the i^{th} IRS element, P_{ps} is the hardware static power of phase-shift circuit, P_c^{Tx} is circuit power at the transmitter and P_c^{Rx} is circuit power at the receiver.

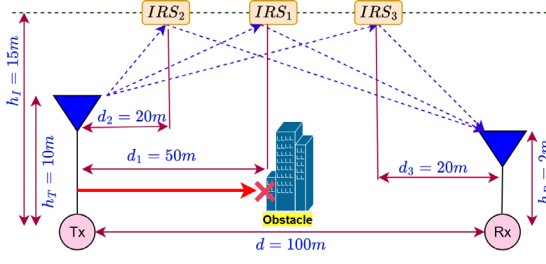
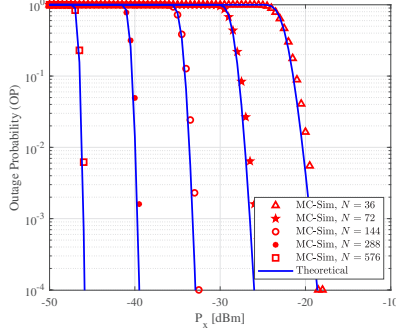


Fig. 2. MC simulation setup.


 Fig. 3. Outage probability versus source transmit power (P_x) with $\eta = 1$, $\mu = 0.5$, $L_t = L_r = 2$, and S-IRS, varying N .

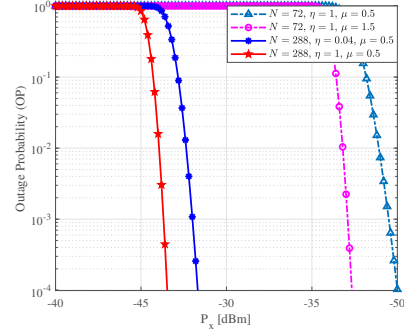
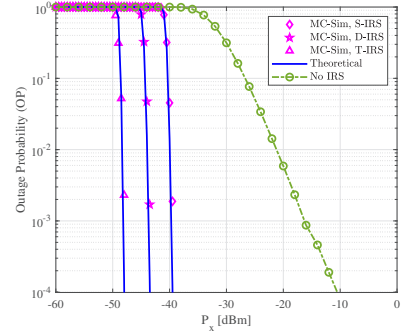
V. RESULTS AND DISCUSSIONS

This section presents the outcomes of the derived OP for the system model under consideration for Single IRS (S-IRS), Double IRS (D-IRS), and Triple IRS (T-IRS). The simulation of MC is established based on the parameters outlined in Table III, aimed at both validating the accuracy of the derived analytical expressions and gaining further insights into the variables influencing overall system performance. The comparison among various IRSs is performed by maintaining an equal number of reflective elements used in S-IRS, D-IRS, and T-IRS systems, specifically $N_1 = 2N_2 = 3N_3 = N$. The simulation configuration is established in accordance with Fig. 2.

A. OP Results

Fig. 3 illustrates the impact of N on the OP. The strong agreement between the simulated and theoretical results once again confirms the accuracy of the obtained formulas in Section III and the OP decreases as N increases. This improved OP performance can be traded with the transmit power. For instance, if an OP of 10^{-4} is deemed appropriate for a specific wireless system, subsequently utilizing $N = 576$ will result in achieving this at approximately $P_x = -46$ dBm, and with $N = 36$, it can be achieved at $P_x = -18.5$ dBm. Therefore, by increasing the number of reflecting elements, it is possible to achieve a gain of 27.5 dBm. Furthermore, as the value of P_x increases, the OP decreases significantly for a fixed N .

Fig. 4 shows the effect of η , μ on the OP performance. The accuracy of the derived expressions is once again confirmed. Based on Fig. 4, it can be deduced that the increases in the η , μ values lead to a decrease in the OP. For instance, when ensuring a constant OP 10^{-2} , an increase in μ from 0.5 to


 Fig. 4. Outage probability versus source transmit power (P_x) with $N = 72, 288$, $L_t = L_r = 2$, and D-IRS, varying η , μ

 Fig. 5. Outage probability versus source transmit power (P_x) for S-IRS, D-IRS, T-IRS, and No IRS cases with $N = 288$, $\eta = 1$, $\mu = 0.5$, and $L_t = L_r = 2$.

1.5 leads to a transmit power improvement of 2.1 dBm, and also, an increase in η from 0.04 to 1 leads to a transmit power improvement of 1.5 dBm.

Fig. 5, illustrates the OP performance for S-IRS, D-IRS, and T-IRS systems, along with a baseline comparison to a conventional MIMO system without IRS (denoted as “No IRS”). The graph demonstrates that the theoretical results derived in Section III are in close agreement with the outcomes obtained through MC simulations. It is evident that an increase in P_x leads to a significant reduction in OP. Furthermore, the OP decreases as the number of IRSs increases. One can note that an increase in P_x results in a noticeable decline in OP. In addition, as the number of IRSs increases the OP decreases. For example, consider an outage probability of 10^{-4} , the transmit power P_x approximately -39 dBm, -43 dBm, and -46 dBm for S-IRS, D-IRS, and T-IRS, respectively. In contrast, the conventional “No IRS” system requires significantly higher transmit power approximately -10 dBm to achieve the same OP, clearly illustrating the efficiency and effectiveness of IRS-assisted communication. These results affirm that multi-IRS deployment not only improves reliability but also enables substantial transmit power savings compared to traditional non-IRS MIMO systems.

The influence of different values of L_t , L_r on the OP is depicted in Fig. 6. The OP results obtained from theoretical analysis for the multi-antenna systems for S-IRS are in complete agreement with the simulated results. Based on the information provided in Fig. 6, it is clear that the outage performance improves significantly as the number of transmitting

Improved Outage Probability using Multi-IRS-Assisted MIMO Wireless Communications in Cellular Blockage Scenarios

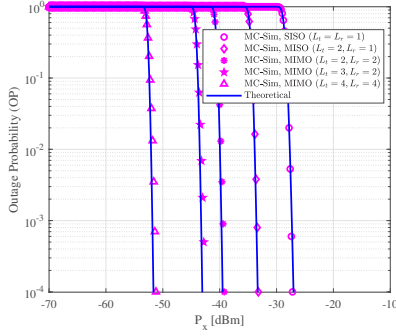


Fig. 6. Outage probability versus source transmit power (P_x) with $N = 288$, $\eta = 1$, $\mu = 0.5$, S-IRS, varying L_t, L_r .

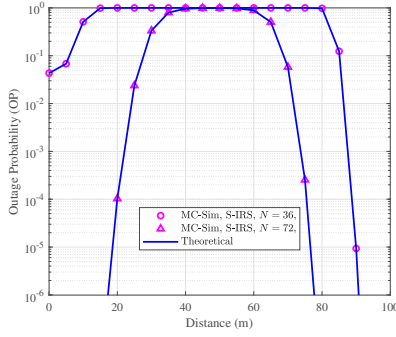


Fig. 7. Outage probability versus distance with $N = 36, 72$, $\eta = 1$, $\mu = 0.5$, S-IRS, $L_t = L_r = 2$, $P_x = 30$ dBm.

and receiving antennas increases in the S-IRS scenario. For example, consider an outage probability of 10^{-4} , the number of Tx and Rx antennas (L_t, L_r) values increases from (1, 1) to (4, 4), the equivalent P_x values changes from -27.2 dBm to -51.2 dBm.

The impact of the user location on the outage probability is depicted in Fig. 7. Firstly, at a fixed P_x , the outage probability is less when the IRS is placed closer to either Tx or Rx. Secondly, as the number of reflecting elements changes from 36 to 72, the outage values decrease drastically. For instance, with $N = 36$, an OP value of 0.0681 is attained when IRS is positioned 5m away from Tx, and with $N = 72$ an OP value of 0.0238 is attained when IRS is 25m away from Tx and also 0.000254 is attained when IRS is 25m away from Rx.

Fig. 8 shows the Outage probability versus Number of reflecting elements (N) with $\eta = 1$, $\mu = 0.5$, S-IRS, varying P_x dBm. The influence of different values of P_x dBm on the outage probability is depicted in Fig. 8, it is clear that the outage performance improves significantly as the source transmit power P_x dBm increases. In addition, as the number of reflecting elements increases the OP decreases.

Fig. 9 compares the OP performance of Random phase shift (RPS), Optimum phase shift (OPS), and Discrete phase shift (DPS) techniques with different quantization levels. The OPS scheme achieves the best performance due to perfect phase alignment, resulting in the lowest OP across all transmit power levels P_x . In contrast, RPS yields the poorest performance due to random phase shifts, resulting in higher OP. The DPS schemes serve as a practical compromise between OPS and

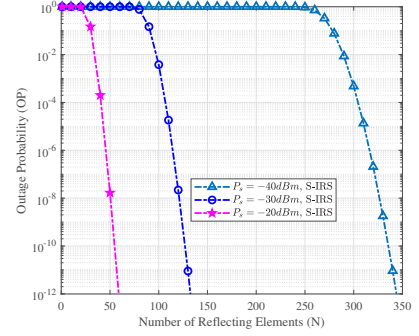


Fig. 8. Outage probability versus Number of reflecting elements (N) with $\eta = 1$, $\mu = 0.5$, S-IRS, varying P_x dBm.

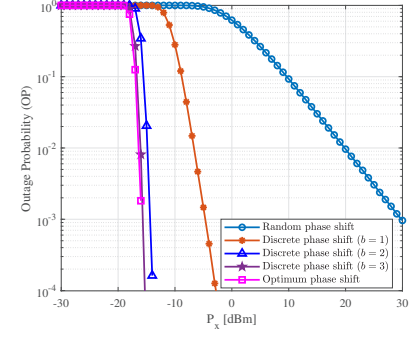


Fig. 9. Outage probability versus source transmit power (P_x) for multiple phases with $L_t = L_r = 1$, $\eta = 1$, $\mu = 1$, $N = 72$, and S-IRS.

RPS. The performance of DPS improves as the number of quantization bits b increases, with 3-bit quantization ($b = 3$) nearly matching the OPS curve. For instance, at an OP of 10^{-3} , the required transmit powers are approximately 30 dBm, -5 dBm, -14 dBm, -16 dBm, and -16 dBm respectively for RPS, DPS ($b = 1$), DPS ($b = 2$), DPS ($b = 3$), and OPS.

Fig. 10 shows the OP performance of a IRS-assisted communication system under $\eta - \mu$ fading for various M -PSK modulation orders ($M \in \{2, 4, 8, 16\}$). As the modulation order increases, the system experiences higher OP at a given transmit power (P_x). For example, to achieve an OP of 10^{-3} , the required transmit power increases from approximately $P_x = -14$ dBm, for $M = 2$ to about $P_x = -1$ dBm for $M = 16$. This trend occurs because higher-order M -PSK schemes have denser constellations with reduced symbol

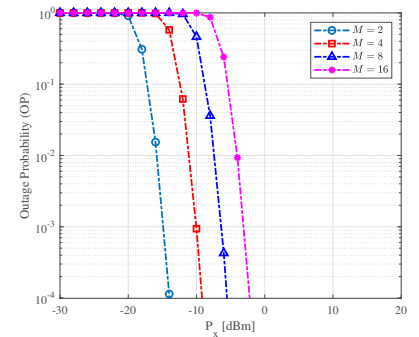


Fig. 10. Outage probability versus source transmit power (P_x) varying M (M -PSK) with $L_t = L_r = 2$, $\eta = 1$, $\mu = 1$, $N = 32$, and S-IRS.

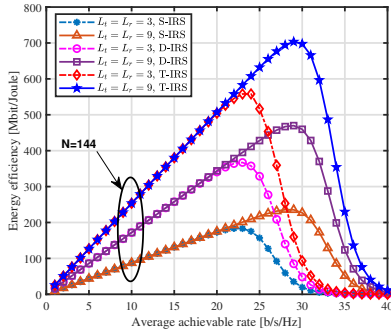


Fig. 11. Energy efficiency versus average achievable rate for varying L_t, L_r with fixed $\eta = 1, \mu = 1$, and $N = 144$.

spacing, making them more susceptible to noise and channel fading. Therefore, more transmit power is required to maintain reliable communication as M increases.

In Fig. 11, we evaluate the EE of S-IRS, D-IRS, and T-IRS systems, all possessing an identical number of reflecting elements, while considering the total power consumption, which encompasses the circuit power dissipation of both the Tx, Rx, and the hardware impairments of IRS elements, in relation to the target spectral efficiency (R) with respect to equation (20). Firstly, the T-IRS system surpasses the D-IRS and S-IRS systems in terms of a fixed average achievable rate. For a rate of 30 b/s/Hz with $(Z_t, Z_r) = (9, 9)$, the EE (in Mbits/Joule) at S-IRS, D-IRS, and T-IRS are 226.94, 459.59, and 697.57, respectively. Furthermore, for elevated achievable rate values, the utilization of multiple IRSs proves more advantageous than the S-IRS system. Secondly, with a constant number of IRSs, an increase in the number of antennas at the BS enhances energy efficiency to attain the same achievable rate. The aforementioned observations are substantiated by the analysis presented in Section IV, which indicates that the T-IRS attains a superior rate relative to D-IRS and S-IRS, allowing for a trade-off between the target average achievable rate and energy efficiency as delineated in equation (20).

VI. CONCLUSIONS

This study examined the outage probability performance analysis of a multi-antenna MIMO communication system supported by multiple IRSs under a generalized η - μ fading environment. The proposed model targets NLOS scenarios, where direct Tx and Rx links are blocked, as commonly seen in urban environments. Using the central limit theorem, we derived the statistical properties of the received SNR for the system in terms of its first and second moments. Consequently, the theoretical expression for the OP utilizing the Q -function was obtained. Simulation results under various parameter configurations validate the accuracy of the theoretical analysis. Energy efficiency showed that the multi-IRS system outperformed the S-IRS system. The findings highlight that system performance is significantly affected by factors such as the number of IRSs, their placement, the count of reflecting elements, channel fading parameters, and the number of antennas at the Tx and Rx. Our findings show that incorporating multiple IRSs significantly improves link

reliability by lowering the OP in blockage-heavy areas. The performance benefits grow with more reflecting elements and higher values of η and μ parameters. Additionally, the use of multiple antennas at both ends further enhances diversity gains. IRS placement also plays a vital role, with closer positioning to the Tx or Rx leading to noticeable improvements. Overall, the study demonstrates that a multi-antenna, multi-IRS MIMO system achieves better performance than a system supported by a single IRS and is also a promising solution for reliable and efficient next-generation wireless communication. The analysis may be expanded to encompass more generalized fading models, such as the α - η - μ distribution, to more accurately reflect real-world variability. Furthermore, integrating machine learning techniques for real-time IRS control can enhance adaptability and performance in dynamic wireless environments.

REFERENCES

- [1] T. Padmavathil, K. K. Cheepurupalli, and R. Madhu, "Evaluation of fbmc channel estimation using multiple auxiliary symbols for high throughput and low ber 5g and beyond communications," *Infocommunications Journal*, vol. 16, no. 2, pp. 25–32, 2024, doi: 10.36244/ICJ.2024.2.4.
- [2] R. Pestourie, C. Pérez-Arancibia, Z. Lin, W. Shin, F. Capasso, and S. G. Johnson, "Inverse design of large-area metasurfaces," *Optics express*, vol. 26, no. 26, pp. 33 732–33 747, 2018, doi: 10.1364/OE.26.033732.
- [3] C. Huang, S. Hu, G. C. Alexandropoulos, A. Zappone, C. Yuen, R. Zhang, M. Di Renzo, and M. Debbah, "Holographic mimo surfaces for 6g wireless networks: Opportunities, challenges, and trends," *IEEE wireless communications*, vol. 27, no. 5, pp. 118–125, 2020, doi: 10.1109/MWC.001.1900534.
- [4] Y. Liu, X. Liu, X. Mu, T. Hou, J. Xu, M. Di Renzo, and N. Al-Dhahir, "Reconfigurable intelligent surfaces: Principles and opportunities," *IEEE communications surveys & tutorials*, vol. 23, no. 3, pp. 1546–1577, 2021, doi: 10.1109/COMST.2021.3077737.
- [5] M. Chaudhary and S. K. Bandari, "A low complex joint optimization model for maximizing sum rate and energy efficiency in an irs-assisted multi-user communication scenario," *Physical Communication*, vol. 63, p. 102 296, 2024, doi: 10.1016/j.phycom.2024.102296.
- [6] M. A. ElMossallamy, H. Zhang, L. Song, K. G. Seddik, Z. Han, and G. Y. Li, "Reconfigurable intelligent surfaces for wireless communications: Principles, challenges, and opportunities," *IEEE Transactions on Cognitive Communications and Networking*, vol. 6, no. 3, pp. 990–1002, 2020, doi: 10.1109/TCCN.2020.2992604.
- [7] Z. Wang, L. Liu, S. Zhang, and S. Cui, "Massive mimo communication with intelligent reflecting surface," *IEEE Transactions on Wireless Communications*, vol. 22, no. 4, pp. 2566–2582, 2022, doi: 10.1109/TWC.2022.3212537.
- [8] S. K. Bandari, V. Mani, and A. Drosopoulos, "Performance analysis of gfdm in various fading channels," *COMPEL: The International Journal for Computation and Mathematics in Electrical and Electronic Engineering*, vol. 35, no. 1, pp. 225–244, 2016, doi: 10.1108/COMPEL-06-2015-0215.
- [9] S. K. Bandari, S. S. Yadav, and V. Mani, "Analysis of gfdm in generalized η - μ fading channel: A simple probability density function approach for beyond 5g wireless applications," *AEU-International Journal of Electronics and Communications*, vol. 153, p. 154 260, 2022, doi: 10.1016/j.aue.2022.154260.
- [10] R. Mallaiah and V. Mani, "A novel ofdm system based on dfrft-ofdm," *IEEE Wireless Communications Letters*, vol. 11, no. 6, pp. 1156–1160, 2022, doi: 10.1109/LWC.2022.3159534.
- [11] W. Mei and R. Zhang, "Joint base station and irs deployment for enhancing network coverage: A graph-based modeling and optimization approach," *IEEE Transactions on Wireless Communications*, vol. 22, no. 11, pp. 8200–8213, 2023, doi: 10.1109/TWC.2023.3260805.

Improved Outage Probability using Multi-IRS-Assisted MIMO Wireless Communications in Cellular Blockage Scenarios

- [12] T. Hou, Y. Liu, Z. Song, X. Sun, Y. Chen, and L. Hanzo, "Reconfigurable intelligent surface aided noma networks," *IEEE Journal on Selected Areas in Communications*, vol. 38, no. 11, pp. 2575–2588, 2020, doi: 10.1109/JSAC.2020.3007039.
- [13] L. Yang, W. Guo, and I. S. Ansari, "Mixed dual-hop fso-rf communication systems through reconfigurable intelligent surface," *IEEE Communications Letters*, vol. 24, no. 7, pp. 1558–1562, 2020, doi: 10.1109/LCOMM.2020.2986002.
- [14] S. Li, B. Duo, X. Yuan, Y.-C. Liang, and M. Di Renzo, "Reconfigurable intelligent surface assisted uav communication: Joint trajectory design and passive beamforming," *IEEE Wireless Communications Letters*, vol. 9, no. 5, pp. 716–720, 2020, doi: 10.1109/LWC.2020.2966705.
- [15] E. Björnson, Ö. Özdogan, and E. G. Larsson, "Intelligent reflecting surface versus decode-and-forward: How large surfaces are needed to beat relaying?" *IEEE Wireless Communications Letters*, vol. 9, no. 2, pp. 244–248, 2019, doi: 10.1109/LWC.2019.2950624.
- [16] A. E. Canbilen, E. Basar, and S. S. Ikki, "Reconfigurable intelligent surface-assisted space shift keying," *IEEE wireless communications letters*, vol. 9, no. 9, pp. 1495–1499, 2020, doi: 10.1109/LWC.2020.2994930.
- [17] W. Zhao, G. Wang, S. Atapattu, T. A. Tsiftsis, and C. Tellambura, "Is backscatter link stronger than direct link in reconfigurable intelligent surface-assisted system?" *IEEE Communications Letters*, vol. 24, no. 6, pp. 1342–1346, 2020, doi: 10.1109/LCOMM.2020.2980510.
- [18] Q. Wu, G. Chen, Q. Peng, W. Chen, Y. Yuan, Z. Cheng, J. Dou, Z. Zhao, and P. Li, "Intelligent reflecting surfaces for wireless networks: Deployment architectures, key solutions, and field trials," *IEEE Wireless Communications*, 2025.
- [19] Y. Du, N. Qi, K. Wang, M. Xiao, and W. Wang, "Intelligent reflecting surface-assisted uav inspection system based on transfer learning," *IET Communications*, vol. 18, no. 3, pp. 214–224, 2024.
- [20] W. R. Braun and U. Dersch, "A physical mobile radio channel model," *IEEE transactions on Vehicular Technology*, vol. 40, no. 2, pp. 472–482, 1991, doi: 10.1109/25.289429.
- [21] M. D. Yacoub, "The κ - μ distribution and the η - μ distribution," *IEEE Antennas and Propagation Magazine*, vol. 49, no. 1, pp. 68–81, 2007, doi: 10.1109/MAP.2007.370983.
- [22] M. H. Samuh, A. M. Salhab, and A. H. A. El-Malek, "Performance analysis and optimization of ris-assisted networks in nakagami-m environment," *arXiv preprint arXiv:2010.07841*, 2020.
- [23] M. A. A. Junior, G. Fraidenraich, R. C. Ferreira, F. A. De Figueiredo, and E. R. De Lima, "Multiple-antenna weibullfading wireless communications enhanced by reconfigurable intelligent surfaces," *IEEE Access*, vol. 11, pp. 107 218–107 236, 2023, doi: 10.1109/ACCESS.2023.3310936.
- [24] R. M. Rafi and V. Sudha, "Ser and outage probability analysis of double ris assisted wireless communication system," *Wireless Personal Communications*, vol. 133, no. 4, pp. 2339–2354, 2023, doi: 10.1007/s11277-024-10869-y.
- [25] R. K. Hindustani, S. Sharma, and D. Dixit, "Outage analysis of multi-irs wireless communication over κ - μ fading with panel selection," in *2024 5th International Conference on Innovative Trends in Information Technology (ICITIT)*. IEEE, 2024, pp. 1–5, doi: 10.1109/ICITIT61487.2024.10580099.
- [26] L. Yang, Y. Yang, D. B. da Costa, and I. Trigui, "Outage probability and capacity scaling law of multiple ris-aided networks," *IEEE Wireless Communications Letters*, vol. 10, no. 2, pp. 256–260, 2020, doi: 10.1109/LWC.2020.3026712.
- [27] Z. Zhang, Y. Cui, F. Yang, and L. Ding, "Analysis and optimization of outage probability in multi-intelligent reflecting surface-assisted systems," *arXiv preprint arXiv:1909.02193*, 2019, doi: 10.48550/arXiv.1909.02193.
- [28] T. N. Do, G. Kaddoum, T. L. Nguyen, D. B. Da Costa, and Z. J. Haas, "Multi-ris-aided wireless systems: Statistical characterization and performance analysis," *IEEE Transactions on Communications*, vol. 69, no. 12, pp. 8641–8658, 2021, doi: 10.1109/TCOMM.2021.3117599.
- [29] S. Penchala, M. Chaudhary, S. K. Bandari, and V. Mani, "Impact of nakagami-m fading channel on reconfigurable intelligent surface symbol error rate in non-line-of-sight scenarios," in *2022 IEEE 19th India Council International Conference (INDICON)*. IEEE, 2022, pp. 1–4, doi: 10.1109/INDICON56171.2022.10040187.
- [30] P. Nayeri, F. Yang, and A. Z. Elsherbeni, *Reflectarray antennas: theory, designs, and applications*. John Wiley & Sons, 2018.
- [31] Q. Wu and R. Zhang, "Beamforming optimization for wireless network aided by intelligent reflecting surface with discrete phase shifts," *IEEE Transactions on Communications*, vol. 68, no. 3, pp. 1838–1851, 2019.
- [32] E. Basar, "Transmission through large intelligent surfaces: A new frontier in wireless communications," in *2019 European Conference on Networks and Communications (EuCNC)*. IEEE, 2019, pp. 112–117, doi: 10.1109/EuCNC.2019.8801961.
- [33] J. G. Proakis and M. Salehi, *Digital communications*. McGraw-hill, 2008.
- [34] M. D. Yacoub, "The κ - μ distribution and the η - μ distribution," *IEEE Antennas and Propagation Magazine*, vol. 49, no. 1, pp. 68–81, 2007.
- [35] M. Simon, "Alouini. ms: Digital communication over fading channels," *A John Wiley Sons.-2005.-900 p*, 2005, doi: 10.1109/TIT.2008.924676.
- [36] A. Annamalai, C. Tellambura, and J. Matyas, "A new twist on the generalized marcum q-function $qm(a, b)$ with fractional-order m and its applications," in *2009 6th IEEE Consumer Communications and Networking Conference*, 2009, pp. 1–5, doi: 10.1109/CCNC.2009.4784840.
- [37] M. H. N. Shaikh, V. A. Bohara, A. Srivastava, and G. Ghatak, "Performance analysis of intelligent reflecting surface-assisted wireless system with non-ideal transceiver," *IEEE Open Journal of the Communications Society*, vol. 2, pp. 671–686, 2021, doi: 10.1109/OJ-COMS.2021.3068866.



Suresh Penchala received his B.Tech and M.Tech degrees in Electronics and Communication Engineering from JNTU Hyderabad, Telangana, India, in 2005 and 2008 respectively. He is currently pursuing a Ph.D. degree at the National Institute of Technology, Meghalaya, India. His research interests include intelligent Reflecting Surfaces, wireless communication beyond 5G, and signal processing for communication. He is an associate member of the Institute of Electronics and Telecommunication Engineers (IETE), India.



Shravan Kumar Bandari received his Ph.D. degree from, NIT Warangal, India in 2018. Received postdoctoral fellowship from GIGA Korea project at Seoul National University, South Korea, 2018–2019. Received ERASMUS-MUNDUS Euphrates scholarship from the European Union under the doctoral mobility programme at TEI of Patras, Greece (now changed to University of the Peloponnese), 2014–2015. Selected candidate to receive Erasmus-Mundus INTERWEAVE scholarship holder at Tallinn University of Technology, Estonia, 2014–2015. He is twice scholarship holder from MHRD India. Since 2019, he has been with the National Institute of Technology Meghalaya, India as an Assistant Professor of Electronics and Communication Engineering. His research interest includes future wireless systems, multi-carrier waveforms for next-generation wireless systems, multi-antenna/multi-user channels, compressed sensing, and cognitive radio.



Venkata Mani Vakamulla received the B.E. and M.E. degrees in electronics and communication engineering from the College of Engineering, Andhra University, India, in 1992 and 2003, respectively, and the Ph.D. degree in electrical engineering from the Indian Institute of Technology Delhi, India, in 2009. She joined the Electronics and Communication Engineering Department, at NIT Warangal, in April 2008, as an Assistant Professor, where she has been a Professor, since July 2022. She has numerous publications in credit in national and international conferences and journals. Her research interests include wireless communication and signal processing for communication. She is a fellow of the Institute of Electronics and Telecommunication Engineers (IETE), India.

Improving the Clipping-based Active Constellation Extension PAPR Reduction Technique for FBMC Systems

Zahraa Tagelsir^{1,2}, and Zsolt Kollár¹

Abstract—Filter Bank MultiCarrier (FBMC) systems are known for their superior spectral efficiency and flexibility of the spectral shape. These benefits are challenged by the high Peak-to-Average Power Ratio (PAPR) and the Complementary Cumulative Distribution Function (CCDF) of the modulated signal, degrading power amplifier efficiency and system performance, causing spectral leakage and increasing Bit Error Rate (BER). This paper investigates the possibility of improvement for clipping-based Active Constellation Extension (ACE) based PAPR reduction techniques for FBMC systems. A modified method is proposed, combining iterative clipping and enclipping. A dynamic Clipping Ratio (CR) selection method is provided to optimize the iterative process. The method also enhances the BER performance by leveraging an enclipping step to increase the transmit signal power. Furthermore, the complexity of the iterative process is reduced by applying efficient and low computational FBMC modulation and demodulation schemes. The simulation outcomes validate that the suggested method efficiently reduces PAPR, optimizes the CCDF and improves the BER performance. Furthermore, the proposed system complexity is also reduced compared to the previous solutions.

Index Terms—FBMC, PAPR, clipping, enclipping, ACE, CR.

I. INTRODUCTION

FILTER Bank MultiCarrier modulation using Offset Quadratic Amplitude Modulation (FBMC-OQAM) system is a multicarrier modulation technology that is a promising physical layer candidate for 5G and 6G communication networks. The central concept of FBMC is to boost the time duration of the symbols in the time domain using a prototype filter. The symbols are designed so that they overlap both in time- and frequency-domain. The overlap in the time domain makes it possible to maintain the original data transmission rate, and the frequency-domain overlap results in a well-localized spectral shape [1].

An inherent challenge in MultiCarrier (MC) transmission schemes like FBMC is the elevated Peak-to-Average Power Ratio (PAPR). This issue arises due to the unique structure of the MC signal, in which numerous independently modulated subcarriers are transmitted in parallel. In certain conditions, the phases of the subcarriers may align constructively, resulting in signal peaks that are significantly higher than the average signal power. While the probability of such extreme peaks is

relatively low, in practice, such signals exhibit considerably higher PAPR values than their single-carrier counterparts [2].

Power Amplifiers (PAs) are designed to operate efficiently near their saturation point, but high PAPR forces them into nonlinear regions, causing significant distortions; the nonlinear effects include spectral spreading, intermodulation, and warping of signal constellations, which degrade system performance. A large input back-off is also required to prevent the PA from operating in its nonlinear region, reducing the amplifier's power efficiency [3]. The nonlinearities result in in-band distortion; the amplified signal experiences different gains, which degrades Bit-Error Rate (BER) performance by introducing additive noise and out-of-band distortion, where intermodulation generates undesired frequency components at frequencies outside the signal range, which leads to interference with adjacent channels [2].

Various studies have explored different aspects of FBMC, including its implementation challenges and optimization techniques. In [4], a lower bound on the achievable PAPR reduction using optimization-based approaches is established.

Numerous schemes have been introduced to reduce PAPR in FBMC systems. These include coding [5], Partial Transmit Sequence (PTS) [6], Companding techniques [7], Selected Mapping (SM) [8], Tone Injection (TI) [9], Tone Reservation (TR) [10] or Active constellation extension [11]. However, while these methods aim to reduce the high PAPR, they have drawbacks, such as:

- increased computational complexity,
- reduced data rate,
- required side information for demodulation,
- increased bit error rate.

This paper investigates the optimization possibilities for the ACE-based PAPR reduction technique for FBMC modulation, as this technique does not reduce the data rate nor require any side information for the demodulation process. The only drawback of this technique is that soft demodulation in the receiver is not possible, as the constellation points are pre-distorted. Thus, the decoding gain for complex channel coding such as Turbo codes, Low-Density Parity Check (LDPC) or Polar codes is lost [12]. Furthermore, ACE can be jointly used effectively with an iterative clipping and filtering-based technique, as presented in [11].

An adaptive Clipping Ratio (CR) optimization approach is proposed, where the CR values are dynamically tuned across multiple iterations to effectively minimize the PAPR of the

¹ Department of Artificial Intelligence and Systems Engineering, Faculty of Electrical Engineering and Informatics, Budapest University of Technology and Economics, Budapest, Hungary.

E-mail: {zahraa, kollarzs}@mit.bme.hu

² Department of Electrical and Electronic Engineering, Faculty of Engineering, Red Sea University, Port Sudan, Sudan.

DOI: 10.36244/ICJ.2025.3.6

signal. Furthermore, a novel hybrid method is introduced, combining clipping and enlipping [13]. This approach combines the advantages of both techniques and improves BER performance by amplifying the signal power in the last enlipping process. Furthermore, a signal processing complexity reduction of the iterative modulation-demodulation of the FBMC signal is also proposed to enable a practical implementation.

The proposed method enhances the gain metric, which balances PAPR reduction and the additional signal power by significantly reducing PAPR with fewer iterations than conventional methods. Moreover, the proposed scheme shows an improvement in complexity by requiring fewer iterations for a comparable or better performance, which reduces the computational overhead of iterative PAPR reduction techniques. Therefore, it enhances the proposed technique's effectiveness regarding performance and is more computationally efficient, addressing a critical trade-off in communication system design.

The primary contributions of this paper are:

- 1) An optimized CR selection strategy across iterations is provided.
- 2) A modified PAPR reduction technique using ACE is proposed, combining iterative clipping and enlipping.
- 3) A reduced computational complexity of the iterative structure is also proposed that requires fewer signal processing steps.

The paper is structured as follows. In the next section, the FBMC system model is introduced, and a statistical analysis of the signal is given, including an evaluation of the Complementary Cumulative Distribution Function (CCDF) of PAPR and the kurtosis for various numbers of subcarriers. The ACE-based PAPR reduction technique, which uses iterative clipping and filtering, is discussed in Section III. Section IV introduces the proposed improvements; the results are verified using simulations. Section V discusses complexity reduction possibilities of the iterative process, and analytical derivations are given for the required number of multiplications. Section VI concludes the paper with key findings and contributions.

II. FBMC MODULATION

FBMC is a multicarrier modulation technique with high potential for the physical layer of future communication systems. In this scheme, the binary data is first converted into complex symbols through Quadrature Amplitude Modulation (QAM). The resulting symbols are then split into real and imaginary parts, each undergoing a phase rotation and passing through a prototype filter. To maintain orthogonality, the real and imaginary parts are transmitted with a half-symbol time offset (OQAM). After filtering, the components are modulated to different subcarriers, and the combined terms form the final baseband output signal [14]. The resulting FBMC-OQAM signal can be mathematically described as follows:

$$x[n] = \sum_{m=-\infty}^{\infty} \sum_{k=0}^{N-1} \left(\theta^k \Re(X_k[m]) f_0[n - mN] + \theta^{k+1} \Im(X_k[m]) f_0[n - mN - N/2] \right) e^{j \frac{2\pi}{N} kn}. \quad (1)$$

where:

- $X_k[m]$: The complex QAM symbol mapped to subcarrier k at time instant m ,
- $\Re(X_k[m]), \Im(X_k[m])$: The real and imaginary parts of the complex symbol $X_k[m]$,
- θ^k and θ^{k+1} : Phase rotation terms: $\theta = e^{j \frac{\pi}{2}}$.
- $f_0[n - mN]$: impulse response of the prototype filter,
- $e^{j \frac{2\pi}{N} kn}$: the complex exponential term responsible for modulating the signal to the k^{th} subcarrier.

One of the significant issues of FBMC systems is the high computational complexity, primarily resulting from the necessity of large-size IFFTs, especially when using the Frequency Spreading (FS) approach to implement the signal modulation/demodulation [1]. The PolyPhase Network (PPN) based transceiver structure eases this problem using polyphase filterbanks with smaller IFFTs [1]. One of the most efficient variants of the PPN-based transmitter is the modified PPN structure presented in [15], which employs a single IFFT combined with two PPN filters and some signal processing blocks. This structure is shown in Fig. 1. The structure significantly reduces the computational complexity compared to traditional implementations. In conventional PPN transmitters, two separate N -point IFFTs are required to process the real and imaginary parts of the signal independently. However, the modified PPN structure leverages a single IFFT to process both components simultaneously. Using a single IFFT block in the PPN structure effectively halves the number of complex multiplications required in the transmitter, making it a computationally efficient solution. This transmitter structure is used in this paper.

Similarly, Fig. 2 illustrates the corresponding PPN-FBMC receiver employing a single FFT. Traditional PPN-based receivers utilize two parallel FFT blocks to process the received signal's real and imaginary components independently. The enhanced PPN architecture substitutes these two FFTs with a single full-size FFT and some additional signal processing. The PPN filtering is initially applied, followed by correcting for the phase rotation factor from the transmitter using a circular shift in the opposite direction [15]. The overlapping real and imaginary segments of the FBMC-OQAM symbols are divided into even and odd components: the even component is derived from the real signal. In contrast, the odd components are received from the imaginary signal. These components are merged later to create complex time-domain symbols, which are subsequently processed by a single FFT to retrieve the transmitted QAM symbols [16]. This receiver structure will be used in this paper.

The statistical properties should also be investigated, as the FBMC signal can be considered multicarrier. The significant dynamic range of the FBMC signal presents a challenge as the signal at the transmitter goes through a PA, which may cause nonlinearities or operate inefficiently if not driven using the entire linear range. Applicable PAs do not preserve linearity across the entire signal dynamic range, leading to uneven amplification of various signal sections. This distortion affects the quality of the signal, resulting in a reduced BER performance [2].

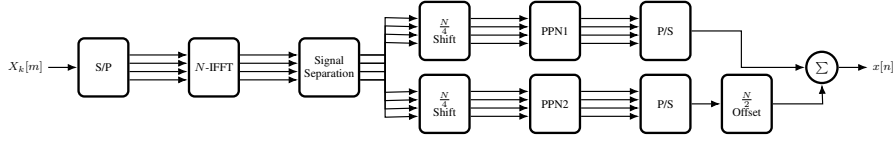


Fig. 1: FBMC transmitter using a single IFFT block.

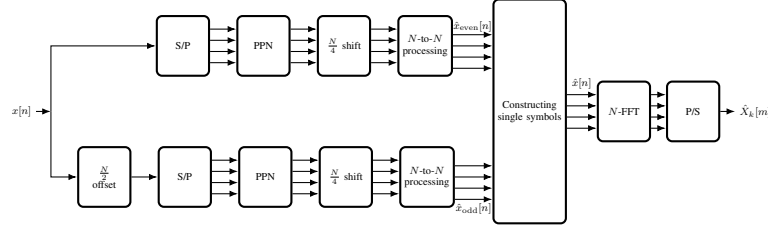


Fig. 2: FBMC receiver using a single FFT block.

To analyze the behaviour of the FBMC signals in the presence of such PAs and to investigate possible PAPR reduction techniques, the statistical metrics such as kurtosis and CCDF of the PAPR are examined for different numbers of subcarriers.

The kurtosis of these components for different subcarriers is evaluated to determine the similarity of the real and imaginary components of the FBMC signal to Gaussian characteristics as the number of subcarriers increases. The kurtosis is a statistical indicator of how often outliers appear. It provides a statistical description of a distribution's shape and relationship to the normal distribution [11]. The kurtosis β of the random variable ζ is defined as:

$$\beta = \frac{\mathbb{E}[\zeta^4]}{(\mathbb{E}[\zeta^2])^2}. \quad (2)$$

It is the ratio of the variable's fourth moment to the square of its average power. A univariate normal distribution has a kurtosis of 3.

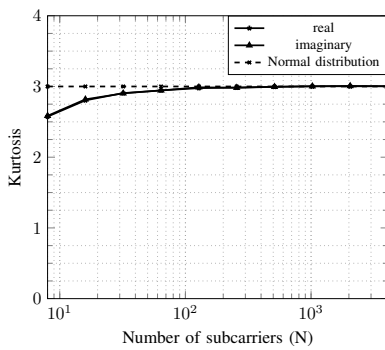


Fig. 3: Kurtosis of the real and imaginary part of the complex FBMC signal in function of the number of subcarriers.

The kurtosis of the real and imaginary parts of the complex baseband FBMC signal in function of the subcarriers can be seen in Fig. 3. As the number of subcarriers grows, the kurtosis values approach 3, indicating convergence toward a normal distribution. Above $N = 256$ subcarriers, it can be concluded that the real and imagery parts of the FBMC signal are Gaussian distributed.

A further metric of the FBMC signal is the PAPR, which measures the difference between a signal's peak and average power. It can be used to characterize the study of the dynamic features of FBMC signals. The PAPR for the m^{th} FBMC symbol can be expressed as:

$$\text{PAPR}(x_m[n]) = \frac{\max_{n \in [0, L-1]} (|x_m[n]|^2)}{\mathbb{E}(|x_m[n]|^2)}. \quad (3)$$

where L is the length of the prototype filter. Furthermore, the CCDF is used for the statistical distribution of the PAPR. The CCDF denotes the likelihood of the PAPR values surpassing a specified limit PAPR_0 :

$$\text{CCDF} = \text{Prob}(\text{PAPR} > \text{PAPR}_0). \quad (4)$$

For the FBMC signal, the CCDF of the PAPR was analyzed for a different number of subcarriers. The CCDF can be seen in Fig. 4. It can be seen that a rising number of subcarriers leads to an increased PAPR due to the signal approaching a Gaussian distribution with a higher number of independent variables, consistent with the central limit theorem. This metric will also be used to assess the efficacy of the clipping-based PAPR reduction method presented in the next section.

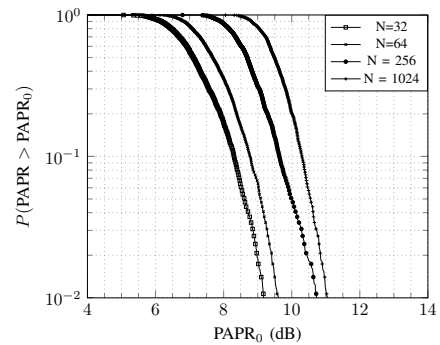


Fig. 4: CCDF of the PAPR of the FBMC symbols for different numbers of subcarriers.

III. CLIPPING-BASED ITERATIVE PAPR REDUCTION VIA ACE

This section briefly introduces the clipping-based PAPR reduction combined with ACE and filtering. This method was introduced and evaluated in [11].

A. Iterative clipping and filtering

Clipping is a non-linear signal processing technique that lowers the amplitude of peaks in a transmitted signal that exceeds a pre-defined threshold. It is done to minimize the signal's dynamic range and thus make it easier to transmit. The process entails clipping such high peaks while trying, as much as possible, to preserve the integrity of the underlying data. The clipped FBMC signal can be expressed as

$$x_c[n] = \begin{cases} x[n] & \text{if } |x[n]| \leq A_{\max} \\ A_{\max} e^{j\phi(x_n)} & \text{if } |x[n]| > A_{\max} \end{cases}, \quad (5)$$

where $x_c[n]$ is the clipped signal, and $A_{\max}$ is the clipping threshold. The clipping threshold is determined by a parameter called the CR, which is mathematically defined as:

$$\text{CR}_{\text{dB}} = 20 \log_{10}(\gamma) \quad (6)$$

with

$$\gamma = \frac{A_{\max}}{\sqrt{P_x}},$$

where P_x is the signal power.

Clipping is simple to implement and computationally efficient, so it's an attractive choice for PAPR reduction. It can considerably increase the efficiency of power amplifiers by properly selecting the threshold value. However, clipping causes both in-band and out-of-band distortions, increasing BER and introducing unwanted distortion in the spectral shape. It can eliminate the out-of-band radiation after demodulating the signal and resetting the unused subcarriers to zero, i.e. filtering them out. Still, it will also increase the peaks of the newly generated signal. The in-band distortion can be handled in several ways; in the case of TR, some carriers are kept so that they are used for PAPR reduction purposes; in the case of ACE, the constellation points are allowed to be distorted as long as the no BER degradation is introduced during demodulation. This repeated interactive method using clipping and filtering with TR and ACE can be used effectively to reduce the PAPR of the FBMC signal. In the next subsection, the ACE method will be described in more detail.

B. Active constellation extension

The concept of ACE was initially introduced in [17] for mitigating the PAPR in OFDM systems that employ QAM constellations. ACE is usually applied along with clipping to minimize the distortions introduced to the data subcarriers. ACE modifies the amplitude and phase of some of the constellation values located at the outer points of the constellation of the signal. These points can be shifted outwards, within permitted regions, to create a new representation of the same symbols. By using ACE on the subcarriers distorted by clipping, the QAM values are restored, and in parallel,

the BER performance is also improved while maintaining the clipped and distorted values of the given subcarrier for PAPR reduction.

A 4-QAM modulation is employed to optimize the constellation-shaping process, where four possible constellation points are evenly distributed across each quadrant of the complex plane, and constellation points are modified within the quarter-plane outside the nominal point. The extension regions for 2-QAM, 4-QAM and 16-QAM are depicted in Fig. 5. For higher-order modulations, the ACE is less effective as fewer QAM values can be extended toward outer regions, thus leading to a smaller improvement in the PAPR reduction performance of the time domain signal. During our investigations, we will use the constellation with 4-QAM.

IV. IMPROVEMENTS TO THE CLIPPING BASED ACE METHOD

A. Improving the clipping-based PAPR reduction scheme

1) *Fixed CR*: The most straightforward approach involves maintaining the CR constant over the iterations. To evaluate the performance, the CCDF of the PAPR is observed at a specific probability after each iteration. The best CR value that achieves the smallest PAPR for a given number of iterations will be selected.

2) *Adaptive CR*: PAPR reduction techniques and iterative methods such as clipping combined with ACE are commonly used to achieve substantial reduction. However, traditional approaches often rely on a fixed CR across all iterations. While effective to an extent, this static approach does not leverage the evolving characteristics of the signal throughout the iterative process. An adaptive CR approach is also suggested in this paper to optimize the clipping thresholds across iterations dynamically. The algorithm optimizes the PAPR reduction process by systematically evaluating all possible combinations of CR values across iterations and selecting the sequence that achieves the minimum PAPR. This exhaustive search ensures the best possible PAPR reduction with controlled signal distortion. Similar to the fixed CR case, the CCDF of the PAPR is observed during the iterative process. During each iteration, the PAPR Reduction algorithm is implemented for different CRs. Once the set number of iterations has been completed, the CR values that achieve the most significant PAPR reduction are selected. The exhaustive search minimizes the PAPR by considering all possible combinations. This ensures better performance than fixed CR approaches.

B. Applying the enlipping technique

The suggested technique presents an innovative way to reduce iterative PAPR in FBMC systems. It combines conventional iterations of clipping and ACE with a final enlipping and ACE step. This method is designed to achieve superior PAPR reduction, minimize signal distortion, and improve BER performance. Enlipping, introduced in [13], involves adjusting the amplitude of each signal to its peak value instead of limiting the high peaks to a predetermined level:

$$x_e[n] = A_{\max} e^{j\phi(x_n)}, \quad (7)$$

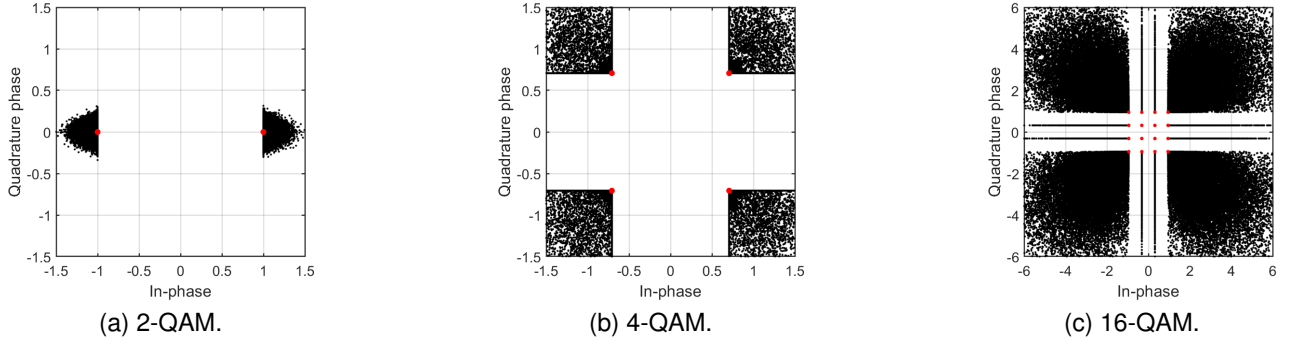


Fig. 5: Extension regions of ACE for 2-QAM, 4-QAM and 16-QAM. Red dots mark the ideal cancellation points.

where $x_e[n]$ is the enlipped signal, $A_{\max}$ is the clipping threshold, and ϕ is the signal phase. Despite the signal samples exceeding or falling below the clipping threshold $A_{\max}$, every signal sample will be modified to align with $A_{\max}$ while preserving its phase value.

The unique feature of this method is found in the adaptive CR used in the clipping iterations and the enlipping process at the end. While clipping reduces the signal peaks to control PAPR, the final enlipping step increases the signal power without increasing the signal peaks, thereby improving BER by improving the SNR. This dual-action methodology navigates the compromise between PAPR minimization and overall system performance.

Clipping limits the signal peaks that exceed a given CR value, effectively reducing PAPR. However, this can distort the signal and cause constellation points to migrate outside their valid decision regions. When this happens, the receiver may misinterpret these points, resulting in significant BER degradation. ACE selectively moves the distorted points back to their legitimate constellation boundaries to solve this problem. The restoration process will ensure the modified symbols stay within acceptable decision regions to maintain the transmitted signal's integrity and ensure minimal BER degradation while maintaining effective PAPR reduction. The iterative method dynamically lowers PAPR, preparing it for the concluding stage. In the final step, we apply enlipping, then ACE, enlipping modifies all signal amplitudes to match a specified CR value, enhancing the average signal power. This stage offsets the losses caused by clipping in previous iterations.

The block diagram of the proposed modified iterative technique extended with enlipping is displayed in Fig. 6.

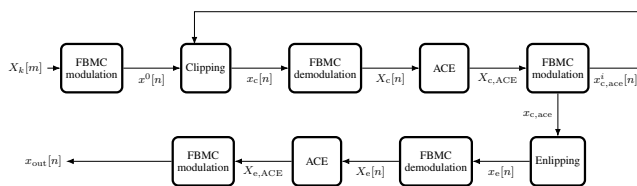


Fig. 6: Block diagram of the proposed iterative PAPR reduction scheme using ACE extended with enlipping for FBMC.

C. Simulation results

The parameters used throughout the simulation are outlined in Table I for the applied FBMC system. From the available 256 subcarriers, 170 were modulated, and an oversampling factor of 4 was applied to ensure that the peaks in the signal were well approximated. A 4-QAM is selected for the modulation alphabet so that the ACE can be well utilized while maintaining a relatively high data rate. During the simulation, the PHYDYAS filter is applied to the prototype filter with an overlapping factor (K) of 4. To have a statically large enough dataset for evaluating the PAPR, 1000 FBMC symbols were applied.

 TABLE I
SIMULATION PARAMETERS.

Parameter	Value
Number of subcarriers (N)	256
Number of data subcarriers	170
Modulation	4QAM
Oversampling	4
Prototype filter	PHYDYAS
Overlapping factor (K)	4
Number of symbols	1000

The choice of the CR value over the iterations significantly influences the overall performance of PAPR reduction. The first step in the case of fixed CR is to identify which CR values result in the maximum PAPR reduction of the FBMC signal [18]. CR values ranging from (CR = 0 dB to 8 dB) with steps of 0.25 dB were tested for 1, 2 and 3 iterations to evaluate the effectiveness of the iterative method in reducing PAPR. The results are shown for different numbers of iterations in Fig. 7. The identified optimal CR values for different iterations are as follows:

- 5.25 dB achieving 9.92 dB PAPR value for 1 iteration,
- 5.25 dB achieving 9.26 dB PAPR value for 2 iterations,
- 5.50 dB achieving 8.81 dB PAPR value for 3 iterations.

Further optimization can be done using adaptive CR. Fig. 8 shows the obtained CR values for 3 iterations for different combinations throughout the iterative process. The results show that the best combination of CR values ($CR_1 = 4.50$ dB, $CR_2 = 5.50$ dB, $CR_3 = 5.75$ dB) was identified by exhaustively searching all possible combinations. The iteration process is also analyzed in details shown in Fig. 9. The

Improving the Clipping-based Active Constellation Extension PAPR Reduction Technique for FBMC Systems

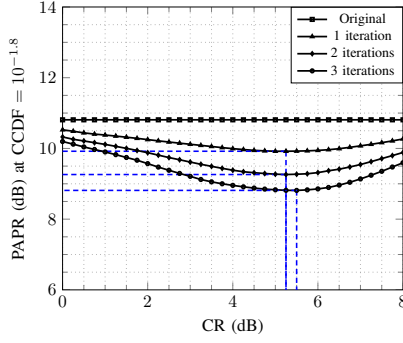


Fig. 7: Resulting PAPR with fixed CR over the iterations. Optimal CR values are indicated by dashed guides.

initial PAPR values are depicted with a dashed line, and a closely linear PAPR improvement throughout the iterations is shown in a solid line. Furthermore, the CR values used in the adaptive clipping with blue color show a slight logarithmic increase through the iterative process achieving a minimum PAPR of 8.79 dB in the target probability, resulting in a slight improvement of 0.02 dB compared to the fixed CR method. The gain can be increased by utilizing higher iteration counts; however, the computation complexity increases exponentially in parallel.

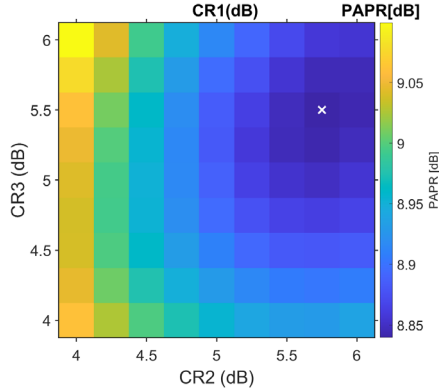


Fig. 8: The achievable PAPR for adaptive CR values. The optimal PAPR value is marked with a white cross. Minimum PAPR = 8.7940 dB with optimal CR values of [4.50, 5.5, 5.75] dB.

The proposed scheme performs clipping with ACE for three iterations using the optimized CR values resulting from Fig. 8 followed by a single enlipping operation combined with ACE, which leads to further improvements in PAPR reduction, BER performance, and power gain. In Fig. 10, the PAPR CCDF shows a notable PAPR reduction of the processed signal curve compared to the original signal. This enhancement is achieved through iterative clipping combined with ACE, effectively reducing the signal's high peaks.

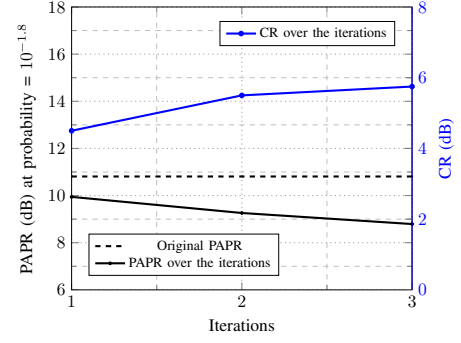


Fig. 9: PAPR and CR values vs clipping iterations for different numbers of subcarriers.

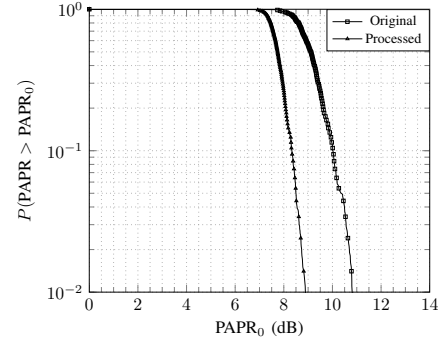


Fig. 10: CCDF of PAPR.

In addition, in Fig. 11a, the BER performance in function of the SNR is shown, and it can be seen that the proposed method produces a gain of 1 dB. This enhancement is due to the enlipping stage, which amplifies the average signal power, consequently improving the BER performance. Furthermore, Fig. 11b shows The PSD of the original and processed signals. The processed signal displayed negligible spectral regrowth, meeting the necessary out-of-band emission limit, and the spectral mask of the FBMC signal was kept.

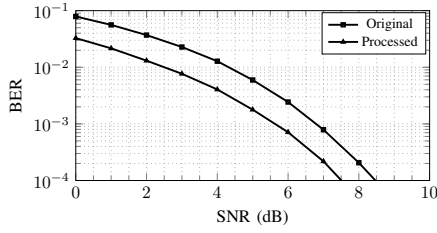
Fig. 12 illustrates the increase in power obtained through the iterative approach. The gain stays low throughout the initial three iterations of clipping and ACE, but a significant increase can be observed after implementing the last step, achieving 3.91 dB power gain. This highlights the technique's capacity to preserve signal integrity while boosting power.

The efficacy of the PAPR reduction method is assessed through a gain metric (Θ), which considers both PAPR reduction and the extra signal power utilized. Following the formulation in [11], the definition of PAPR gain is as follows:

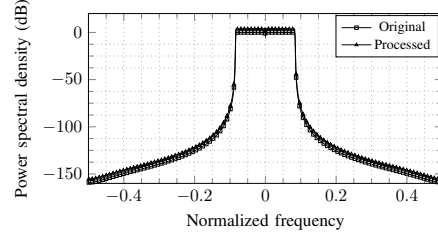
$$\Theta = \max_m \{ \text{PAPR}(x[m])_{\text{dB}} \} - \max_m \{ \text{PAPR}(x_{\text{new}}[m])_{\text{dB}} \} + \Delta P_s, \quad (8)$$

where

- $\max_m \{ \text{PAPR}(x[m])_{\text{dB}} \}$: The maximum PAPR of the original signal.
- $\max_m \{ \text{PAPR}(x_{\text{new}}[m])_{\text{dB}} \}$: The maximum PAPR of the processed signal.
- $\Delta P_s = 10 \cdot \log_{10} \left(\frac{\text{Power of processed signal}}{\text{Power of the original signal}} \right)$.



(a) BER vs SNR.



(b) Power Spectral Density (PSD).

Fig. 11: Statistical analysis of the proposed scheme.

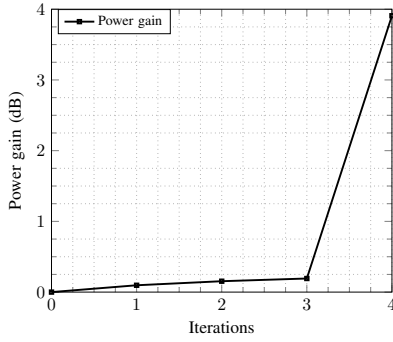


Fig. 12: Power gain.

After 3 iterations of clipping and ACE, a PAPR reduction gain of 2.24 dB is achieved, combining the power gain with the PAPR gain results in an overall gain metric of 6.15 dB. The proposed method achieves a higher Θ than the iterative clipping and ACE method in [11] as we can see in Fig. 13, requiring fewer iterations to reach superior performance. Although the traditional approach consistently enhances Θ with an increasing number of iterations, the additional step in the proposed method dramatically elevates Θ , minimizing the requirement for further iterations and enhancing efficiency in complexity. This renders the proposed method more robust and computationally efficient for PAPR reduction.

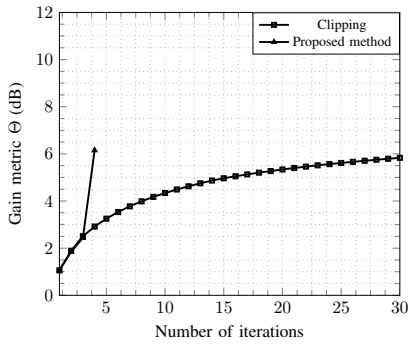


Fig. 13: Gain metric.

TABLE II

NOTATION USED IN COMPLEXITY ANALYSIS

Symbol	Description
M_{mod}	Number of multiplications in one FBMC modulation
M_{demod}	Number of multiplications in one FBMC demodulation
M_{total}	Total number of multiplications of the proposed scheme
I	Number of iterations
N	Number of subcarriers
K	Overlapping factor

V. ANALYSIS OF THE COMPUTATIONAL COMPLEXITY OF THE ITERATIVE PAPR REDUCTION SCHEME

In this section, the computational complexity of the proposed approach is derived; the analysis focuses explicitly on the number of multiplications required in the FBMC modulator and demodulator, as other steps introduce negligible overhead. The proposed method depicted in Fig. 6 include the following signal processing steps:

- 1) Initial modulation, generating the original FBMC signal.
- 2) Iteration steps consisting of clipping, demodulation, ACE, and modulation operations.
- 3) A final step for enclipping, demodulation, ACE, and modulation operations to generate the FBMC signal with reduced PAPR and enlarged signal power.

The main symbols used in the complexity analysis are summarized in Table II. Using this notation, the total number of multiplications required can be expressed as follows:

$$M_{\text{total}} = M_{\text{mod}} + I \cdot (M_{\text{demod}} + M_{\text{mod}}) + M_{\text{demod}} + M_{\text{mod}}, \quad (9)$$

where the required number of multiplications of the modulator and demodulator can be expressed based on [19] for PPN using a single IFFT/FFT based on the number of subcarriers N , the overlapping factor K , and the number of iterations I as

$$M_{\text{mod}} = N(\log_2 N - 3) + 4 + 4KN, \quad (10)$$

$$M_{\text{demod}} = N(\log_2 N - 3) + 4 + 4KN. \quad (11)$$

As a result, the number of multiplications required for the proposed scheme in the function of the iterations can be given based on (9), (10), and (11) as

$$M_{\text{total}} = (2I + 3)(N(\log_2 N - 3) + 4 + 4KN). \quad (12)$$

Improving the Clipping-based Active Constellation Extension PAPR Reduction Technique for FBMC Systems

TABLE III
COMPARISON OF PAPR-REDUCTION METHODS.

Method	Side information	Bit-rate loss	Complexity	Power increase
PTS (Partial Transmit Sequences) [20]	Yes (phase index)	Yes (SI bits)	High	No
SLM (Selected Mapping) [20]	Yes (sequence index)	Yes (SI bits)	High	No
TI (Tone Injection) [21]	No	No	High	Yes
TR (Tone Reservation) [10]	No	Yes (reserved tones)	Low	Yes
Clipping (+ ACE) [22]	No	No	Low	No
Proposed Method	No	No	Low	Yes

The required number of multiplications based on (12) for 3 iterations is shown Fig. 14. The plot depicts the overall computational complexity of the suggested approach in the function of the number of subcarriers. The figure also reveals the number of multiplications required in the case of the standard FS- and PPN-based approaches. The analysis shows a linear growth in complexity with a growing number of subcarriers. The results indicate that the FS and PPN approaches exhibit significantly higher computational complexity than the suggested PPN approach.

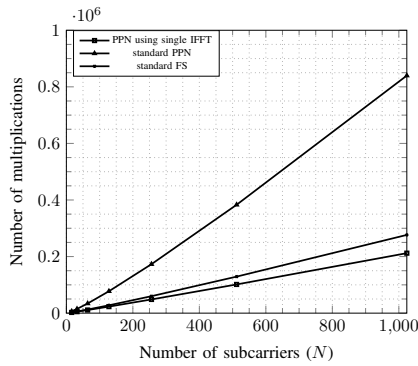


Fig. 14: Required number of multiplications of the iterative PAPR reduction method with 3 iterations using FS, PPN, and PPN with single IFFT as modulator and demodulator.

To provide a structured view of common PAPR-reduction techniques, Table III summarizes qualitative trade-offs for side information, bit-rate loss, complexity, and power increase for the proposed method and representative baselines.

VI. CONCLUSIONS

In this paper, a statistical analysis of FBMC signals was conducted, including evaluating the CCDF of PAPR and kurtosis for different numbers of subcarriers. A novel method for minimizing PAPR in FBMC systems was presented. The proposed scheme applies three iterations of (clipping and ACE) to suppress the most significant peaks, followed by one iteration of (enclipping and ACE) to increase average signal power and improve BER performance. The CR values were optimized for each iteration by searching a range to determine the CR schedule that minimizes the resulting PAPR. The introduced enclipping aids improve the BER performance by raising the average signal power. The proposed scheme performed better in PAPR reduction and BER performance than the conventional schemes while maintaining computational efficiency.

ACKNOWLEDGMENTS

This study receives partial funding from the National Research Development and Innovation Office (NKFIH) via the OTKA Grant (K 142845) and the János Bolyai Research Grant (BO/00042/23/6) provided by the Hungarian Academy of Sciences.

REFERENCES

- [1] Z. Kollár and H. Al-Amaireh, "FBMC transmitters with reduced complexity," *Radioengineering*, vol. 27, no. 4, pp. 1147–1154, 2018, doi: 10.13164/re.2018.1147.
- [2] M. Rahim, T. Stitz, and M. Renfors, "Analysis of clipping-based PAPR-reduction in multicarrier systems," in *VTC Spring 2009-IEEE 69th Vehicular Technology Conference*. IEEE, 2009, pp. 1–5, doi: 10.1109/VETECS.2009.5073391.
- [3] E. Costa, M. Midrio, and S. Pupolin, "Impact of amplifier nonlinearities on OFDM transmission system performance," *IEEE Communications Letters*, vol. 3, no. 2, pp. 37–39, 1999, doi: 10.1109/4234.749355.
- [4] B. Horváth and P. Horvath, "Establishing lower bounds on the peak-to-average-power ratio in filter bank multicarrier systems," in *Infocommunication Journal*, 2015.
- [5] P. Mishra and M. Amin, "PAPR reduction in OFDM using various coding techniques," *International Journal of Wireless and Mobile Computing*, vol. 15, no. 1, pp. 16–20, 2018, doi: 10.1504/IJWMC.2018.10015847.
- [6] N. Shi and S. Wei, "A partial transmit sequences based approach for the reduction of peak-to-average power ratio in FBMC system," in *2016 25th Wireless and Optical Communication Conference (WOCC)*. IEEE, 2016, pp. 1–3, doi: 10.1109/WOCC.2016.7506550.
- [7] S. Ramavath and U. C. Samal, "Theoretical analysis of PAPR companding techniques for FBMC systems," *Wireless Personal Communications*, vol. 118, no. 4, pp. 2965–2981, 2021, doi: 10.1007/s11277-021-08164-1.
- [8] C. Yang, R. Liu, L. Zhao, and Y. Wei, "Modified SLM scheme of FBMC signal in satellite communications," *IET Communications*, vol. 13, no. 11, pp. 1702–1708, 2019, doi: 10.1049/iet-com.2018.5101.
- [9] R. Gopal and S. K. Patra, "Combining tone injection and companding techniques for PAPR reduction of FBMC-OQAM system," in *2015 Global Conference on Communication Technologies (GCCT)*. IEEE, 2015, pp. 709–713, doi: 10.1109/GCCT.2015.7342756.
- [10] M. Laabidi, R. Zayani, and R. Bouallegue, "A new tone reservation scheme for PAPR reduction in FBMC/OQAM systems," in *2015 International Wireless Communications and Mobile Computing Conference (IWCMC)*. IEEE, 2015, pp. 862–867, doi: 10.1109/IWCMC.2015.7289196.
- [11] Z. Kollár, L. Varga, B. Horváth, P. Bakki, and J. Bitó, "Evaluation of clipping based iterative PAPR reduction techniques for FBMC systems," *The Scientific World Journal*, vol. 2014, no. 1, p. 841 680, 2014, doi: 10.1155/2014/841680.
- [12] B. Tahir, S. Schwarz, and M. Rupp, "Ber comparison between convolutional, turbo, ldpc, and polar codes," in *2017 24th International Conference on Telecommunications (ICT)*, 2017, pp. 1–7, doi: 10.1109/ICT.2017.7998249.

- [13] Z. Tagelsir, H. Al-Amaireh, and Z. Kollár, "Active constellation extension with enlipping technique for PAPR reduction in FBMC-OQAM systems," in *2024 34th International Conference Radioelektronika (RADIOELEKTRONIKA)*, 2024, pp. 1–5, doi: 10.1109/RADIOELEKTRONIKA61599.2024.10524051.
- [14] M. Bellanger, D. Le Ruyet, D. Roviras, M. Terré, J. Nossek, L. Baltar, Q. Bai, D. Waldhauser, M. Renfors, T. Ihalainen *et al.*, "FBMC physical layer: a primer," *Phydyas, January*, vol. 25, no. 4, pp. 7–10, 2010, technical report. [Online]. Available: <http://www.ict-phydyas.org>
- [15] L. Varga and Z. Kollár, "Low complexity FBMC transceiver for FPGA implementation," in *2013 23rd International Conference Radioelektronika (RADIOELEKTRONIKA)*. IEEE, 2013, pp. 219–223, doi: 10.1109/RadioElek.2013.6530920.
- [16] H. Al-Amaireh, "Optimization of FBMC-OQAM transceiver architectures," Ph.D. dissertation, *Budapest University of Technology and Economics*, Budapest, Hungary, 2021. [Online]. Available: <https://repozitorium.omikk.bme.hu/items/ee124ef3-fabb-4ceb-8a4e-a5ecc9aebc96>
- [17] B. Krongold and D. L. Jones, "PAR reduction in OFDM via active constellation extension," *IEEE Transactions on broadcasting*, vol. 49, no. 3, pp. 258–268, 2003, doi: 10.1109/ICASSP.2003.1202695.
- [18] B. Horvath and B. Botlik, "Optimization of tone reservation-based PAPR reduction for ofdm systems," *Radioengineering*, vol. 26, no. 3, pp. 791–797, 2017, doi: 10.13164/re.2017.0791.
- [19] H. Al-Amaireh and Z. Kollár, "Complexity comparison of filter bank multicarrier transmitter schemes," in *2018 11th international symposium on communication systems, networks & digital signal processing (CSNDSP)*. IEEE, 2018, pp. 1–4, doi: 10.1109/CSNDSP.2018.8471795.
- [20] T. Jiang and Y. Wu, "An overview: Peak-to-average power ratio reduction techniques for OFDM signals," *IEEE Transactions on broadcasting*, vol. 54, no. 2, pp. 257–268, 2008, doi: 10.1109/TBC.2008.915770.
- [21] M. C. P. Paredes and M. Garcia, "The problem of peak-to-average power ratio in OFDM systems," *arXiv preprint arXiv:1503.08271*, 2015, doi: 10.48550/arXiv.1503.08271.
- [22] Y.-C. Wang and Z.-Q. Luo, "Optimized iterative clipping and filtering for PAPR reduction of OFDM signals," *IEEE Transactions on communications*, vol. 59, no. 1, pp. 33–37, 2010, doi: 10.1109/TCOMM.2010.102910.090040.



Zahraa Tagelsir received her MSc degree in Electrical Engineering in 2018 from the University of Khartoum. She started her PhD studies in 2022 at the Budapest University of Technology and Economics (BME). Her research interests include digital signal processing and wireless communication.



Zsolt Kollár is an Associate Professor in the Department of Artificial Intelligence and Systems Engineering at the Budapest University of Technology and Economics (BME). His research focuses on wireless communication systems, signal processing, and cognitive radio technologies.

Convolution Assisted Polar Encoder with Flexible Iterative Decoding for Forward Error Correction in Wireless Communication Incorporated in Fpga

T. Ranjitha Devi^{*1}, and Dr. C. Kamalnathan²

Abstract—The effectiveness of Forward Error Correction (FEC) approaches depends on the coding scheme and code length, which result in increasing processing latency and delay. Hence, a novel Convolution Assisted Polar Encoder with Flexible Iterative Decoding (CAPE-FID) is proposed to improve communication efficiency by reducing latency and retransmission. Existing approaches use longer block length codes for error correction and seamless communication, but it unnecessarily increases bits, affecting bandwidth and data rate. Hence, a novel Polarized Convolutional Encoder (PCE) is introduced, which uses Adaptive Frozen Polar Coding (AFPC) and Convolution coding to eliminate unnecessary high redundancy from the input data, improving both the useful data rate and the efficient use of available bandwidth. Existing CRC algorithms' lower-degree polynomials during divisional detection cause collisions, thus incorrectly identifying error-free data, causing overhead and affecting data transmission overall. Hence, a novel Flexible Turbo Decoder (FTD) is introduced, which uses Reed-Solomon Euclid (RSE) Code and Sequential Concatenated Turbo coding to reduce packet loss and improve data transmission quality and network congestion handling. Finally, the error-corrected data is decoded using Adaptive Polar Coding with reversed polarization transformation (APC-Rpt). The results show that the proposed method has improved efficiency and reduced retransmission rate, latency and error rates.

Index Terms—Forward Error Correction, Convolution codes, Polar Codes, Euclidean algorithm, Reed-Solomon code, Turbo coding

I. INTRODUCTION

Forward Error Correction (FEC) is a key technique in wireless communication that improves data accuracy and reliability by adding redundant bits at the transmitter. These bits enable the receiver to detect and correct errors caused by fading, noise, interference, and signal attenuation, without requiring retransmission. By encoding data with error-correcting codes and transmitting both the data and redundancy, FEC mitigates channel-induced errors and ensures more consistent performance [1-4].

FEC enhances wireless communication reliability by reducing bit error rate and mitigating channel defects, enabling data recovery even with weak or corrupted signals. This is vital for real-time applications or scenarios with power or latency constraints that preclude retransmission. Common FEC

techniques include polar, Reed-Solomon, Low-Density Parity-Check (LDPC), and convolutional codes, each offering different trade-offs between error correction, complexity, and efficiency. Widely used in digital broadcasting, Wi-Fi, cellular networks (3G, 4G, and 5G), satellite communication, and wireless sensor networks, FEC ensures stable and robust connections even in challenging wireless environments [5-8].

Retransmission addresses the challenges of defective wireless channels by recovering lost or corrupted packets, improving reliability, and ensuring correct data delivery. However, excessive retransmissions can increase delays and reduce throughput, making process optimization vital for balancing reliability and efficiency. Protocols like Automatic Repeat reQuest (ARQ) trigger retransmission when ACK or NAK responses are not received. Error correction systems must adapt to unstable channel conditions, adjusting decoding settings in real time, which is a challenge under channel variability. Multipath propagation and interference from neighboring signals can further reduce effectiveness, requiring advanced decoding to separate interference from the intended signal. Achieving a balance between error correction capabilities and processing efficiency is difficult, as stronger techniques often demand significant resources [9-12].

High-speed communication systems use FEC to maintain reliability at high data rates, but FPGA (Field Programmable Gate Array)-based implementations face challenges in managing clock domains, serialization/deserialization, and synchronization. Compared to ASICs (Application-Specific Integrated Circuit), FPGAs generally consume more power, requiring optimization techniques such as reducing unnecessary computations, using low-power resources, and implementing power-gating. Validating FEC accuracy and performance on FPGAs is also complex due to their concurrent nature, demanding extensive test benches, simulations, and fault scenario validation [13-15]. Despite ongoing research, many aspects of wireless communication still need improvement for effective data transfer without losses, highlighting the need for a novel FEC-based approach. The main contributions of this paper are:

- To reduce the unnecessary high redundancies added to the input data, a novel PCE is introduced, which improves the useful data rate and effective usage of available bandwidth.

- To avoid packet loss and to improve the overall data transmission quality, a novel FTD is introduced, which detects, locates and corrects the errors based on the sequential manner.

The proposed work is divided into five sections, the first is

¹* G. Pullaiah College of Engineering and Technology (GPCET), Department of Electronics and Communication Engineering, Kurnool, Andhra Pradesh, India (e-mai: tranjithadevicee@gpcet.ac.in)

² Department of Electrical Electronics and Communication Engineering, GITAM School of Technology, Bengaluru, India (e-mai: kchandra@gitam.edu)

an introduction and the second is a review of the literature. Section 3 encapsulates the proposed method discussing sub modules. Performance and comparative analysis are covered in Section 4. Section 5 delves deeply into the work's conclusion.

II. LITERATURE SURVEY

Vinodhini et al. [16] suggested a Transient Error Correction (TEC) coding technique for a low-power data link layer in Network on Chip (NoC) for error correction capability with minimal hardware complexity. Validated through simulations and realistic traffic patterns, it achieved more reliable transmission at lower link swing voltages and reduced power consumption compared to the Hamming product code, albeit with a slight increase in codec delay and router latency. However, it did not fully balance error correction with bandwidth efficiency.

Syed Mohsin Abbas et al. [17] outlined the GRAND-MO algorithm's hardware architecture for decoding linear block codes in memory-based communication channels, which was enhanced for simplifying hardware implementation. The improved VLSI design achieved higher throughput and decoding gains compared to both the GRANDAB and BCH decoders, though it still requires adaptation to varying channel conditions for optimal error correction.

Konda Nandan Kumar et al. [18] proposed a Matrix-based error correction and parity-sharing approach for detecting and rectifying adjacent and certain multiple-bit faults. By computing specific bits using a DNA-shaped curve within the matrix, the approach reduced power, area, and latency, making it suitable for memory systems, though some residual errors remained uncorrected.

Qilin Zhang et al. [19] created a deep learning-based, SPA-based, data-driven method for decoding CRC codes that coupled internal trainable parameter optimization with learning. A loss function based on a weighted average of negentropy was employed for evaluating the Gaussianity and binary cross-entropy (BCE) of the decoder output. Though the proposed strategy with deep unfolding and negentropy-aware loss function enhanced the bit error rate, it failed to ensure interoperability between different systems.

Linfang Wang et al. [20] presented cyclic redundancy check (CRC) aided probabilistic amplitude shaping (PAS) on a trellis-coded modulation (TCM) to achieve the short-block length random coding union (RCU) condition. The transmitter's distribution matcher first encoded similar message bits to produce amplitude symbols with the necessary distribution, then encrypted using a CRC, and finally, encoded and modulated using Ungerboeck's TCM scheme. However, there was a trade-off between the level of redundancy and the achievable error correction capability.

Svitlana Matsenko et al. [21] evaluated indivisible codes for binary symmetric channels using the average probability method and implemented a fractal decoder in FPGA with multilayer PAM and grey code mapping for short-reach optical interconnects. While effective for end-to-end data management with uniform structure and high speed, the approach was inefficient in handling burst errors and required improved code mapping.

Bai et al. [22] suggested a brand-new photonic spectrum-spreading phase-coding-based joint radar and communication (JRC) system operating at millimeter waves. The convergence of spectrum-spreading multiplexing and photonic microwave phase-coding techniques was the key to the system. By orthogonalizing the transmission data and the radar signal, spectrum-spreading multiplexing significantly reduces mutual interference. Furthermore, spectrum-spreading multiplexing enhance communication's anti-jamming, anti-noise performance as well as radar's peak sidelobe ratio. However, nonlinear distortions were produced by the radar's power amplifier (PA).

Bhargavi et al. [23] provided a technique for utilizing the H-Matrix as a detection and repair mechanism for several instances of double adjacent faults. Three distinct H-Matrix based codes with varying quantities of parity bits were presented and compared for various parameters using three different parity check matrices. To minimize the complexity of encoder and decoder circuits, the H-Matrix was built with fewer ones and less redundancies. However, it only focuses on correcting errors of up to 8 bits caused by double adjacent errors.

Prathyusha et al. [24] introduced a Matrix based technique for error detection and correction for memories, using power and stoppage to concentrate on the equality bits. The approach was effective in scenarios with limited speed and parity bits, employing both row and column decoders for address control during read and write operations. However, the system's complexity resulted in higher power consumption and longer processing delays.

Wang et al. [25] proposed deep error-correcting output codes, integrating deep, online, and ensemble learning. Incremental SVMs acted as weighted links between successive layers, with each ECOC module representing a layer pair. Class labels were used in pre-training to initialize the network, enabling effective performance, particularly for large-scale tasks. However, the method does not fully address scalability or computational efficiency for more complex datasets.

The literature review shows that prior works faced various limitations: [16] struggled to balance error correction and bandwidth efficiency; [17] underperformed across channel conditions; [18] failed to reduce residual errors; [19] lacked interoperability; [20] faced a redundancy-error correction trade-off; [21] inadequately handled burst errors and required better code mapping; [22] suffered from radar's power amplifier induced nonlinear distortions; [23] was limited to correcting small double-adjacent errors; [24] had high power use and delays due to complexity; and [25] lacked scalability and computational efficiency for large datasets. These gaps highlight the need for a novel approach to improve wireless communication effectiveness.

III. CAPE-FID

The amount of redundant bits added by the chosen coding scheme and the code length determine how well FEC approaches repair errors. The current FEC methods increase processing latency and delay by requiring more computational resources and a lower usable data rate. Hence, a novel

Convolution Assisted Polar Encoder with Flexible Iterative Decoding for Forward Error Correction in Wireless Communication Incorporated in Fpga

Convolution Assisted Polar Encoder with Flexible Iterative Decoding is developed to solve the drawbacks of the current FEC approaches used in wireless communication and enhance its communication efficiency by lowering latency and retransmission. The current methods employ longer block length codes in an attempt to increase error correction capacity and facilitate smooth system-to-system communication, but doing so unnecessarily takes up more bits overall, which ultimately reduces available bandwidth and the required useful data rate. Hence, a novel Polarized Convolutional Encoder is presented, in which the original data is first encoded using Adaptive Frozen Polar Coding in which the polarization transformation approach encodes the original input bits along with some added frozen bits only based on the length of the original data thereby avoiding unnecessary block lengths. After this, the transformed bits from the polar code undergo additional encoding using convolution codes that employ generator polynomials to control the output bit length and provide redundancy to the data stream, which facilitates error identification and correction at the receiver end. With this PCE, the unnecessary high redundancies added to the input data are reduced thus the useful data rate and effective usage of available bandwidth are improved.

Moreover, in congested network conditions, retransmission occurs due to packet loss, causing data loss at the receiver end. Existing error approaches, like CRC algorithm, which struggles to handle congestion, resulting in packet drop and lost data. Since different data sequences can produce the same CRC value in the existing CRC algorithms, the lower-degree polynomials used during divisional detection at the receiver end are more likely to collide. As a result, they mistakenly identify error-free data as erroneous, adding overhead to the error correction technique in the next step and affects the data transmission overall. Hence, a novel Flexible Turbo Decoder is presented at the receiver end, whereby Turbo coding schemes and modified Reed-Solomon (RS) code are used for error detection and correction. Once the data is received at the receiver end, the Reed-Solomon Euclid (RSE) Code is used to identify errors. This reduces the incorrect identification of error-free data and eliminates the need for error correction overhead by calculating syndromes based on the received codeword using the generator polynomial of the RS code and locating the error using the Euclidean algorithm. Following error detection, the error correction process uses Sequential Concatenated Turbo coding, which uses an iterative error correction approach in interaction with the error detection technique (RSE) in the previous stage and corrects the erroneous bits in the incoming bits in a sequential manner. This prevents packet loss and enhances the overall data transmission quality with improved network congestion handling. Finally, the error-corrected data is decoded using APC-Rpt.

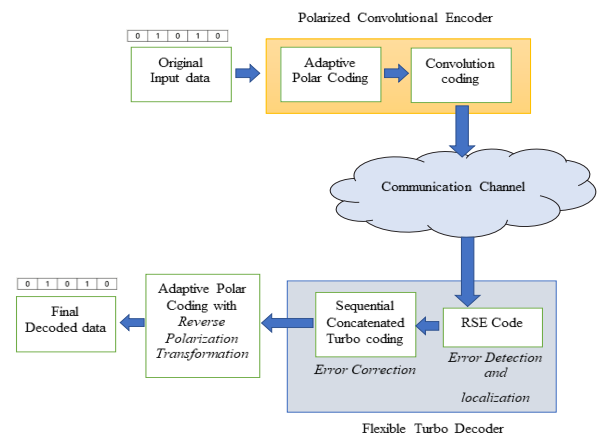


Fig. 1. Block Diagram of the proposed system

Fig. 1. displays the block diagram of the proposed model, which includes the sender side, the communication channel, and the receiver end. Initially, the original input data is sent to the PCE, which includes Adaptive polar coding and convolution coding for encoding the original data by this the unnecessary high redundancies added to the input data are reduced. Then, the encoded data are sent to the FTD placed at the receiver end through the communication channel for error detection and correction. First, the RSE Code is employed to detect errors followed by Sequential Concatenated Turbo coding for error correction. Finally, the error-corrected data is decoded using APC-Rpt.

A. Polarized Convolutional Encoder

In wireless communication, FEC is a technique that increases dependability by adding redundant information to the sent data. This way, faults can be corrected by receivers without requiring retransmission. This aids in the fight against problems like interference and noise in wireless channels. Initially, the original input data is sent to the PCE, which consists of AFPC and convolution codes, is used to reduce the unnecessary high redundancies added to the input data thus the useful data rate and effective usage of available bandwidth are improved.

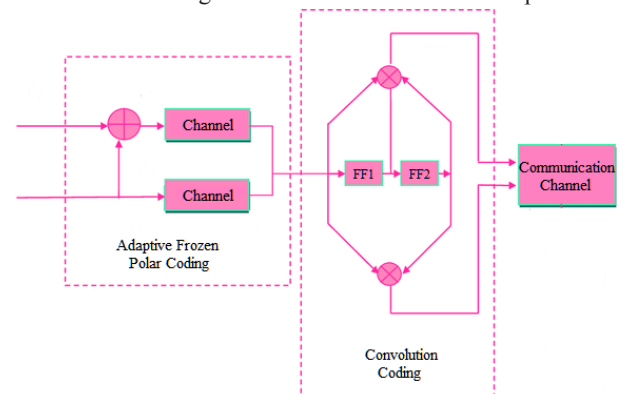


Fig. 2. Polarized Convolutional Encoder

Fig. 2. displays the PCE of the proposed system. The output from the polar encoder is applied to the convolutional encoder during data transmission to create an output at the transmitter

side. Additionally, the convolutional encoder's output is sent over a communication channel, and the recipient side decodes the data to obtain it.

First, the AFPC is used to encode the initial input data. Using a process called polar coding, information is converted into polarized bits, some of which are labelled as "frozen" and some as "information". By selectively activating or deactivating specific bits in the original data sequence, the polarization transformation essentially produces a polarized version of the data. The adaptive feature of frozen bit insertion is the main novelty in this case. The encoder dynamically decides the number of frozen bits based on the length of the original data, as opposed to employing a fixed set of frozen bits for all data lengths. The encoding process is represented by the equation (1).

$$x = u \times F + f \quad (1)$$

where, u is the original input data bits, x is the encoded polarized bits, f is the frozen bits, and F is the transformation matrix that polarizes the original data. This adaptive approach helps avoid unnecessary block lengths, optimizing the encoding process for each specific data sequence.

Following the polarization transformation, convolutional codes are used to further encode the polar code's transformed bits. Generator polynomials are used by convolutional codes to produce redundant bits and control the output bit stream's length. The encoder controls the amount of redundancy introduced to the data stream by varying the generating polynomials. The convolutional encoding process is represented as in equation (2).

$$c = x' \times G \quad (2)$$

where, x' is the encoded polarized bits after polar coding, c is the convolutional encoded bits, and G is the generator polynomial matrix. Convolutional coding adds redundancy, which facilitates error detection and correction at the receiver end. Convolutional codes are appropriate for wireless communication systems because of their well-known capacity to identify and correct errors in noisy communication channels. The overall encoding process involves both polar coding and convolutional coding is given in equation (3).

$$c = (u \times F + f) \times G \quad (3)$$

Thus, the PCE reduces excessive redundancy that are introduced to the input data that aren't necessary. This reduction in redundancies translates to improved useful data rate and more effective use of available bandwidth. By adding only necessary redundancy to the data stream, the communication system's overall efficiency is improved. Next, the encoded data is sent to the FTD at the receiver end through the communication channel for error detection and correction, as explained in section 3.2.

B. Flexible Turbo Decoder (FTD)

A FTD is a type of decoder designed to handle FEC in wireless communication systems, in which the error detection and correction are done with modified Reed-Solomon (RS) code and Turbo coding schemes. By employing a unified design, it is able to decode a variety of FEC codes, such as Turbo, Low-Density Parity-Check (LDPC), and Polar codes. In

Convolution Assisted Polar Encoder with Flexible Iterative Decoding for Forward Error Correction in Wireless Communication Incorporated in Fpga

various communication circumstances, this adaptability enables dependable and efficient data delivery.

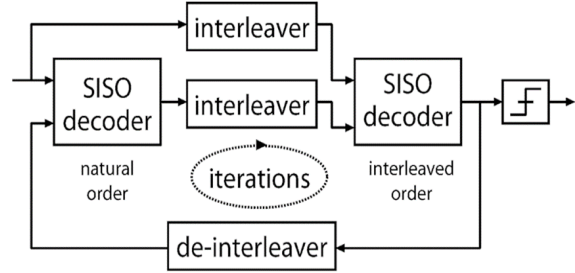


Fig. 3. Turbo Decoder

Fig. 3. displays the Turbo Decoder of the proposed system. The iterative decoding process is the essential component of a Turbo Decoder. This involves two iteratively exchanging Soft-Input Soft-Output (SISO) component decoders. Based on the output of the previous decoder and the received signal, each decoder improves its estimation of the transmitted data. With each iteration, this back-and-forth procedure improves the precision of the decoding.

When the transmitted data is received at the receiver end, it undergoes error detection using the RSE Code. The generating polynomial of the RS code is used by the RSE code to compute syndromes based on the received codeword. These syndromes basically serve as error detectors, indicating whether errors are present in the received data. The codeword $c(x)$ of RS code can be expressed mathematically as in equation (4).

$$c(x) = g(x) \times m(x) + e(x) \quad (4)$$

where, $g(x)$ is the generator polynomial of an RS code, $m(x)$ is the message polynomial, and $e(x)$ is the error polynomial. Based on the determined syndromes, the Euclidean method is then applied to detect and identify errors in the received data. The steps of the Euclidean algorithm involve repeated division with remainder until the remainder becomes zero. The receiver distinguish between data that has to be corrected and data that is error-free by successfully detecting errors at this stage. Mathematically, the algorithm is expressed as in equation (5).

$$gcd(a, b) = gcd(b, a \bmod b) \quad (5)$$

This lowers the possibility that error-free data will be mistakenly identified and also saves needless overhead in the subsequent error correction process.

After error detection using the RSE code, the next step is error correction using Sequential Concatenated Turbo Coding, an iterative error correction scheme that works in conjunction with the error detection method used in the previous stage (RSE). Using a step-by-step methodology, it corrects the erroneous bits in the incoming data stream, continuously improving the error correction procedure. Sequential Concatenated Turbo Coding can efficiently handle packet loss and enhance overall data transmission quality by iteratively refining error correction based on prior detections, especially in situations with network congestion or other unfavourable conditions.

After error correction is finished, APC-Rpt is used for the final decoding of the error-corrected data. The error-corrected

Convolution Assisted Polar Encoder with Flexible Iterative Decoding for Forward Error Correction in Wireless Communication Incorporated in Fpga

data is decoded using Adaptive Polar Coding, which modifies the decoding procedure according to the incoming signal's properties. In order to guarantee a correct reconstruction of the original data, reversed polarization transformation is used to reverse the polarization transformation that was applied during the encoding stage. This is expressed in equation (6).

$$u' = x' \times F' \quad (6)$$

The accurate decoding of the error-corrected data is ensured by this last decoding step, making it suitable for transmission or additional processing.

Thus, the Flexible Turbo Decoder scheme combines error detection with RSE Code, iterative error correction with Sequential Concatenated Turbo Coding, and final decoding with Adaptive Polar Coding to effectively handle errors and improve data transmission quality in wireless communication systems. Each step contributes to reducing latency, avoiding packet loss, and enhancing the overall efficiency of the communication system.



Fig. 4. Overall Architecture of the proposed model

Fig. 4. illustrates the overall architecture of the proposed model. The original data is first delivered to the PCE, which uses convolution coding and adaptive polar coding to encode the original data. The encoded data is then transmitted via the communication channel (Noisy channel) to the FTD, which is positioned at the receiver end, where RSE Code is used for error detection and Sequential Concatenated Turbo coding is used for error correction. Overall, the proposed error correction approach enhanced the useful data rate with available bandwidth, and data transmission efficiency, and also reduced retransmission.

IV. RESULT AND DISCUSSION

This section includes a thorough discussion of the implementation results, as well as the performance of the proposed CAPE-FID method, and a comparison to ensure that the proposed system successfully performed FEC in wireless communication incorporated in FPGA.

A. System configuration

The proposed system is simulated in MATLAB R2023a with the HDL Coder extension for FPGA implementation on a Xilinx Zynq-7000 series platform [30], with Windows 10 (64-bit) OS, Intel(R) Core (TM) i3-4130 CPU @ 3.40GHz processor, and 8GB RAM.

B. Performance metrics of the proposed system

In this section, a detailed explanation of the effectiveness of the suggested technique and the achieved outcome were

explained. Implementation results are analyzed in terms of throughput and latency, optimized for the target device's parallel processing and efficient logic resources.

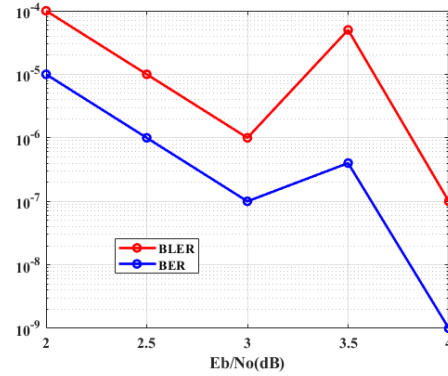


Fig. 5. BER, and BLER of the proposed system

The Bit Error Rate (BER) and Block Error Rate (BLER) performance of the proposed system, shown in Fig. 5. is evaluated across various SNR per bit (E_b/N_0) values. The BER achieves a maximum value of 10^{-5} at 2 dB and reaches a minimum of 10^{-9} at 4 dB, as the APC-Rpt enhances decoding efficiency, maximizing the likelihood of accurately recovering transmitted information and contributing to a lower BER. Meanwhile, the BLER reaches a maximum of 10^{-4} at 2 dB and decreases to a minimum of 10^{-7} at an SNR of 4 dB. This reduction in BLER at higher SNR values is largely due to the effectiveness of the RSE Code for error detection, which minimizes false positives by accurately identifying errors and reducing the likelihood of mistakenly flagging error-free blocks as erroneous.

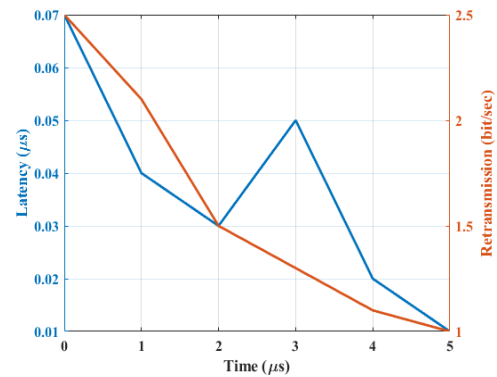


Fig. 6. Latency and retransmission of the proposed system

Fig. 6 illustrates the latency and retransmission performance of the proposed system across varying time periods. At a time interval of 5 μs , the latency reaches a minimum of 0.01 μs , while a maximum latency of 0.04 μs is observed at a time interval of 2 μs . This latency behavior of alternation between 2 to 4 μs is due to the system's CAPE-FID which optimizes both encoding and decoding processes, thereby reducing overall processing time, minimizing latency. Meanwhile, the retransmission rate peaks at 2.1 bits/sec at 1 μs , and decreases

to 1 bit/sec at $5 \mu\text{s}$. Sequential Concatenated Turbo Coding, employed for iterative error correction, allows the decoder to refine its estimation of the transmitted data progressively over multiple iterations. The reduction in retransmissions enhance the reliability of the received data and lead to more efficient use of the communication channel.

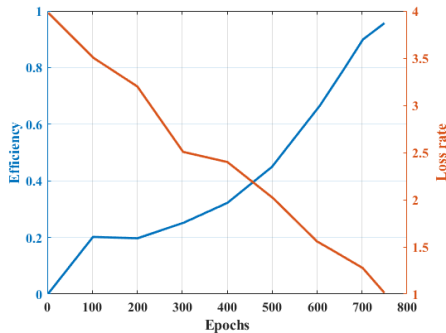


Fig. 7. Efficiency and loss rate of the proposed system

Fig. 7 presents the efficiency and loss of the proposed method across varying numbers of epochs ranging from 0 to 700. The system achieves a maximum efficiency of 0.94% at 700 epochs and a minimum efficiency of 0.45% at 500 epochs. This trend is attributed to CAPE-FID, which optimizes bandwidth utilization and by fine-tuning the encoding and decoding, it minimizes redundancy and increases efficiency. Meanwhile, on epoch 500, the system experiences a maximum loss rate of 2%, and a minimum of 1% at epoch 700. The loss reduction is primarily attributed to the PCE which enhance the error correction capabilities of the encoding process. Thus with increasing epochs, performance is increased with high efficiency and reduced loss rate.

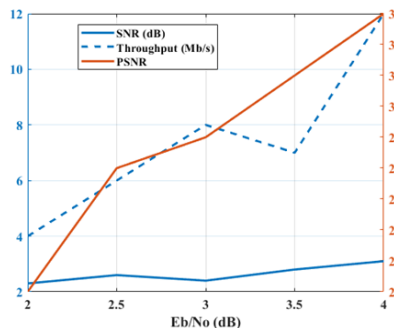


Fig. 8. SNR, PSNR and throughput of the proposed system

Fig. 8 analyzes the Signal-to-Noise Ratio (SNR), Peak Signal-to-Noise Ratio (PSNR) and throughput of the proposed method across varying SNR per bit values. The system achieves a maximum SNR of 3.1 dB at SNR per bit (E_b/N_0) of 4 dB, and a minimum of 2.3 dB at 2 dB. This improvement is due to the FTD, which minimizes packet loss and enhances data reliability, increasing SNR. In case of PSNR, a maximum PSNR of 33 is achieved at 4 dB SNR per bit, and a lower PSNR of 24 at 2 dB. This improved PSNR is largely recognized by modified RS codes in the proposed system which detects and corrects errors, reducing the error rate, resulting in higher PSNR values, hence a clearer and more accurate received signal.

On the other hand, at SNR per bit 2 dB, the system achieves a minimum throughput of 4 Mb/s, and at 4 dB, it reaches a maximum of 12 Mb/s. This increase in throughput is largely due to the APC-Rpt, which enhances the efficiency of decoding error-corrected data which minimizes retransmissions and optimizes data flow, directly contributing to higher throughput.

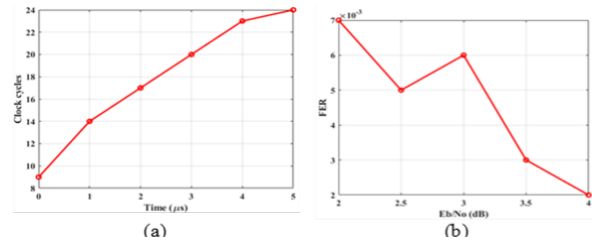


Fig. 9. (a) Clock cycle (b) FER of the proposed system

Fig. 9 (a) shows the clock cycle count of the proposed system corresponding to CPU execution time. At $5 \mu\text{s}$, the clock cycle count reaches a maximum of 24, while at $1 \mu\text{s}$, it decreases to 14 clock cycles. This reduction in clock cycles with shorter time periods is primarily due to the PCE and FTD, which optimize data processing that reduce the clock cycle and increases data encoding and decoding efficiency. The higher transistor switching speeds and more efficient data handling, contribute to an overall increase in processing speed, enabling the CPU to execute instructions with fewer clock cycles as time intervals decrease. Fig. 9(b) illustrates the Frame Error Rate (FER) performance of the proposed system across varying SNR per bit levels. At an SNR per bit of 4 dB, the FER reaches a minimum of 0.001, while at 2 dB, it achieves a maximum FER of 0.007. This reduction in FER at higher SNR levels, with a 2.5 to 3 dB variation is primarily due to the FTD, which effectively identify and correct errors in the received data, thus, decreasing the possible frame errors, even under challenging SNR conditions.

C. Comparison of Proposed Model with Previous Models

This section emphasizes the effectiveness of the proposed model by comparing it with the outcomes of existing techniques such as Polar [26] [27], Low-density parity check (LDPC), Tailbiting convolutional codes (TBCC), Turbo [26], Reed-Solomon codes (RS) [27], Unequal Error Protection scheme (UEP), Equal Error Protection (EEP) [28], Conventional Adaptive FEC, and Packet-Level FEC [29] to highlight its performance advantages over existing approaches.

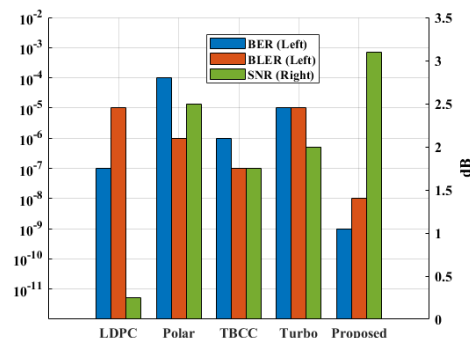


Fig. 10. Comparison of the BER, BLER and SNR

Convolution Assisted Polar Encoder with Flexible Iterative Decoding for Forward Error Correction in Wireless Communication Incorporated in Fpga

Fig. 10 shows the BER, BLER and SNR of the proposed model compared to the existing models. For the current models LDPC, Polar, and Turbo, the BER values are 10^{-7} , 10^{-6} , and 10^{-5} , respectively, the BLER values are 10^{-5} , 10^{-4} , 10^{-6} , and 10^{-5} and the corresponding SNR are 0.25dB, 2.5dB, 1.75dB, and 2dB. However, the proposed model achieves a lower BER of 10^{-9} , low BLER of 10^{-7} , and a high SNR of 3.1dB in comparison to the existing models.

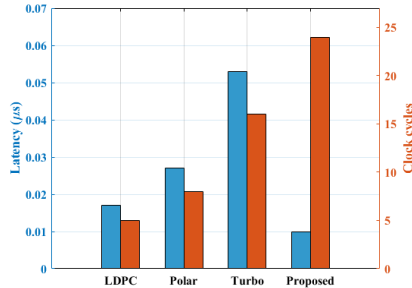


Fig. 11. Comparative analysis of the latency and clock cycle

Fig. 11 illustrates the comparison of the latency and clock cycle of the proposed model with existing models. The existing models such as LDPC, Polar, and Turbo achieves a latency of 0.017μs, 0.027μs, and 0.053μs and clock cycles of 5, 8, and 16 respectively. Compared with existing models the proposed model achieves a less latency of 0.01μs, and a maximum clock cycle of 24.

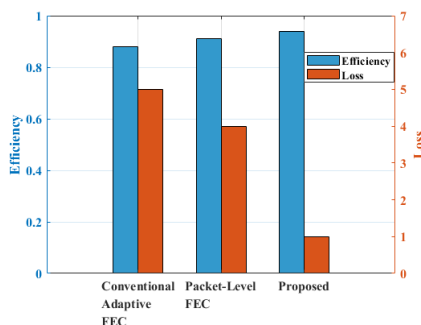


Fig. 2. Polarized Convolutional Encoder

The efficiency and loss of the suggested model is compared with existing models, as seen in Fig. 12. The efficiency of 0.88 and 0.91 and loss rates of 5% and 4% is achieved by the current models, Conventional Adaptive FEC and Packet-Level FEC. In comparison to current models, the proposed model attains a higher efficiency of 0.94 and a reduced 1% loss rate.

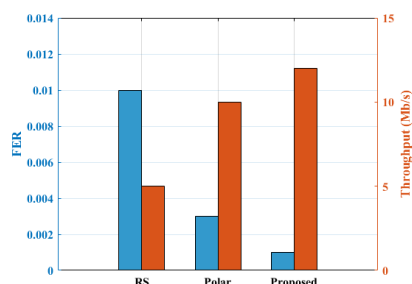


Fig. 12. Comparison analysis of the Efficiency and loss

Fig. 13 shows a comparison of the FER and throughput of the proposed model with the existing models. The FER for the current models such as RS and Polar are 1 and 0.003, respectively, and the corresponding throughput values of 5Mb/s and 10Mb/s. Compared with existing models, the proposed model achieves minimum FER of 0.001, and maximum throughput of 12Mb/s.

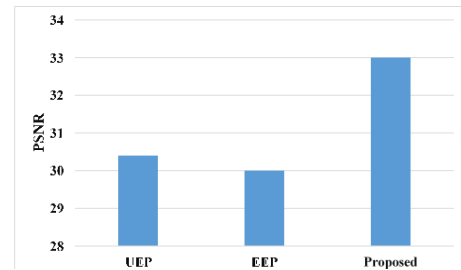


Fig. 14. Comparison of the PSNR of the proposed model

A comparison of the suggested model's PSNR with that of the current models is presented in Fig. 14. The current models' PSNR values are 30.4 and 30 for UEP and EEP. In contrast to current models, the suggested model attains a maximum PSNR value of 33.

V. CONCLUSION

In conclusion, the proposed CAPE-FID presented a promising solution to address the limitations of existing FEC techniques in wireless communication. By introducing a PCE and aFTD, the system effectively reduces latency, minimizes retransmission, and enhances communication efficiency. Through AFPCand Sequential Concatenated Turbo coding, the approach optimizes error detection and correction processes, thereby improving data transmission quality and network congestion handling. The utilization of Matlab modeling and FPGA implementation underscores the practicality and feasibility of the proposed mechanisms in real-world applications. Thus, the results of the simulation demonstrated that the proposed model has a low BER of 10^{-9} , BLER of 10^{-7} , FER of 0.001, Latency of 0.01μs, Loss of 1%, and high Clock cycles of 24, Efficiency of 0.94%, SNR of 3.1dB, PSNR of 33, and Throughput of 12Mb/s when compared to the previous models. Overall, this innovative approach offers significant advancements in wireless communication systems, adopting improved data rates, enhanced bandwidth utilization, and improved transmission efficiency in diverse network environments.

REFERENCES

- [1] M. Luvisotto, Z. Pang, and D. Dzong, "High-performance wireless networks for industrial control applications: New targets and feasibility," *Proc. IEEE*, vol. 107, no. 6, pp. 1074–1093, Jun. 2019, doi: 10.1109/JPROC.2019.2898993.
- [2] O. Seijo, J. A. López-Fernández, and I. Val, "w-SHARP: Implementation of a high-performance wireless time-sensitive network for low latency and ultra-low cycle time industrial applications," *IEEE Trans. Ind. Inform.*, vol. 17, no. 5, pp. 3651–3662, May 2021, doi: 10.1109/TII.2020.3007323.
- [3] G. J. Sutton, et al., "Enabling technologies for ultra-reliable and low latency communications: From PHY and MAC layer perspectives," *IEEE Commun. Surveys Tuts.*, vol. 21, no. 3, pp. 2488–2524, Jul.–Sep. 2019, doi: 10.1109/COMST.2019.2897800.

- [4] M. Zhan, Z. Pang, D. Dzong, and M. Xiao, "Channel coding for high performance wireless control in critical applications: Survey and analysis," *IEEE Access*, vol. 6, pp. 29 648–29 664, 2018, **doi:** 10.1109/ACCESS.2018.2842231.
- [5] F. Rezaei, D. Galappaththige, C. Tellambura, and S. Herath, "Coding Techniques for Backscatter Communications-A Contemporary Survey," *IEEE Communications Surveys & Tutorials*, 2023, **doi:** 10.1109/COMST.2023.3259224.
- [6] R. Lahman, and H. M. Kwon, "Low-Density Parity Check-Code in DVB- S2 versus Polar Code under SATCOM Fading," In *2023 IEEE Aerospace Conference*, pp. 1-6, 2023, March. IEEE, **doi:** 10.1109/AERO55745.2023.10115816.
- [7] A. S. Karar, A. R. E. Falou, J. M. H. Barakat, Z.N. Gürkan, and K. Zhong, "Recent Advances in Coherent Optical Communications for Short-Reach: Phase Retrieval Methods," In *Photonics*, vol. 10, no. 3, p. 308, 2023, March. MDPI, **doi:** 10.3390/photonics10030308.
- [8] D. Cavalcanti, J. Perez-Ramirez, M. M. Rashid, J. Fang, M. Galeev, and K. B. Stanton, "Extending accurate time distribution and timeliness capabilities over the air to enable future wireless industrial automation systems," *Proc. IEEE*, vol. 107, no. 6, pp. 1132–1152, Jun. 2019, **doi:** 10.1109/JPROC.2019.2903414.
- [9] V.K. Joshi, and T. Kavitha, "Efficient methods to improve power bandwidth measurement of ETSI low bandwidth digital mobile radio," *Measurement: Sensors*, vol. 27, p. 100 703, 2023, **doi:** 10.1016/j.measen.2023.100703.
- [10] A. Perdomo-Campos, I. Vega-González, and J. Ramírez-Beltrán, "ESP32 Based Low-Power and Low-Cost Wireless Sensor Network," In *The conference on Latin America Control Congress*, pp. 275–285, 2023. Springer, Cham, **doi:** 10.1007/978-3-031-26361-3_24.
- [11] M. A. El-Bendary, O. Al-Badry, and A. E. Abou-El. Azm, "Implementation of Novel Block and Convolutional Encoding Circuit Using FS-GDI," *IETE Journal of Research*, pp. 1–14, 2023, <http://dx.doi.org/10.1080/03772063.2023.2181876>.
- [12] X. Yu, H. Zhang, Z. Yang, Z. Lyu, H. Yang, Y. He, S. Liu, N. Li, O. Ozolins, X. Pang, and L. Zhang, "Photonic-wireless Communication and Sensing in the Terahertz Band," In *2023 Optical Fiber Communications Conference and Exhibition (OFC)*, pp. 1–3, 2023, March. IEEE, **doi:** 10.1364/OFC.2023.W4J.1.
- [13] Z. R. M. Hajiyat, A. Sali, M. Mokhtar, and F. Hashim, "Channel coding scheme for 5G mobile communication system for short length message transmission," *Wireless Pers. Commun.*, vol. 106, no. 2, pp. 377–400, 2019, **doi:** 10.1007/s11277-019-06167-7.
- [14] T. Xie, and J. Yu, "4Gbaud PS-16QAM D-Band Fiber-Wireless Transmission over 4.6 km by Using Balance Complex-Valued NN Equalizer with Random Oversampling," *Sensors*, vol. 23, no. 7, p. 3655, 2023, **doi:** 10.3390/s23073655.
- [15] O. Ferraz, *et al.*, "A survey on high-throughput non-binary LDPC decoders: ASIC, FPGA, and GPU architectures," *IEEE Commun. Surveys Tuts.*, vol. 24, no. 1, pp. 524–556, Jan.–Mar. 2022, <http://dx.doi.org/10.1109/COMST.2021.3126127>.
- [16] M. Vinodhini, N. S. Murty, and T. K. Ramesh, "Transient error correction coding scheme for reliable low power data link layer in NoC," *IEEE Access*, vol. 8, pp. 174 614–17 4628, 2020, **doi:** 10.1109/ACCESS.2020.3025770.
- [17] S. M. Abbas, M. Jalaleddine, and W. J. Gross, "High-throughput VLSI architecture for GRAND Markov order," In *2021 IEEE Workshop on Signal Processing Systems (SiPS)*, pp. 158–163, 2021, October. IEEE, **doi:** 10.1109/SiPS52927.2021.00036.
- [18] K. N. Kumar, N. A. Reddy, P. Shanmukh, and M. Vinodhini, "Matrix based Error Detection and Correction using Minimal Parity Bits for Memories," In *2020 IEEE International Conference on Distributed Computing, VLSI, Electrical Circuits and Robotics (DISCOVER)*, pp. 100–104, 2020, October. IEEE, **doi:** 10.1109/DISCOVER50404.2020.9278030.
- [19] Q. Zhang, S. Ibi, T. Takahashi, and H. Iwai, "Deep Unfolding-Aided Sum- Product Algorithm for Error Correction of CRC Coded Short Message," In *2022 Asia-Pacific Signal and Information Processing Association Annual Summit and Conference (APSIPA ASC)*, pp. 1737–1743, 2022, November. IEEE, **doi:** 10.23919/APSIPAASC55919.2022.9979875.
- [20] L. Wang, D. Song, F. Areces, and R.D. Wesel, "Achieving short-blocklength RCU bound via CRC list decoding of TCM with probabilistic shaping," In *ICC 2022-IEEE International Conference on Communications*, pp. 2906–2911, 2022, May. IEEE, **doi:** 10.1109/ICC45855.2022.9838498.
- [21] S. Matsenko, O. Borysenko, S. Spolitis, A. Udalcovs, L. Gegere, A. Krotov, O. Ozolins, and V. Bobrovs, "FPGA-Implemented Fractal Decoder with Forward Error Correction in Short-Reach Optical Interconnects," *Entropy*, vol. 24, no. 1, p. 122, 2022, **doi:** 10.3390/e24010122.
- [22] W. Bai, X. Zou, P. Li, J. Ye, Y. Yang, L. Yan, W. Pan, and L. Yan, "Photonic millimeter-wave joint radar communication system using spectrum-spreading phase-coding," *IEEE Transactions on Microwave Theory and Techniques*, vol. 70, no. 3, pp. 1552–1561, 2022, **doi:** 10.1109/TMTT.2021.3138069.
- [23] C. Bhargavi, D.V. Nishanth, P. Nikhita, and M. Vinodhini, "H-matrix based error correction codes for memory applications," In *2021 International Conference on Advances in Electrical, Computing, Communication and Sustainable Technologies (ICAECT)*, pp. 1–5, 2021, February. IEEE, **doi:** 10.1109/ICAECT49130.2021.9392574.
- [24] A. B. Prathyusha, and P. Sivadurgarao, "Optimizing Parity Bits for Error Detection and Correction for Memories Using Matrix based Technique," *Specialusis Ugdymas*, vol. 1, no. 43, pp. 10 597–10 607, 2022.
- [25] L. N. Wang, H. Wei, Y. Zheng, J. Dong, and G. Zhong, "Deep Error-Correcting Output Codes," *Algorithms*, vol. 16, no. 12, p. 555, 2023, **doi:** 10.3390/a16120555.
- [26] L. Fanari, E. Iradier, I. Bilbao, R. Cabrera, J. Montalban, P. Angueira, O. Seijo, and I. Val, "A survey on FEC techniques for industrial wireless communications," *IEEE Open Journal of the Industrial Electronics Society*, vol. 3, pp. 674–699, 2022, **doi:** 10.1109/OJIES.2022.3219607.
- [27] E.D. Spyrou, and V. Kappatos, "Application of Forward Error Correction (FEC) Codes in Wireless Acoustic Emission Structural Health Monitoring on Railway Infrastructures," *Infrastructures*, vol. 7, no. 3, p. 41, 2022, **doi:** 10.3390/infrastructures7030041.
- [28] A. Legrand, B. Macq, and C. De Vleeschouwer, "Forward error correction applied to jpeg-xs codestreams," In *2022 IEEE International Conference on Image Processing (ICIP)*, pp. 3723–3727, 2022, October. IEEE, **doi:** 10.1109/ICIP46576.2022.9897287.
- [29] A. Nafaa, T. Taleb, and L. Murphy, "Forward error correction strategies for media streaming over wireless networks," *IEEE Communications Magazine*, vol. 46, no. 1, pp. 72–79, 2008, **doi:** 10.1109/MCOM.2008.4427233.
- [30] https://github.com/git-hubuser2024/Source-code/tree/main/fec_fpga_zynq



T. Ranjitha Devi. I completed my Diploma and Graduation in Electronics and Communication Engineering (ECE), followed by a Post-Graduation in VLSI Design. Currently, I am pursuing my Ph.D. at GITAM University with research interests in VLSI and Communication Systems. I am working as an Assistant Professor at G. Pullaiah College of Engineering and Technology (GPCET), Kurnool, with 15 years of teaching experience.



C. Kamalanathan received his Ph.D. from Anna University, Chennai. Has good experience in teaching and research. Working as Associate professor at GITAM for many years in the Department of EECE at GST, Bengaluru. His research interests are in Wireless Communication and IoT. In addition to regular teaching, he has published a good number of research articles in peer-reviewed journals and guided research scholars. Also working as an Assistant Director in the Directorate of Academic Affairs.

Improving Technical Reviews: Metrics, Conditions for Quality, and Linguistic Considerations to Minimize Errors

Gabriella Tóth

Abstract—This paper presents a systematic methodology for evaluating the quality of technical reviews in software development, aiming to address the issue of ineffective reviews and their impact on product quality. The methodology, grounded in literature review, identifies key factors for high-quality reviews, including planning, reviewer commitment, and meaningful feedback. It emphasizes linguistic clarity and accuracy, drawing lessons from the Therac-25 case, where poor documentation contributed to fatal accidents. The paper analyzes specific linguistic issues like negative structures, passive voice, and terminology, highlighting their impact on comprehension. The methodology's effectiveness is validated through a Lean Six Sigma project in a large software development company, resulting in significant improvements. These include a 60% improvement in the Technical Review Quality KPI, a 29% reduction in customer-reported faults related to reviews, and the elimination of TL9000 non-conformities. This case study demonstrates the practical applicability of the framework and its potential for significant impact. The paper concludes by highlighting the importance of linguistic considerations in ensuring safer and more effective software products. Future research directions include extending the methodology to other types of artifacts and exploring textual analysis of reviewer feedback for deeper insights into the review process.

Index Terms—Documentation, Grammar, Manuals, Quality management, Reviews, Writing

I. INTRODUCTION

HAVING any kind of text or other artefacts checked by (at least) a second pair of eyes is a crucial step in the software development process. This does not happen differently in technical documents either. We can even say that for technical documents this is particularly critical as inaccuracies in content or unclear language and style can lead to misunderstandings, improper product use, system failures, or safety risks.

There are several terminologies used for the activity that include the observation and evaluation of software or documentation conducted by experts other than the author. In this paper I am using the term “technical review” in accordance with the terminology defined in [1]. According to this standard,

a technical review is a “systematic evaluation of a software product by a team of qualified personnel that examines the suitability of the software product for its intended use and identifies discrepancies from specifications and standards. Technical reviews may also provide recommendations of alternatives and examination of various alternatives”. To clarify further, a “software product” is: “(A) A complete set of computer programs, procedures, and associated documentation and data. (B) One or more of the individual items in (A)”. Therefore, technical documents are software products, and they are checked in technical reviews, as opposed to “audits”, “inspections” or “walk-throughs”.

The concept of “error” permeates various domains, often used interchangeably with terms like “fault,” “failure,” “defect,” or “mistake.” For the purposes of this paper, an error is defined as any incorrect, misleading, ambiguous, inconsistent, outdated, or missing information in a technical document that can negatively impact user understanding, product use, or downstream processes.

There are several examples in the literature that highlight the problem of errors escaping the review or testing phase in product or system development. Some of these are in industries where such errors might lead to very severe consequences (severe injuries or even loss of life).

One example is the case of the Therac-25 software-controlled radiation therapy machine where software errors and poor documentation played a role in tragic accidents [2]. Although reviews were not specifically mentioned but at the time of the events there was a lack of code or documentation review practices. We can safely assume that with conducting technical reviews the most significant errors could have been eliminated and the machine operators would have had proper support for operating the machine.

Therefore, ensuring the quality of technical reviews is essential. Organizations frequently report that their review processes are ineffective, expressing concern over the recurrence of undetected errors that are ultimately present in the published documentation. These documents may serve both internal stakeholders (such as software developers or testers) and external audiences (including client organizations or end users of technical products). While acknowledging these concerns is important, identifying concrete strategies for improving the effectiveness of technical reviews is even more critical.

Submitted on 29.01.2025.

Gabriella Tóth is with the Doctoral School of Linguistics, University of Debrecen; also with the Department of English Linguistics, Institute of English and American Studies, University of Debrecen; and with Nokia Solutions and Networks, Budapest, Hungary. (e-mail: gabriella.toth.szakal@gmail.com).

Similarly to other data-driven analysis and improvements facts are needed first to understand what the problem is. As in [3], "Measurement is the first step that leads to control and eventually to improvement. If you can't measure something, you can't understand it. If you can't understand it, you can't control it. If you can't control it, you can't improve it."

"Lord Kelvin's statement that "one does not understand what one cannot measure" is at least as true for software engineering as it is for any other engineering disciplines." [4].

Despite their critical role technical reviews have rarely been systematically measured, analyzed, or improved. This paper addresses this significant gap by introducing a novel framework for evaluating and enhancing the quality of technical reviews, specifically in the context of technical documentation. Through a Lean Six Sigma project, this research demonstrates how a data-driven approach can transform review processes, substantially reduce errors, and improve compliance with industry standards. The implementation of a Technical Review Quality KPI, combined with mindset-shifting initiatives like gamification, led to measurable improvements: a 60% increase in the composite review quality score and a 29% reduction in customer-reported faults. These outcomes underscore the practical value of applying engineering rigor to a process often overlooked in systematic quality improvement efforts. The findings and methodologies presented here offer actionable insights and a reusable approach for organizations aiming to elevate the quality of their review processes.

II. BACKGROUND AND PREVIOUS LITERATURE

One of the most renowned experts in information design, has extensively discussed the importance of technical reviews and quality assurance in documentation [5]. She argues that effective technical reviews should not only focus on the accuracy of content but also on the clarity and accessibility of language, ensuring documents are comprehensible to their target audience.

The following aspects are emphasized in [5] related to the quality of technical reviews: error detection rate and rate of missed errors, review coverage (the percentage of documents that have undergone formal review out of the total number of documents produced. It indicates how thoroughly the review process has been applied across a given documentation set), timeliness, reviewer expertise, consistency and compliance with standards and user feedback. However, while these factors underscore the multifaceted nature of review quality, [5] does not offer concrete guidance on how such elements should be systematically measured or operationalized. There is little elaboration on how metrics could be designed, collected, or applied in practice to support consistent evaluation across projects.

In addition to [5], several authors emphasized the importance of reviews for technical documents, similarly to inspections of software code that started already in the 1970s [6], [7].

There was a significant contribution to technical reviews in [8], particularly in the field of software engineering, with his

development of the Fagan inspection process in the 1970s. His method was one of the first formalized, structured approaches for software review processes and focused on improving product quality and reducing defects through systematic reviews.

The concept of defect categorization for code reviews was first included in [8] to some extent, but no details were presented. While this foundational process is relevant as context, the present study does not aim to elaborate on formal inspection methodologies. Instead, it builds on their legacy by focusing specifically on the measurement of review quality in technical documentation, an area where structured metrics are still underdeveloped.

A work in 2009, [9], looked at the literature of code review error classification and concluded that before their work there had not been a complete system of code review defect categorization and the focus had been on defect counts instead of type of errors. They defined a taxonomy of defects identified during code reviews, distinguishing between functional defects - such as logic, interface, or timing issues - and evolvability defects, which affect maintainability, readability, and clarity.

Measuring the effectiveness of technical reviews already appeared in [8], both for code and document reviews, and there have been several studies on the topic but they mostly concentrated on Defect Detection Efficiency or Defect Density (e.g. [8], [10], [11]).

I will also refer to previous literature as part of the following sections regarding specific topics.

A. Errors Escaping the Technical Review Phase

Before we can even start thinking about possible improvements, we must look at the following facts:

1. How many errors escape from the technical review phase?
2. What is the quality of technical reviews?

Errors that escape the technical review phase can be found when one analyzes the errors that are reported by users. For each of the reported errors, it must be defined which phase of the development process the error should have been found. The technical review phase should be considered a suitable point for addressing this.

In the rest of the paper, I will not cover further details of error analysis and categorization but will focus on technical review quality.

B. Defining Review Quality

It is a relatively challenging task to have a universally agreed measurement of the quality of technical reviews. It might be impacted by the industry, technology, type of artefacts (technical documents, software code, project plans, test plans, etc.) or other factors.

I am using "effectiveness" instead of "quality" for the metrics quoted in earlier literature (Defect Detection Efficiency or Defect Density) as I consider these basic metrics to reflect the observable outcomes or measurable results of the review

process - what could be termed the “facts” of the review but not the quality of the review.

Defect Detection Efficiency is usually defined as the ratio of defects found during reviews and the total number of defects in the reviewed software product. Defect Density is usually measured per 1000 lines of code, reviewing hour or document page. In our experience, these do not show how well the review “performs” in reducing the errors escaping from the review phase. First, it is difficult to measure the total number of defects (are those the ones found later before delivery or after delivery, by the user?), also argued in [4] and second, the increase in defect density alone does not necessarily result in better reviews (it is also important to check what kind of defects are found). However, these metrics are useful for following up and later planning the needed efforts for reviews, which is an important aspect for Project Managers.

It was proposed in [4] that review time and duration are also considered in the measurement, as part of the cost: engineering time benefits and costs. This study did not give exact definitions but in the current context I would consider as following. Reviewer time is the total time invested by reviewers in planning, preparation, and conducting document reviews. Engineering time benefits are the estimated time savings in activities like reduced rework, fewer errors, or faster development cycles.

In [12], reviewer expertise was addressed through perspective-based reviews, where individuals evaluate the document from roles aligned with their specific competence—such as user, developer, or tester perspectives—to improve review effectiveness.

While the above studies discussed quality metrics for reviews, none of them proposed a systematic way of measurement that would include aspects other than defect density.

My objective in this paper is to show how such a measurement can be set up for documents that are created during software development.

C. Recent Developments in Automated and AI-Assisted Technical Reviews

In the past decade, technical documentation review practices have evolved significantly through the adoption of automation and documentation-as-code (DaC) workflows. Traditional peer review methods—where subject matter experts manually assess user guides or interface documents—are increasingly supplemented by CI/CD-integrated pipelines and version-controlled documentation. In this approach, documentation is written in lightweight markup languages (e.g., AsciiDoc), stored in Git, and reviewed via pull requests, similar to software code. In [13] the authors demonstrated a DaC-based pipeline for managing Interface Control Documents (ICDs), integrating CI tools that validate document structure, automate glossary generation, and enforce writing style rules through custom quality gates. Their work showcases how such pipelines can reduce documentation errors, ensure uniformity, and improve

maintainability in complex engineering environments. These developments have improved traceability and efficiency, especially in agile and DevOps contexts. Review feedback is increasingly structured and tracked within issue systems, and changes can be linked directly to product development cycles. At the same time, Natural Language Processing (NLP) techniques are increasingly applied to support high-level documentation review tasks such as detecting vagueness, redundancy, or omissions in technical texts. A set of machine learning and deep learning models was developed and evaluated - including BERT - to automatically identify five types of “documentation smells” in API references, such as overly vague descriptions, overly technical detail, and fragmented or tangled content [14]. While these tools can significantly reduce reviewer workload and enhance consistency, the authors emphasize that they are best used in combination with human judgment, especially when domain-specific knowledge is required.

It is important to note that the present study does not focus on these recent technological developments. The analysis and findings reflect the current documentation review practices implemented within the specific organizational context in which the study was conducted. As such, the study concentrates on established manual review methods, as supported by existing literature. Nonetheless, the integration of automated and AI-assisted review processes represents an important and timely area for future research, particularly as these tools continue to evolve and gain adoption across technical domains.

III. CONDITIONS FOR A GOOD TECHNICAL REVIEW

It is important to explain the scenario for which the measurement is described. In this study, the term “technical documents” refer specifically to product-related documentation authored by technical writers—such as user guides, operating manuals, and reference documentation—that support key product milestones. The measurement framework focuses on technical reviews of such documents, typically organized and facilitated by technical writers, and conducted with input from subject matter experts.

There is no golden rule for defining what makes a technical review good or of expected quality. One approach is to select the key factors that need to be satisfied and measure the quality based on how much those conditions are fulfilled. Below is a list of conditions that proved to be the most relevant within the scope of this study. The conditions are grouped into pre-review and review conditions.

A. Pre-Review Conditions

Table I includes the criteria that are included in the pre-review requirements.

TABLE I
PRE-REVIEW REQUIREMENTS

#	Requirement
1	Technical reviews are part of the project plan, with clearly defined roles and responsibilities (like moderator, reviewer, approver)
2	Mandatory and optional review participants are selected
3	There are at least two mandatory reviewers
4	(At least) the mandatory reviewers have committed resources for participation

The prerequisite for the pre-review conditions is that the technical writing team has a plan for all documents that need to be reviewed for a milestone. They need to provide the expected schedule and the proposed review participants (including the moderator and approver) based on the scope and topic of the documents. The latter is important in defining the responsibility of technical writers, who must share it with the reviewers in the development and product management teams. The quality of technical reviews does not only depend on the reviewers; reviewers can only commit to and participate in reviews if they receive a proper plan.

We defined at least two mandatory reviewers (in addition to the author) as a pre-review requirement. There are different views on how many reviewers are needed, some look at it from the expertise point of view, some from effort/cost point of view. In [10] several studies were included that had mentioned 3-5 reviewers as optimal, while they also cite an experiment that concluded that reducing the number of reviewers to two may significantly reduce effort without increasing review time or reducing effectiveness.

In addition to providing a plan for reviewers to estimate the needed effort for the review and to reserve the time to be able to participate in the reviews, there are best practices that can enhance the participation rate of reviewers. It was discussed in [15] how positive framing can increase the effectiveness of technical reviews. She explains that “The technical review process often suffers from a mismatch of priorities and expectations within a cross-functional team”. This often results in reviewers not considering participating in technical reviews important enough. She argues that document authors “need to reframe the technical review and “sell” the frame to engage reviewers”.

There were many studies cited in [10] that emphasized the need to “advertise” reviews and show that reviews work and that errors are inevitable, so it is completely fine to search for and find errors.

While this is an important social aspect to planning technical reviews, priority should be given to the proper planning and execution of reviews based on institutionalized processes and quality systems.

B. Review Conditions

Table II includes the review requirements.

TABLE II
REVIEW REQUIREMENTS

#	Requirement
1	At least all mandatory reviewers participate ^a in the review process
2	Comments and findings are provided to the author in the given timeframe
3	At least all mandatory reviewers provide comments or findings ^b
4	At least one of the findings is related to the technical content of the document
5	The approver approves or rejects the document review based on the participants’ findings
6	The author provides the needed corrections based on the findings

^a Participation means that the reviewer either approves the document as such or provide comments or findings, out of which at least one is related to the technical content of the document.

^b Providing comments or findings include both the listing of findings and explicitly stating that the document is approved as such.

C. Comments, Findings and Review Decisions

In the proposed framework, a review is considered of good quality if all mandatory reviewers provide comments or findings, and the approver makes a documented decision (either approval or rejection) based on the participants’ input. This means that in some cases reviewers can fulfil their roles as reviewers by simply stating that the document is good as such, and it is approved by the reviewer without any other comment. This is completely fine as it is not a “silent” approval (when the reviewer simply does not provide any comment but does not state his approval for the document as such). The practice of “silent” approvals is not preferred as it remains ambiguous and gives no added value to the review.

However, accepting the document as such might also have some hidden risks of relying on other reviewers’ opinions or “votes”.

In [16] research was described in which they analyzed software code reviews and how the decision “votes” for accepting or rejecting a software patch of the previous reviewers might influence the vote of a reviewer. He describes several other factors (like relationship between reviewers, status, seniority etc.), claiming that such review dynamics might cause more error-prone code: “we find that the proportion of reviewers who provided a vote consistent with prior reviewers is significantly associated with the defect-proneness of a patch” [16]. Although the study was conducted for code reviews, the general statements about why code reviews are important and how code reviews are currently carried out using collaborative tools that make all reviewers’ comments and votes visible to all other reviewers are all valid for technical document reviews as well.

However, influence of reviewer interaction dynamics on decision-making is not within the scope of this paper, but I wanted to show what other perspectives could also be considered when creating a review quality measurement concept.

D. Type of Comments as an Indicator for Review Quality

This condition may be considered one of the more debated factors influencing review quality. A common assumption is that comments addressing technical content carry greater value than those focused on language-related issues. However, this perspective is not universally valid. Language-related feedback can be equally important, particularly when it affects clarity, usability, and the accurate interpretation of information. It is also important to recognize that the relative importance of comment types may vary depending on the nature and purpose of the documentation under review. For reference-type documents for which the template is very controlled or there is little running text or graphics, this condition works well. On the other hand, for operating manuals, product descriptions or other non-referential documents, this condition is not one to be used.

Language-related comments can significantly influence the overall quality of a review. If aspects such as language, style, or layout are overlooked—or if reviewers notice issues but fail to provide feedback to the author—the intended message of technically accurate content may still be misinterpreted or overlooked by the user.

A detailed taxonomy for document issues based on a wide empirical study was defined in [17]. They listed five main categories (Completeness, Up-to-dateness, Usability, Tool-related, and Readability) and several sub-categories for issues collected from comments from various sources (i.e., emails, issues and pull requests of open-source projects, and Stack Overflow threads). Although these sources might considerably differ from the technical documents the current paper describes a framework for, the taxonomy in [17] supports the relevance of linguistic considerations and language issues.

Among the categories defined in [17], Readability—which includes subtypes such as too abstract, too technical, too verbose/noisy, and simple typos—is particularly relevant to the current paper’s focus on review quality from a linguistic perspective. While Readability was not identified as the most frequent issue by the original authors, it is emphasized here due to its central role in evaluating the clarity and accessibility of reviewed technical content.

There was no further analysis on the exact issue types but the finding that “issues related to lack of clarity represented more than half (55%) of these problems” highlights the importance of including linguistic considerations into technical reviews.

IV. LINGUISTIC CONSIDERATIONS DURING TECHNICAL REVIEWS

In this section I will describe why the language related reviewer perspective during reviews can help with identifying significant problems before certain technical documents are published for users. As I elaborated in the previous section, linguistic considerations do not have the same priority for different documents. Documents that provide step-by-step instructions for operating or troubleshooting technical

equipment must include the review of the language, while for example, product descriptions or parameter lists might not need such reviewing aspect.

There might be several linguistic structures where wrong language use can cause misunderstanding for users. I selected three structures that I consider critical for understanding technical documentation.

1. Negative structures may be wrongly used and thus, it becomes confusing whether something is stated or negated

2. The usage of passive and active voice substantially impacts the understanding of the text’s intention

3. Incorrect use of terminology might cause the misunderstanding of the technical content

The earlier cited work, [2], summarized the findings of the tragic Therac-25 case, where part of the root causes was poor documentation. The study did not go into the details of the documentation problems, and the original machine documentation is not available for researchers. The only root causes mentioned related to documentation were vague error messages, insufficient instructions for what to do in situations when an error occurs. Based on this, I can only assume that the listed language issues might have played a part in this specific case.

A. Problems with Negative Structures

Structures with “does not”, “cannot”, “will not”, “is not”, “was/were not” are the most unambiguous, unless the contracted forms (“doesn’t”, “can’t”, “won’t”, “isn’t”, “wasn’t”, “weren’t”) are used. Contracted forms can be easily misinterpreted by non-native speakers of English or users less familiar with the language. To the best of my knowledge, there are currently no academic studies that directly investigate the impact of contracted forms on comprehension in written technical texts. However, numerous professional style guides and technical writing resources discourage the use of contractions, particularly in contexts aimed at international or non-native audiences (for example, [18]). The guide emphasizes that technical texts should prioritize clarity and minimize potential misunderstanding for global readers. This caution is well-founded, as many contracted forms are ambiguous. For instance: “I’d” could mean “I would” or “I had”, “Who’s” could mean “who is”, “who has”, or even be misinterpreted as “who was” depending on context.

An additional problem might arise when contracted (or even uncontracted) forms are used together with a negative word in the same sentence. Or, vice versa, if there is no negative structure in the sentence, however, there is another expression that has a negative meaning.

The below examples are samples from a technical text corpus used by the Sketch Engine online tool [19]. It is important to note that the publicly available corpora that the tool uses are final versions of texts, meaning that (ideally) they already went through some kind of review. However, the corpus might also include technical text not from product documents but from online guidelines or instructions that might not go through any reviews. Therefore, they will only show the structures

described, but not necessarily ones that are unclear or confusing.

Table III shows the function of the highlighted items.

TABLE III
EXAMPLES FOR PROBLEMS WITH NEGATIVE STRUCTURES

Example sentence	Type of negation
"But it is not a rare feature in desktop CPUs anymore."	DIRECT NEGATION + INDIRECT NEGATOR
"Installed despite Kaspersky AV not liking the set up file."	CONSESSION + DIRECT NEGATION
"The unidirectional level converter is in there, but the pins do not match."	CONTRAST + DIRECT NEGATION
"For instance, sending to anything not in the mozillaessaging.com domain will fail ."	DIRECT NEGATION + IMPLIED NEGATION
"Our tests in Chicago didn't show great performance from 'nationwide' 5G."	CONTRACTED DIRECT NEGATION
"Very important: Don't ever try to call any method before the call of InitializeComponent method."	CONTRACTED DIRECT NEGATION + INTENSIFIER + NEGATION-SENSITIVE DETERMINER
"This is frequently caused by the chroot directory not being a mountpoint when the chroot is entered."	DIRECT NEGATION
"Make sure you have 50% or more battery life otherwise the update won't download and install."	DIRECT NEGATION

Examples were manually selected from a concordance using the English Web 2021 (enTenTen21) public corpus, with filtering for topic=Technology & IT, word/lemma="not" and sorting based on GDEX, in sentence format. GDEX provides Good Dictionary Examples, identifying identify sentences which are easy to understand and illustrative enough [19].

The above examples show that having more than one DIRECT NEGATOR in a sentence makes it more complex and more difficult to understand it. The sentences themselves are not incorrect but would not be recommended to be used in English for Specific Purposes, in this case, for technical texts.

B. Problems with Active/Passive Voice

The incorrect usage of active vs. passive voice in technical texts might create unclarity and might result in serious consequences in situations when, for example, the text provides instructions for carrying out a specific task. Improper use of passive voice in technical documentation can reduce clarity, especially when the actor performing an action is not specified. This issue is particularly problematic in requirements documents, where missing agents can lead to ambiguity in system behavior descriptions. The author of [20] highlights that passive voice often signals omitted information in use case scenarios, as it obscures who is responsible for performing actions. This becomes a serious problem in early project stages, where lacking detail can hinder requirements analysis. This work supports the view that clarifying agency through explicit language—including avoiding unnecessary passive voice—contributes to more complete and interpretable documentation.

Table IV shows examples with the verb "install", which is a frequent action in technical documents like installation or operating manuals.

TABLE IV
EXAMPLES FOR PROBLEMS WITH ACTIVE/PASSIVE VOICE

Example sentence	Issue
PICSRules: Specifies an interchange format for filtering preferences, so that preferences can be easily installed or sent to search engines.	Who is responsible for the installation? End user or software system or an administrator?
So P3P could be installed on major Server implementations like Apache, Jigsaw, Netscape-Server or Internet Information Server from Microsoft.	Who is responsible for the installation? Server administrator or developers or automated tools?
Once the real font has been successfully downloaded and temporarily installed , it replaces the temporary font, hopefully without the need to reflow.	Who is responsible for the installation? A user or an automated system or a specific application?
Magnetic flow meters can be installed in horizontal or vertical piping so long as the pipe remains full at the point of measurement.	The "actor" is not specified, however, if the focus is on the flow meter's installation capabilities rather than on the installer, this may suffice for high-level overviews or specifications.
After the extension is installed , you will find a new section has been added to the permissions window.	Who is responsible for the installation? The user or is the extension pre-installed by an administrator or automated process? The sentence includes 2 passive structures. The second one ("a new section has been added") is less problematic because the focus is on the result (the new section in the permissions window), and the actor who added it is irrelevant to the reader's understanding of the outcome.

C. Problems with Terminology

The whole idea of terminology is to make sure that there is only one specific term used for a concept. If wrong terms are used in technical documents or if the terms are ambiguous (for example, because the abbreviation is not resolved, and the document does not include a glossary or index), the intention of the product specification or operating instructions might undermine the specificity of the technical text.

A thoroughly studied major incident included a basic terminological error. While the 1999 NASA Mars Climate Orbiter case was more of a software issue related to unit conversion, it is closely tied to documentation and communication errors between teams. NASA's Mars Climate Orbiter was lost due to a mismatch between metric and imperial units in the software documentation. One team used the imperial system while another used the metric system, and this was not clearly communicated in the software documentation. This ambiguity in the language and instructions led to a navigational error, causing the spacecraft to enter Mars' atmosphere at the wrong altitude, leading to its destruction [21].

Framework for Intrusion Detection in IoT Networks: Dataset Design and Machine Learning Analysis

In Table V there are additional examples of hypothetical scenarios developed to illustrate common themes in wrong terminology usage. These are not direct citations but inspired by text I came across during my work.

TABLE V
EXAMPLES FOR PROBLEMS WITH TERMINOLOGY

Issue type	Example	Issue
Unclarity about whether the two terms are the same or not	"Refer to the <i>firmware update guide for instructions.</i> " vs. "Download the latest driver update to ensure compatibility."	Is the "firmware update" the same as the "driver update," or are they separate processes?
	"The configuration file is located in the /config directory." vs. "The settings file can be found in the /config folder."	Are the "configuration file" and "settings file" the same, or are they different?
Inconsistent use of terms	"Press the <i>power button</i> to turn the device on." vs. "Hold the <i>on/off switch</i> to start the device."	Mixing "power button" and "on/off switch" for the same component may confuse users about whether they are referring to different physical controls.
	"Upload the document to the <i>cloud storage</i> service." vs. "Save the file to the <i>online repository.</i> "	Using "cloud storage service" and "online repository" inconsistently without specifying if they are the same or different systems can create unnecessary complexity.

D. Language-Focused Review Comments and Their Automation Potential

Language issues—such as ambiguity, excessive passive voice, or inconsistent terminology—pose significant risks in international and regulated documentation. Rule-based tools, including classical grammar checkers and language profilers, are effective at identifying superficial issues like punctuation, agreement errors, and overly complex sentences. For example, a systematic review of automated grammar checking tools notes these systems can flag syntax errors, missing articles, and punctuation problems, but often struggle with deeper stylistic nuances or domain-specific inconsistencies [22]. ML-based approaches, on the other hand, have begun to demonstrate greater sophistication. For instance, tools like GrammarTagger use transformer-based models to profile grammatical features across languages, while recent work in "documentation smells" detection employs pretrained models like BERT to identify vague phrasing, structural omissions, and incoherent logic in API documentation [14]. These AI-enhanced methods offer promise in highlighting subtle quality issues that rule-based tools typically miss. Nevertheless, human review remains essential to interpret context, validate technical relevance, and resolve ambiguities that automation alone cannot adequately address.

It should be emphasized, however, that the present study does not analyze or implement these automated approaches. The

language-related observations and categorizations discussed in this chapter are based on manual review practices documented within the organizational setting of the current research. While this provides a grounded understanding of real-world review behavior, the exploration of automated or AI-assisted solutions for language-related feedback remains outside the present scope. Future research will seek to incorporate and assess such technologies more systematically.

V. DEFINING A REVIEW QUALITY MEASUREMENT FRAMEWORK

A. Designing a Meaningful Metric

After defining the criteria for good quality technical reviews, a usable and meaningful metric can be designed so that the As-Is and To-Be situations are identified.

This paper describes an option where:

- Each document is evaluated as compliant or non-compliant with the criteria
- The milestone- or project-based document sets are evaluated based on the document-level compliance

Tables VI and VII include illustrative examples of how the metric calculation is done.

TABLE VI
DOCUMENT-LEVEL COMPLIANCE SCORING

Criterion	Doc1	Doc2	Doc3
Number of mandatory reviewers defined and committed	4	2	1
Number of mandatory reviewers providing valid comments	3	2	1
At least 2 mandatory reviewers defined and committed before the start of the review: YES/NO	YES	YES	NO
Percentage of mandatory reviewers who gave valid comments during the reviews or in the mail review periods: %	75%	100%	100%
Were comments late? YES/NO fill in only if column F is <100%	NO	NO	NO
Overall compliance YES/NO	NO	YES	NO

In Table VII, you can see additional calculations on top of the overall score. It shows that once you have the figures, you can yourself define which additional factor you would like to focus on or bring attention to the development teams or management. In this case we used the number of reviewers and the timeliness of comments as these two were the fundamental issues in our assumptions.

TABLE VII
DOCUMENT SET LEVEL SCORING

DOCUMENT SET percentage of documents meeting all criteria for good quality technical reviews		33%
percentage of documents with only 1 defined and committed reviewer		33%
percentage of deviations due to late comments		0%

The examples show a scenario where the document set has a Technical Review Quality of 33%. Once you have the overall score, it is up to your project, team or organization to set a target score and to define actions for what should happen when the score goes below the target or if there is a negative trend over several months.

B. Automation of Data Collection

Similarly to any kind of metric, in ideal case, all data about the review quality conditions is automatically collected and analyzed. For most of them it is possible and advisable to find the right toolkit that can handle all aspects of reviews: a database of review items with metadata, the draft documents themselves, the comments, findings and decisions.

The illustrative example in Tables VI and VII is entered and calculated simply in a spreadsheet, with manual data entries. This is the easiest way to collect data but in case of complex products with even hundreds of artefacts to review spreadsheets are not an option. As the number of reviewed items and reviewers increases, manual data collection becomes time-consuming and error-prone, raising concerns about scalability and reliability.

Moreover, reliance on manual data entry can introduce bias, both in what data is recorded and how it is interpreted. Reviewers may inadvertently omit information, apply inconsistent labeling, or interpret review outcomes differently depending on their experience or focus. These inconsistencies can distort the assessment of review quality and undermine efforts to compare performance across projects or teams.

To mitigate these issues, future implementations of the framework should explore semi-automated or fully automated solutions that can extract review data directly from collaborative platforms such as Git repositories, issue trackers, or review tools. For example, metadata about review duration, number of reviewers, and comment volumes can be parsed automatically from pull requests or change logs. Additionally, Natural Language Processing (NLP) techniques could be integrated to categorize review comments (e.g., distinguishing between content-related and language-related feedback) and to identify recurring issues such as ambiguity, terminology misuse, or incomplete instructions. Such methods would significantly enhance data completeness, reduce manual effort, and enable consistent and scalable quality measurement across documentation projects.

C. Addressing Quality Deviations

Measuring the quality of technical reviews is only the first step: gathering data as opposed to having a “feeling” about how technical reviews go. When you have facts in detail, you can easily pinpoint concrete problems towards management or the stakeholders.

In the following sections, I will present possible actions you can take to improve the quality of the reviews, and eventually, to reduce the errors found by the user.

Our approach was to first make the defined **quality conditions** into **criteria lists** or **checklists**.



Fig. 1. Quality conditions made into criteria lists

As we earlier defined the most crucial factors that impact the quality and efficiency of technical reviews, the next step is to make sure that these are checked early enough to avoid deviations afterwards by intervening as quickly as possible.

Checklists are a good means that technical writers can use before and after reviews. Of course, having a checklist does not necessarily mean any change unless there is a clear agreement and guidelines for its usage. To further strengthen the integration of linguistic feedback into structured review practices, selected categories from the documentation issue taxonomy proposed in [17] - such as Readability and Completeness - could be incorporated directly into the checklist and quality conditions used during reviews. While these categories were not part of the original implementation, they align well with the current framework’s emphasis on clarity, accuracy, and consistency.

In our case a drastic action was that technical reviews cannot even start before the criteria on the entry checklist are fulfilled. Not starting the reviews on time might result in delays with the approval of the documents and eventually, delays in the documentation delivery and ultimately, delays in software deliveries.

Similarly, reviews cannot be closed before the review exit criteria are fulfilled. Again, this might result in delays with the projects, causing huge problems for the product delivery.

An important change this brought was in the mindset. It used to be a general understanding that problems with the reviews should be owned and managed by technical writers or the technical reviewers, depending on which organization you work in. It was very rare when both teams understood that this is their common responsibility. Technical writers cannot claim that everything is on the shoulder of the reviewers, and they depend on them, while reviewers cannot claim either that the reviews are the sole responsibility of the technical writer team, and they will attend and contribute only if they have time.

With these changes they both understand that it is a joint effort that will bring success or problems to both.

Another improvement action that can be taken is to build the technical review quality into the project completion criteria. As with any other milestone criterion that needs to be passed in order to be able to decide about milestone approval, technical review quality can be a checkpoint, and it provides an opportunity for the Quality Manager and Program Manager not to approve the milestone due to problems with review quality. This includes any delays as in most cases when there is a problem with review quality, the result is some kind of delay, and this is eventually shown in the milestone approval.

As one of the entry criteria is that all the defined mandatory reviewers are committed to taking part in the reviews, there must be a proper plan based on the input from the document authors (i.e. Technical Writers). In addition to the planning of the right mandatory reviewers and the timeframe for the review, effort estimation must be done per reviewer to make sure they will really have time allocated to review the assigned documents. Without such a good plan it is not realistic to expect a good-quality review.

VI. CASE STUDY

A. Background

The measurement concept and improvement actions described earlier in this paper were implemented in the documentation department of a large software development company as part of a Lean Six Sigma project. This initiative focused on addressing the consequences of issues with technical document reviews, which included:

- 1) Customer-reported faults caused by inadequate reviews. These refer to documentation errors reported by customers that should have been detected during the internal review process. Examples include ambiguous instructions, incorrect parameter names, outdated references, or mismatches between the documented behavior and the actual product.
- 2) Program milestone rejections or concessions linked to review problems. In several cases, software program milestones were delayed or conditionally accepted because documentation did not meet the required quality standards. Issues were considered critical enough to block or delay software release approvals until documentation was revised.
- 3) Non-conformities to TL9000 requirements due to review-related issues. These non-conformities were identified during quality audits and related specifically to deficiencies in the review process, such as missing review plans, lack of documented reviewer roles, or absence of traceability between review findings and corrective actions.

This case study outlines the structured approach taken, the execution of improvement actions, and the tangible results achieved. While certain aspects such as reviewer engagement were still assessed subjectively due to the early phase of gamification features, the initiative aimed to increase reviewer participation - measured informally through the number of active reviewers and volume of review comments.

B. Approach and Execution

1) Collection of Customer Faults Data

Data on customer-reported faults related to variations in technical review quality was gathered systematically, providing a foundation for identifying improvement areas.

2) Design and Implementation of a New KPI

A novel metric, the Technical Review Quality KPI, was designed and introduced. This KPI, along with a detailed methodology, enabled the department to measure and monitor review quality objectively.

3) Root Cause Analysis

A thorough Root Cause Analysis (RCA) was conducted to pinpoint recurring problems in the review process. This analysis identified gaps in reviewer engagement, entry and exit criteria, and adherence to best practices.

4) Education and Training

Dedicated training sessions were conducted for technical writers and expert review teams. These sessions focused on enhancing their understanding of the review process, setting clear expectations, and emphasizing their roles in achieving high-quality reviews.

5) Quality Checkpoints Implementation

Several review-related quality checkpoints were introduced, including:

- Entry and Exit Criteria Checklist: Ensuring that reviews met predefined standards at both the start and end of the process.
- Reviewer Commitment: Establishing accountability for reviewers in the process.
- Integration into R&D Definition of Done Criteria: Embedding review quality into the organization's standard development lifecycle.

6) Mindset and Behavior Change through Gamification

Recognizing the critical role of motivation in achieving sustainable improvement, a gamification project was introduced. Drawing inspiration from [10], a potential game was designed to reward participation and contributions to the success of technical reviews. This initiative aimed to shift the mindset of reviewers, fostering a sense of ownership and enthusiasm for the review process. The game has not yet been widely taken into use yet but in the first small-scale trial reviewers earned points for timely participation, adherence to review standards, and quality feedback, which could be exchanged for recognition within the organization. The gamification approach created a competitive yet collaborative environment, further enhancing the team's commitment to continuous improvement.

C. Novelty of the Approach

This project marked a significant departure from traditional practices. Previously, technical review quality had never been measured, analyzed, or systematically improved. The development of a quantifiable KPI and the comprehensive methodology provided a reusable framework for enhancing review quality. Additionally, the integration of gamification to drive mindset change represented an innovative way to ensure that improvements were not only procedural but also cultural. While the focus of this initiative was on technical documentation, the concept can be easily adapted to other review-intensive areas, such as software code reviews or product design assessments.

D. Results

The implementation of this framework yielded remarkable improvements in the documentation department's review process:

- 60% Improvement in Technical Review Quality: The KPI demonstrated a significant enhancement in the effectiveness of reviews.
- 29% Reduction in Customer Faults: The number of customer-reported faults attributed to issues with technical reviews dropped substantially.
- Zero TL9000 Non-Conformities: All review-related non-conformities to TL9000 requirements were eliminated.

- **Improved Reviewer Engagement:** The gamification initiative led to increased participation and a noticeable shift in reviewer behavior, with teams actively striving for better outcomes.

Due to confidentiality agreements within the organization where the study was conducted, exact quantitative data - including baseline values, detailed metric calculations, and project-specific documentation artifacts - could not be disclosed. The review processes evaluated were part of internal quality improvement initiatives linked to ongoing product development efforts, and the associated metrics were derived from proprietary review records and audit outcomes. As a result, while the case study outlines the structure and application of the measurement framework, the focus remains on illustrating the methodology and its practical relevance, rather than presenting fully open, reproducible datasets.

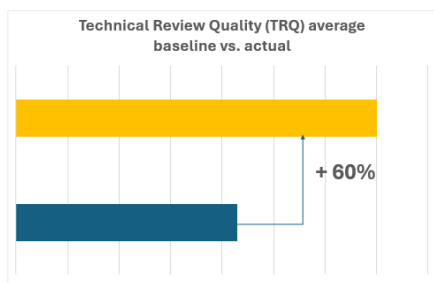


Fig. 2. Result: improvement in Technical Review Quality

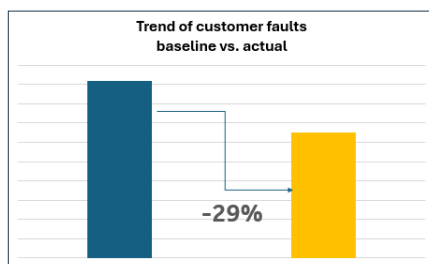


Fig. 3. Result: customer fault reduction

VII. CONCLUSION

The systematic measurement, analysis, and improvement of technical review quality in this project demonstrate the potential for significant impact on documentation processes. By addressing root causes, educating teams, embedding quality checkpoints into workflows, and fostering a mindset of continuous improvement through gamification, the initiative achieved sustainable and cultural changes. The methodology and outcomes serve as a model for other organizations seeking to enhance review processes in documentation or other domains.

To conclude, this study provides the following practical recommendations for professionals seeking to improve the quality of their review processes:

1. **Measure review quality:** Implement a systematic approach to measuring technical review quality, focusing on pre-review and review conditions. Establish clear metrics, such as a Technical Review Quality KPI, to track and monitor improvements effectively.

2. **Prioritize linguistic factors:** Pay attention to linguistic considerations during reviews, ensuring clarity, accuracy, and comprehensibility in technical documentation. Training and resources should emphasize the importance of precise language use to minimize errors.

3. **Use checklists and automation:** Employ checklists to standardize the review process, ensuring consistency and thoroughness. Consider automating data collection and analysis to streamline the identification of problem areas and track progress over time.

4. **Foster collaboration:** Promote a collaborative mindset between technical writers and reviewers by creating a culture of shared responsibility for review quality. Encourage open communication and teamwork, making the review process a constructive, cooperative effort.

5. **Encourage and incentivize mindset changes:** Leverage gamification or similar initiatives to boost motivation and engagement among reviewers. Reward participation, adherence to review standards, and quality contributions to foster a sense of ownership and enthusiasm for achieving high review standards.

Limitations and Future Research

The methodology and the implementation focus on technical reviews for documents created for users and therefore, it might not work the same way for other type of documents or technical documents in other industries or areas.

Potential future research will widen the scope of the methodology to other type of artefacts (software code, multimedia elements, internal specifications, etc.), technical documents in other STEM (Science, Technology, Engineering, Mathematics) fields. In addition, the impact of distinct groups of authors (software developers, engineering students), different review methodologies and the role of automation in review quality measurement will be studied.

While the proposed framework was developed with a primary focus on documentation authored by professional technical writers, many contemporary documentation practices - particularly in agile and DevOps environments - involve collaborative authoring with developers, engineers, or other subject matter experts. In such contexts, the nature and consistency of content, review expectations, and language quality may vary more widely. Future research should explore how the defined quality criteria perform across these more heterogeneous authoring scenarios.

This study did not extensively explore the use of automation or AI-assisted tools in documentation review, as its primary focus was on evaluating and improving existing manual review practices within a specific organizational context. While recent developments in natural language processing and review automation were acknowledged in earlier sections, their implementation was beyond the scope of the current work. Future research should investigate how such technologies, particularly AI-based comment analysis and automated quality checks, could be integrated into the proposed.

Framework for Intrusion Detection in IoT Networks: Dataset Design and Machine Learning Analysis

Another promising avenue for future research is the application of textual analysis methods to reviewer comments in technical documentation, as demonstrated in studies analyzing game reviews (e.g., [23]). The intent here is not to equate the subject matter of game reviews with that of technical documentation, but to illustrate how similar analytical methods can be adapted across domains. Techniques such as word frequency analysis, word associations, and sentiment analysis could provide deeper insights into the focus areas, concerns, and linguistic patterns in technical review feedback. By adapting these methods, it would be possible to identify recurring themes, better understand reviewer priorities, and even uncover implicit biases in the review process. Such an approach could inspire new strategies to further refine technical reviews and enhance their effectiveness.

Ultimately, this paper represents a significant step forward in closing the gap in understanding the effectiveness of technical reviews and establishing a systematic, structured methodology for measuring their quality. By integrating practical solutions, such as gamification to drive mindset change and actionable recommendations for practitioners, this study demonstrates the real-world impact of improving review processes. These principles not only deliver measurable benefits but also provide a foundation for organizations to extend this approach to other areas of quality management. Looking ahead, the insights and methodologies presented here offer a pathway for future research and innovation, helping industries align with evolving demands and set new standards for excellence in technical review quality.

REFERENCES

- [1] IEEE Standard for Software Reviews. IEEE Std 1028-1997, pp.1–48, 1998. [doi: 10.1109/IEEESTD.1998.237324](#)
- [2] N. G. Leveson, C. S. Turner, "An investigation of the Therac-25 accidents," *Computer*, vol. 26, no. 7, pp. 18–41, 1993. [doi: 10.1109/MC.1993.274940](#)
- [3] H. J. Harrington, "Measure and Control" in *Business Process Improvement: The Breakthrough Strategy for Total Quality, Productivity, and Competitiveness*, McGraw-Hill, 1991. ISBN: 978-0-07-026768-2
- [4] B. Freimut, "A Measurement Framework for Software Inspections in the Quasar Context", *IESE Report* No. 118.03/E, a publication by Fraunhofer IESE, 2003.
- [5] J. A. T. Hackos, "Conducting technical reviews" in *Management of Documentation Projects*, In Wiley Encyclopedia of Electrical and Electronics Engineering, J.G. Webster (Ed.), 1999.
- [6] J. Zuchero, "Using inspections to improve the quality of product documentation and code", *Technical Communication: Journal of the Society for Technical Communication*, 42(3), 426–435., 1995.
- [7] M. E. Raven, "What kind of quality are we ensuring with document draft reviews? *Technical Communication*, 42(3), 399–408. 1995.
- [8] M. Fagan, "Design and code inspections to reduce errors in program development", *IBM Systems Journal*, vol. 15, no. 4, pp. 182–211., 1976.
- [9] M. V. Mäntylä, C. Lassenius, "What types of defects are really discovered in code reviews?", *IEEE Transactions on Software Engineering*, vol. 35, no. 3, pp. 430–448, May-June 2009. [doi: 10.1109/TSE.2008.71](#)
- [10] O. Laitenberger, J. M. DeBaud, "An encompassing life cycle centric survey of software inspection", *Journal of Systems and Software*, 50(1), 5–31., 2000. [doi: 10.1016/S0164-1212\(99\)00073-4](#)
- [11] K. E. Wiegers, "Metrics" in *PeerReviews in Software: A Practical Guide*, Addison-Wesley, 2002.
- [12] F. Shull, et al., "What we have learned about fighting defects" in *Proceedings of the 27th International Conference on Software Engineering (ICSE '02)*, 249–258., 2002. [doi: 10.1109/METRIC.2002.1011343](#)
- [13] H. Cadavid, et al., "Documentation-as-code for interface control document management in systems of systems: A technical action research study." in *European Conference on Software Architecture* (pp. 19–37). Cham: Springer International Publishing. 2022. [doi: 10.1007/978-3-031-16697-6_2](#)
J. Y. Khan, et al., "Automatic detection of five api documentation smells: Practitioners' perspectives" in *2021 IEEE International Conference on Software Analysis, Evolution and Reengineering (SANER)*, pp. 318–329, IEEE, 2021. [doi: 10.1109/SANER50967.2021.00037](#)
J. Holdaway, "Technical reviews: framing the best of the best practices", *IEEE International Professional Communication Conference*, Waikiki, HI, USA, pp. 1–13, 2009. [doi: 10.1109/IPCC.2009.5208713](#)
- [14] P. Thongtanunam, "Review dynamics and their impact on software quality", *IEEE Transactions on Software Engineering*, Vol. 47, No. 12., 2021. [doi: 10.1109/TSE.2020.2964660](#)
- [15] E. Aghajani et al., "Software documentation issues unveiled", *2019 IEEE/ACM 41st International Conference on Software Engineering (ICSE)*, Montreal, QC, Canada, 2019, pp. 1199–1210, 2019. [doi: 10.1109/ICSE.2019.00122](#)
- [16] Progress Software. (n.d.). DevTools style guide. Retrieved June 10, 2025, from <https://docs.telarik.com/style-guide/>
- [17] Sketch Engine [Online tool, academic use]. <https://www.sketchengine.eu/>
- [18] L. Kof, "Treatment of passive voice and conjunctions in use case documents", in *International Conference on Application of Natural Language to Information Systems*, pp. 181–192. Berlin, Heidelberg: Springer Berlin Heidelberg, 2007. [doi: 10.1007/978-3-540-73351-5_16](#)
"Mars Climate Orbiter mishap investigation board phase report, 1999, https://llis.nasa.gov/llis_lib/pdf/1009464main1_0641-mr.pdf
- [19] B. Raad, et al., "Exploring the Profound Impact of Artificial Intelligence Applications (Quillbot, Grammarly and ChatGPT) on English Academic Writing: A Systematic Review", *International Journal of Integrative Research (IJIR)* 1.10 599-622., 2023.
- [20] T. Guzsvinecz, and J. Szűcs, "Textual Analysis of Virtual Reality Game Reviews", *Infocommunications Journal, Joint Special Issue on Cognitive Infocommunications and Cognitive Aspects of Virtual Reality*, 2024, pp. 84–91, [doi: 10.36244/ICJ.2024.5.10](#)



Gabriella Tóth is pursuing a PhD degree at the University of Debrecen, Institute of English and American Studies, Department of English Linguistics, Hungary. Her doctoral research is carried out under the Doctoral School of Linguistics at the same university. She holds an MSc in Information Science and an MA in English Linguistics from the University of Debrecen. During her doctoral studies, she conducted research in computational linguistics, corpus linguistics, and information retrieval at the Lister Hill National Center for Biomedical Communications, U.S. National Library of Medicine, Bethesda, MD, USA. She also has professional experience at Nokia Solutions and Networks in Hungary, where she worked as a technical writer and information designer, with a particular interest in technical reviews and their linguistic aspects. Her current research focuses on the linguistic analysis of technical texts and the pragmatics of technical communication, with a special emphasis on the formulation and interpretation of error messages.

Attitude-driven Simultaneous Online Auctions for Parking Spaces

Levente Alekszejenkó and Tadeusz Dobrowiecki

Abstract—Among various parking assignment methods, auction-based procedures have emerged as a relatively simple and flexible mechanism to solve parking assignment and parking pricing problems. In the present contribution, we extend the usual purely financial bidder objective function with a mixed utility that also involves the walking distance to parking lots.

We demonstrate numerous interesting properties of the new parking scheme; some of them are provable analytically; while others are traceable in simulation. At the end of the paper, we also present some practically useful examples.

Index Terms—auctions, parking, parking assignment, parking pricing, parking simulation

I. INTRODUCTION

While cars offer convenient transportation, finding available parking is often difficult. Since Shoup's 2005 study, it has been commonly cited that 30% of urban traffic results from drivers searching for parking; though actual rates vary between 8–74% depending on location and time [1]. As parking occupancy nears 100%, search times can rise to 10–13 minutes [1].

An automated parking assignment system could reduce this inefficiency, saving time, easing congestion, and lowering environmental impact. With smartphones, cyber-physical systems, and IoT, deploying such a system (either as a new app or as part of a navigation tool) is feasible. Drivers could set their destination, parking budget, and walking preference, and be guided to a suitable spot. For instance, a Vickrey-Clarke-Groves (VCG) auction-based app was proposed in [2].

Beyond assigning parking spaces, auctions can dynamically adjust pricing based on demand. In 2011, San Francisco's *SFPark* initiative optimized parking fees using real-time occupancy data [3]. The program successfully maintained ideal occupancy levels by influencing behavior (encouraging long-term and budget-conscious drivers to choose cheaper, more distant spots). However, field studies [4], [5] show many drivers, especially on errands or commutes, are unwilling to forgo close parking. Preferences also vary by age, income, trip type, and number of passengers.

Motivated by these differing preferences, we propose a simultaneous online auction system that assigns and prices parking based on driver attitudes toward cost and walking distance. A smartphone app would automatically bid on behalf

of users, reflecting these preferences. We demonstrate that this system aligns with user expectations through theoretical analysis and simulations using Eclipse SUMO [6].

To broaden the understanding of auction-based parking assignment, we present a short literature review in section II, an abstract parking supply and demand models are described in section III, and the handling of drivers' attitude is introduced in section IV. We formally describe *how this algorithm influences parking itself and what would be the experience of the stakeholders of the system* (e.g., drivers, municipalities, and parking lot operators) from a monetary (section V) and a monetary and distance related (section VI) point of view. Moreover, by explaining its implementation details in section VII, we also provide demonstrations for the theoretical results in section VIII, including an Eclipse SUMO-based simulation. After a short discussion of the applied method in section IX, section X concludes this paper.

II. RELATED WORKS

In 1969, parking lots occupied 87% of U.S. land-use coverage [7]. Today, European cities increasingly integrate parking policies into broader transport strategies to enhance accessibility, stimulate local economies, and improve quality of life. Common measures include limiting supply and using real-time dynamic pricing [8]. Criteria for evaluating such policies—like cost, walking time, capacity, and search time—are outlined in [9], which also classifies models by scale, from individual lots to city-wide systems.

For example, queueing models such as [10] analyze parking behavior within a single facility, while broader models like [11] examine city-wide impacts of pricing, congestion, and travel time, finding that real-time occupancy data can improve travel times by 4% and dynamic pricing can lead to socially optimal outcomes.

Dynamic pricing, first proposed in the 1950s, gained traction with large-scale implementations in the 2010s [12]. A wide range of techniques could be applied, including numerical optimization [13], Gale-Shapley algorithm-based matching techniques [14], [15]. Besides centralized approaches, distributed matching algorithms can also serve intelligent parking [16]. However, most of these solutions require additional steps to compute parking prices. The posted pricing algorithm also solves the parking charging problem [17].

Besides matching algorithms, dynamic programming, and game theory-based approaches [18] lead to auction-based systems. VCG auctions, originally designed for divisible goods with few participants [19], have been adapted for parking

L. Alekszejenkó and T. Dobrowiecki were colleagues of the Department of Artificial Intelligence and Systems Engineering, Budapest University of Technology of Economics, Budapest, Hungary e-mail: {alekszejenkó, dobrowiecki}@mit.bme.hu

Project no. TKP2021-EGA-02 has been implemented with the support provided by the Ministry of Culture and Innovation of Hungary from the National Research, Development and Innovation Fund, financed under the TKP2021-EGA funding scheme.

Attitude-driven Simultaneous Online Auctions for Parking Spaces

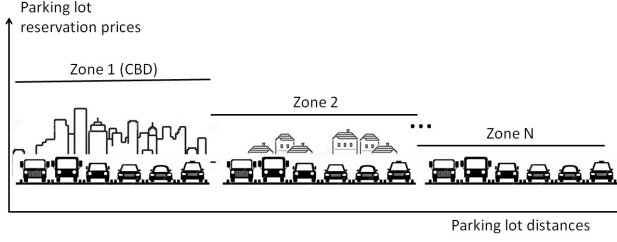


Fig. 1. Modeling parking in a city. Closer parking lots are more demanded and more expensive ones than further away alternatives.

allocation and pricing [20]–[23], although they often involve complex and time-consuming integer programming. To address this, [24] proposes more efficient alternatives.

In addition to VCG, ascending-bid English auctions are also viable. Bansal and Garg [25] introduce Simultaneous Independent Online Auctions (SIA), with Local Greedy Bidding (LGB). This LGB strategy ensures a bidder only bids on the items maximizing its utility, with auction prices increasing incrementally by ϵ .

Literature on parking mechanisms often emphasizes theoretical guarantees, such as individual rationality, budget balance, incentive compatibility, and asymptotic efficiency [20]–[23]. However, in-depth analysis of how these systems affect the stakeholders (e.g., drivers or parking operators) are rarely carried out. For example, [23] addresses user privacy, [18] discusses pricing-related properties, and [22] demonstrates expected utilities.

In this paper, we use the SIA/LGB approach of [25] to implement a solution similar to that in [26], in which all parking lot operators can organize an auction to sell their free parking spaces. Parking operators can set predefined starting prices¹ to be able to define a minimum parking costs, i.e., corresponding to traditional parking fees. The auctions shall run simultaneously, and they can increase actual bids by a predefined ϵ amount. With an attitude factor, drivers can express whether they prefer cheaper or closer parking lots. Our main contribution is the analytical and numerical analysis, including an Eclipse SUMO [6] traffic simulation-based demonstration of the SIA-based parking assignment and pricing system.

III. PARKING MODEL

This paper uses a simplified model of parking supply and demand, assuming negligible differences in driving times, no traffic congestion or road tolls. Parking costs and walking distance are considered the primary factors influencing parking choice.

Parking demand is shaped by human activity and peaks at predictable times, such as just before working hours or some special events. We assume this demand is concentrated in the central business district (CBD), following some probability distribution.

¹Starting prices are usually referred as reserve prices [27], but in this paper, we will use the term ‘starting prices’ to avoid confusion.

People vary in their tolerance for walking. Some, like plumbers or delivery drivers, require parking directly at their destination. Others, such as leisure visitors, may accept longer walks in exchange for lower fees.

Supply mirrors this pattern: parking prices are highest in the CBD and decrease toward the outskirts. Despite equal aerial distances, perceived walking distances can differ due to obstacles like rivers or railways, which help define parking zones [7]. These zones often reflect the same gradient in parking prices, see Fig. 1.

We propose an auction mechanism to assign parking spots based on drivers’ preferences for cost and walking distance. In the following, we assume that there is an idealistic SIA with LGB algorithm-based application through which the drivers can reserve parking lots. This idealistic application has perfect information of free parking spaces, it is tamper-proof, and strictly follows the protocols of the proposed auction scheme. Additionally, drivers also obey the rules and occupy exactly those parking spots that the SIA-based algorithm assigns to them. For this study, we also assume that the parking supply and demand (including drivers’ attitudes) are externally defined.

IV. PARKING ASSIGNMENT WITH SIMULTANEOUS INDEPENDENT ONLINE AUCTIONS

In classical auctions, surplus has a monetary definition; however, when we use SIAs for parking assignment, the monetary surplus definition might not consider a natural behavior of drivers. Traditionally, when drivers arrive at their destination, they start cruising around until they can find a suitable parking place. As drivers generally turn [28], or generally drive in circles [29], trying to minimize the walking distance between the parking space and their destination. On the other hand, it is rational that one would like to minimize its parking expenses.

In traditional parking searches, it doubles the challenge to optimize both parking fees and walking distances. However, in an automated parking assignment method, we can expect to solve (at least partially) these optimization problems. Therefore, in this paper, we define a suitable preference function for the SIA with LGB strategy fusing monetary costs with walking distances to the parking lots.

Considering that there are $N \in \mathbb{N}^+$ parking lots in the area and each driver j , $j \in \{1, 2, \dots, M\}$, $M \in \mathbb{N}^+$ has a limited monetary parking budget p_{v_j} corresponding to the *valuation* concept of classical auctions. Moreover, a driver agent j can compute the $d_{j,i}$ walking distance between the i th parking lot and its destination, where $i \in \{1, 2, 3, \dots, N\}$. Let p_i denote the actual parking cost at the i th parking lot, and let $d_{j,\max} = \max_i d_{j,i}$. Reflecting the attitude towards walking distance and parking prices, we propose to define a utility function $U_j(i)$ of each vehicle j for a parking lot i that is a combination of a monetary $0 \leq U_{p,j}(i) \leq 1$ and a distance-related $0 \leq U_{d,j}(i) \leq 1$ utility component corresponding to the state of the auctions:

$$U_j(i) = \beta_j U_{p,j}(i) + (1 - \beta_j) U_{d,j}(i), \quad (1)$$

using an attitude factor $\beta_j, 0 < \beta_j \leq 1$. By setting a low β_j , for example, a furniture deliveryman can prefer the closest parking lots. On the other hand, if we plan some leisure activity where walking is desired, we can set a high β_j and leave our vehicle in more distant, yet cheaper parking lots.

The monetary utility component can be: $U_{p,j}(i) = \frac{p_{\max} - p_i}{p_{\max}}$, where $p_{\max} = \max_i p_i$ corresponds to the maximum price achievable in auctions (e.g., the maximum valuation that drivers can afford). Similarly, the distance-related utility component can be: $U_{d,j}(i) = \frac{d_{j,\max} - d_{j,i}}{d_{j,\max}}$, where $d_{j,\max}$ is the distance between the most remote parking lot from the j th driver's destination.

Substituting these utility components with (1), we get the $U_j(i)$ utility function:

$$U_j(i) = \beta_j \frac{p_{\max} - p_i}{p_{\max}} + (1 - \beta_j) \frac{d_{j,\max} - d_{j,i}}{d_{j,\max}} = \quad (2)$$

$$= 1 - \beta_j \frac{p_i}{p_{\max}} - (1 - \beta_j) \frac{d_{j,i}}{d_{j,\max}} = 1 - c_{j,i}, \quad (3)$$

where $c_{j,i}$ is the overall cost of the j th vehicle parking at the i th parking lot.

Therefore, we can define the Π_j auction with highest utility for the LGB strategy as:

$$\Pi_j = \arg \max_i U_j(i) = \arg \min_i c_{j,i}. \quad (4)$$

Naturally, since the j th driver's parking budget is limited by p_{v_j} , they shall only bid for the Π_j th parking lot, if $p_{\Pi_j} \leq p_{v_j}$. Consequently, at each moment, driver agent j can select a Π_j parking lot that appears to be suitable for bidding. As there could be multiple parking lot auctions that fulfill the criterion of (4), in the following, we will treat Π_j as a set of appropriate parking auctions. The following Theorem proves that the proposed algorithm will stop after receiving a sufficient amount of bids.

Theorem 1. *If $\beta_j > 0$ and $p_{v_j} > 0$ are finite (individual) constants for all bidders, then the SIA algorithm will terminate.*

Proof. The SIA method terminates if and only if it finds a good assignment. As current prices increase by ϵ upon receiving a new bid, sooner or later the current prices will either reach a state that is a solution or they reach the p_{v_j} valuation of the bidders. Denoting $p_{v_{\max}} = \max_j p_{v_j}$ and $p_{\min}^{(0)} = \min_i p_i^{(0)}$, where $p_i^{(0)}$ stands for the starting price of auction i , the $t_{\max}$ number of bids to reach the valuation of all bidders in all N auctions is bounded by:

$$t_{\max} = \left\lceil \frac{p_{v_{\max}} - p_{\min}^{(0)}}{\epsilon} \right\rceil N. \quad (5)$$

Consequently, the SIA method terminates after receiving $t \leq t_{\max}$ bids. $\square$

In the following, to analyze the proposed system, we demonstrate some theoretically provable properties, illustrate them by numerical simulations, and also provide some application use cases.

V. DISTANCE-INDEPENDENT PROPERTIES OF THE PARKING SIA

Now, we will examine the SIA mechanism for parking lot assignment with the utility function defined in (2), analytically if it is possible or by simulations. By varying the β_j attitude factor, we favor smaller walking distances or lower parking costs.

If a driver chooses an $\beta_j \approx 0.0$ value, it results in a classical parking search algorithm, when people want to park near their destination, regardless of parking fees. In this case, the SIA algorithm is only expected to provide an assignment between parking lots and vehicles without optimizing parking costs, which in over-demand situations ($M > N$) can reach p_{v_j} .

On the other hand, $\beta_j \approx 1.0$ can also be a relevant selection if there are multiple parking lots, possibly having different parking prices, in such a small area, in which drivers will not perceive significant differences in walking distance. Supposing a small area and an $\beta_j = 1.0$ setting, the SIA method can purely optimize parking costs in addition to assigning vehicles and parking spaces. The following lemmas and theorems help us to understand the effect of SIAs on parking prices.

A. Expected Parking Prices

First of all, let us check how prices change during the execution of SIA. Let $\mathbf{p}^{(t)}$ parking prices vector collect all the actual parking prices $p_i^{(t)}$ that received bids after t steps of auctions. Then, Lemma 1. expresses the price changes in SIA.

Lemma 1. *$\exists T$, such that after running SIA for $t > T$ steps, the difference between the elements of the current price vector $\mathbf{p}^{(t)}$ will be $\leq \epsilon$.*

Proof. The proof is given in Appendix A. $\square$

Naturally, the rate of supply and demand for parking lots define parking costs in market-driven pricing scheme. As SIA reacts to the actual parking situation, we can expect that the obtained parking prices obtained also reflect it. Theorem 2. refers to the situation in which the supply is either in equilibrium with the demand or exceeds it.

For a compact notation, let us define two symbols. Firstly, let us create an *ascending list* of starting prices of the auctions. Here, $p_{M-}^{(0)}$ will denote the M th element of the list, corresponding to the M th *lowest* starting price. Secondly, we create a *descending list* of valuations of the bidders. Here, $p_{v,N+}$ will denote the N th element of the list, corresponding to the N th *highest* valuation.

Theorem 2. *Assuming that there is no over-demand for parking ($N \geq M$) and $\beta_j = 1.0$, each element of the assigned, won parking prices vector $\mathbf{p}$, $\mathbf{p} \in \mathbb{R}^M, M > 1$ provided by the SIA method will be approximately equal to the M th lowest $p_{M-}^{(0)}$ starting price: $p_{M-}^{(0)} \leq p_i \leq p_{M-}^{(0)} + \epsilon < p_{v_j}$ for $p_i \in \mathbf{p}, j \in \{1, 2, \dots, M\}$.*

Proof. With the assumptions, the utility function of each j driver agent simplifies to $\Pi_j^{(t)} = \arg \max_i \left(1 - \frac{p_i^{(t)}}{p_{\max}}\right)$ in the t th step of the auction. Consequently, following the LGB strategy, each j driver will bid for one of the cheapest parking

Attitude-driven Simultaneous Online Auctions for Parking Spaces

lots. Moreover, $p_i^{(t+1)} = p_i^{(t)}$ or $p_i^{(t+1)} = p_i^{(t)} + \epsilon$, depending on whether or not someone has bid for it.

If the starting prices are identical, $p_i^{(0)} = p_{i+1}^{(0)}, \forall i : i \in \{1, 2, \dots, N-1\}$, and there is no over-demand, the statement is a natural consequence, and the resulting prices will be the original starting prices.

If the starting prices are different, we shall prove that the M vehicles will compete for the cheapest M parking lots until the resulting prices will not differ more than ϵ , and each vehicle will have been winning exactly one parking space auction. Lemma 1. proves this case. $\square$

On the other hand, the demand can also exceed the supply of parking lots. Theorem 3. shows what would happen in this scenario when we use the SIA algorithm.

Theorem 3. *Assuming that there is an over-demand for parking ($N < M$) and $\beta_j = 1.0$, each element of the winning parking prices vector $\mathbf{p}$, $\mathbf{p} \in \mathbb{R}^N$, $N > 1$ provided by the SIA method will be approximately equal to the $(N+1)$ th highest $p_{v,(N+1)+}$ valuation: $p_{v,(N+1)+} \leq p_i \leq p_{v,(N+1)+} + \epsilon$ for $i \in \{1, 2, \dots, N\}, j \in \{1, 2, \dots, M\}$.*

Proof. The proof is given in Appendix B. $\square$

In summary, due to Lemma 1., prices increase rationally. In a non-over-demand scenario, this process leads to an assignment on the price of the M th cheapest parking lot according to Theorem 2. In contrast with this, in an over-demand scenario, the richest N drivers (those who have the highest p_{v_j} valuations) can find parking spaces for themselves because of Theorem 3.

B. Optimal Parking Lot Size

From another perspective, parking lot operators might face the problem of having to decide how many parking spaces are required to maximize their revenue. The following corollary helps find the optimal parking lot size to maximize parking incomes if the parking pricing system is driven by SIA with the LGB strategy.

Corollary 3.1. *Taking into account the constant demand of $M > 2$ vehicles with $\beta_j = 1.0$, $p_{v_j} > p_i^{(0)}$ for $i \in \{1, 2, \dots, N\}, j \in \{1, 2, \dots, M\}$, an optimal parking lot, which maximizes parking incomes by running SIA with the LGB strategy, provides $N = M - 1$ parking spaces.*

Proof. Let $Np_v = \sum_{j=1}^N p_{v,j}$. Then, in the over-demanded ($N < M$) case, the operator can obtain an income of Np_v , while in the $N \geq M$ case, it can earn $Np_i^{(0)}$. As $p_v > p_i^{(0)}$, at the $N = M$ setting ($Np_v > Np_i^{(0)}$), the operator would lose a significant amount of money.

As Np_v is a strict monotonic function of N , let us check whether the total income in the $N = M - 1$ exceeds the $N = M$ setting:

$$(M-1)p_v >? Mp_i^{(0)} \quad (6)$$

$$M >? 1 + \frac{p_i^{(0)}}{p_v - p_i^{(0)}} \quad (7)$$

As $p_v > p_i^{(0)}$, the right hand side converges to 1.0; hence, if $M > 2$, then the statement is true. $\square$

VI. DISTANCE-DEPENDENT PROPERTIES OF THE PARKING SIA

To make driver decisions more flexible, we introduced a mixed cost and distance-aware utility (2) in section IV. We anticipate a more cost-aware driver (who bids with a high attitude factor β_j) would win a cheaper yet more distant parking lot, and vice versa.

As drivers bids for parking lots maximizing their momentary utility, the following lemma clarifies the relation between the maximum utility and the driver's awareness of costs over distances.

Lemma 2. *Assuming that two parking lots i and i' are equally useful for a vehicle j , at some β_j : $U_j(i) = U_j(i')$. Without loss of generality, let us also assume that $p_i > p_{i'}$ and $d_{j,i} < d_{j,i'}$.*

Then, by changing β_j with $\Delta\beta > 0$, vehicle j will prefer either parking lot i or i' as follows:

- (I) *If $\beta_1 = \beta_j - \Delta\beta$, then vehicle j will prefer the closer but more expensive parking lot, i.e.: $U'_j(i) > U'_j(i')$.*
- (II) *If $\beta_2 = \beta_j + \Delta\beta$, then vehicle j will prefer the more distant but cheaper parking lot, i.e.: $U'_j(i) < U'_j(i')$.*

Proof. We check the two cases:

- (I) If $\beta_1 = \beta_j - \Delta\beta$ then $U'_j(i) > U'_j(i')$.

$$U'_j(i) >? U'_j(i') \quad (8)$$

$$U_j(i) + \Delta\beta \left(\frac{p_i}{p_{\max}} - \frac{d_{j,i}}{d_{j,\max}} \right) >? \quad (9)$$

$$>? U_j(i') + \Delta\beta \left(\frac{p_{i'}}{p_{\max}} - \frac{d_{j,i'}}{d_{j,\max}} \right) \quad (10)$$

Following the assumption $U_j(i) = U_j(i')$:

$$0 >? \Delta\beta \left(\frac{p_{i'} - p_i}{p_{\max}} + \frac{d_{j,i} - d_{j,i'}}{d_{j,\max}} \right) \quad (11)$$

$p_i > p_{i'}$ and $d_{j,i} < d_{j,i'}$ yield that the right-hand side of the above expression is negative. Hence, the first statement is true.

- (II) If $\beta_2 = \beta_j + \Delta\beta$ then $U'_j(i) < U'_j(i')$. Analogously to (I), and following the assumption $U_j(i) = U_j(i')$, we can get:

$$0 <? \Delta\beta \left(\frac{d_{j,i'} - d_{j,i}}{d_{j,\max}} + \frac{p_i - p_{i'}}{p_{\max}} \right) \quad (12)$$

As $p_i > p_{i'}$ and $d_{j,i} < d_{j,i'}$, the right-hand side of the above expression is positive. Hence, the second statement is also true. $\square$

Lemma 2 only presents an elementary step of the bidding process. Following analytically the distribution of the winners' attitudes is perhaps inextricable considering the parallel bidding and the randomness of the driver attitudes. However, we expect that a certain overall regularity might emerge (on the average) in the bidding process.

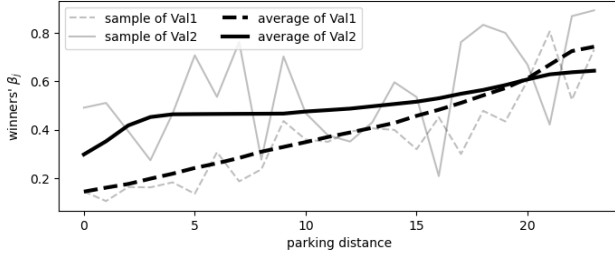


Fig. 2. Distribution of winners' attitudes in function of distance of the won parking lots.

To illustrate it, we present in Fig. 2. the attitude distribution among the winner drivers as a function of the parking lot distance, for an individual auction (*sample*), and for the auction results averaged over 10000 auction runs (*average*). In the simulation (see section VII), $M = 35$ drivers competed for $N = 24$ parking lots with the same starting price ($\forall i : p_i^{(0)} = p_{r, const}$), evenly spaced with increasing distance from the center. The drivers' attitudes were sampled from a uniform distribution $\beta_j \sim \mathcal{U}(0.1, 0.9)$. The drivers' validations were constant i.e.: $\forall j : p_{v_j} = V$ in the Val1 scenario; and in Val2, they were sampled from a uniform distribution $p_{v_j} \sim \mathcal{U}(\frac{2}{3}V, V)$ where V was the maximal validation and $\frac{2}{3}V > p_{r, const}$.

We can observe in Fig. 2. that sample runs are increasing erratically; however, the expected values are smoothly monotonic. Indeed, with the introduced utility mechanism, drivers who opt for cheaper but more distant parking lots win on the average accordingly, and in a distance-dependent proportional way. The difference in behavior of the two cases is most certainly due to the uncorrelated character of random validations (Val2) and the random attitudes. When a strong negative correlation is introduced between random validation and random attitudes (i.e., the higher validation pairs with a lower attitude; those who have more money want to park closer), the sample and average curves are similar to the Val1 case.

VII. ON THE IMPLEMENTATION OF PARKING SIA

Originally, SIA algorithm can be executed in a highly parallel way, in which each parking space can run its own auction server and all parking-seeking drivers (or automated bidder agents impersonating them) can send bids asynchronously. Unfortunately, simulating this parallelism on the PC-based research architecture is not efficient because of the frequent context changes during parallel execution.

Consequently, we serialize the execution of the SIA algorithm. It requires a slight modification of the original behavior of the bidders. Instead of asynchronous execution, the bidder will execute an event-driven program, following the state machine in Fig. 3, of which more complex functions are described by Algorithm 1.

Upon request from a particular auction a_i , a bidder computes whether it is willing to bid for it or not following Algorithm 1. In parallel execution, this function would be

Algorithm 1 Bidders' main functions

```

1: function BIDDER.ASK_BID( $a_i$ )
2:    $bids \leftarrow \text{False}$ 
3:   recall  $state$ 
4:   recall  $p$   $\triangleright$  all price values
5:    $\mathcal{F} \leftarrow \{a_j : p_j \leq p_v, j \in \{1, 2, \dots, N\}\}$   $\triangleright$  feasible  $a_j$ 's
6:    $\Pi = \arg \min_k c_k$   $\triangleright$  computing preferences as in (4)
7:   if ( $state = \text{'overbid'}$ )  $\wedge$  ( $p_i \leq p_v$ ) then
8:      $\mathcal{P} \leftarrow \{a_j : a_j \in \mathcal{F} \cap \Pi\}$   $\triangleright$  preferred auctions
9:      $bids \leftarrow (a_i \in \mathcal{P})$ 
10:    if  $bids$  then
11:      store  $state \leftarrow \text{'winning'}$ 
12:    end if
13:  end if
14:  return  $bids$ 
15: end function
16: function BIDDER.INFORM_PRICE( $p_i$ )
17:   store  $p_i$ 
18:   recall  $state$ 
19:   if ( $\nexists j : p_j \leq p_v$ )  $\wedge$  ( $state \neq \text{'winning'}$ ) then
20:     store  $state \leftarrow \text{'out of budget'}$ 
21:   end if
22: end function
    
```

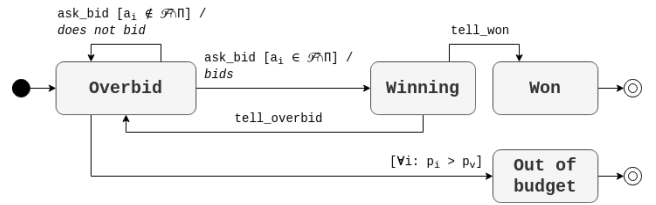


Fig. 3. States machine representing the Bidder agent.

similar; the only difference is that some scheduler algorithm would call this subroutine, and it would return a random element from the $\mathcal{F} \cap \Pi$ set to bid for.

The main difference between SIAs and the serialized simulation is described in detail in Algorithm 2. Instead of running many algorithms in parallel, we iterate over them and over all the bidder agents in line 6 and line 7. The code inside these iterations executes the event-driven calculations.

As stated above, we modified SIA in numerous points to be able to simulate it on simple PCs. Theorem 4. justifies that a fully parallel SIA is the generalization of the serialized auction implementation. Consequently, all theorems, lemmas, and corollaries of sections IV, V, and VI are also necessarily true for serialized auctions.

Theorem 4. *The fully parallel execution of the SIA is a generalization of the serial execution of the auction method.*

Proof. The proof is given in Appendix C. $\square$

VIII. EXAMPLES

By running the serialized SIA defined in section VII, we made experiments to demonstrate some of the theorems, lemmas, and corollaries of section V and VI. Furthermore, we provide some application examples of the proposed method.

Algorithm 2 Auctioneers' algorithm

```

1: function RUN_AUCTIONS( $A, B, \mathbf{p}_s, \epsilon, r_{\max}$ )
2:    $\mathbf{w} \leftarrow \emptyset$   $\triangleright$  init: no winner
3:    $\mathbf{p} \leftarrow \mathbf{p}_s$   $\triangleright$  init: price = starting price
4:    $B_j.state \leftarrow \text{'overbid' } \forall B_j \in B$   $\triangleright$  init: state of bidders
5:   while  $\exists l : B_l.state = \text{'overbid'}$  do  $\triangleright \exists$  running bidder
6:     for  $i \in \{1, 2, \dots, N\}$  do  $\triangleright$  iterating on auctions
7:       for  $j \in \{1, 2, \dots, M\}$  do  $\triangleright$  iterating on bidders
8:          $bids \leftarrow B_j.ASK\_BID(a_i)$ 
9:         if  $bids$  then
10:          if  $w_i \neq \emptyset$  then  $\triangleright$  inform previous winner
11:             $w_i.TELL\_OVERBID$ 
12:          end if
13:           $w_i \leftarrow B_j$   $\triangleright$  new winner
14:           $p_i \leftarrow p_i + \epsilon$   $\triangleright$  increase price
15:          for  $k \in \{1, 2, \dots, M\}$  do
16:             $B_k.INFORM\_PRICE(p_i)$ 
17:          end for
18:        end if
19:      end for
20:    end for
21:  end while
22:  for  $i \in \{1, 2, \dots, N\}$  do
23:     $w_i.TELL\_WON$ 
24:  end for
25:  return  $\mathbf{w}, \mathbf{p}$ 
26: end function
    
```

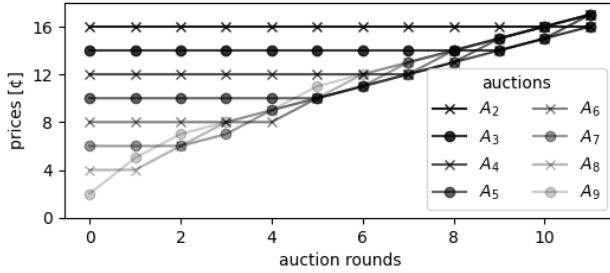
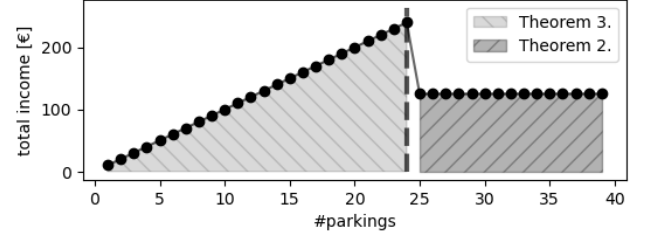


Fig. 4. Parking prices during auctions.

A. Illustration of Lemma 1

According to Lemma 1, prices will be approximately equal after a number of auction rounds. To demonstrate this, we created a simple simulation with $N = 10$ parking spaces and $M = 8$ vehicles. The starting prices at the corresponding $A_0, A_1, A_2, \dots, A_9$ auctions were 20, 18, 16, $\dots$, 2 €. For a simple illustration, we assumed here that drivers do not differentiate parking spaces by distance. We used a bid step of $\epsilon = 1.0$ €, and the valuation of each $j \in \{1, 2, 3, \dots, 8\}$ vehicle was $p_{v_j} = 1000$ €.

Fig. 4. shows the resulting parking prices after each auction round. An auction round corresponds to the state achieved after a run of iteration of line 6 in Algorithm 2. The results show that the current prices at each of (the cheapest M) auctions increase together. The final winning prices are at most ϵ apart from each other.


 Fig. 5. Optimal sizing of a parking lot to maximize the operator's income at a demand of $M = 25$ vehicles.

B. Illustration of Corollary 3.1

As Corollary 3.1 combines the consequences of Theorem 2 and Theorem 3, we created a simulation to demonstrate them together. Hence, we assume a parking demand of $M = 25$ vehicles, and after running SIAs we computed the total income of the parking lot operator for various $N \in \{1, 2, 3, \dots, 40\}$. The attitude factor was $\beta = 1.0$, the bid step was $\epsilon = 1.0$ €, and the valuation of each $j \in \{1, 2, 3, \dots, 8\}$ vehicle was $p_{v_j} = 1000$ €. The starting prices for each auction were 500 €.

On the left-hand side of Fig. 5., when $M < 25$, we can find the demonstration of Theorem 3 regarding the parking prices obtained in the over-demanded case. On the right-hand side of Fig. 5., when $M \geq 25$, we can observe the consequences of Theorem 2 regarding parking prices in a not over-demanded scenario. Finally, Fig. 5. also shows that the operator of the parking lot can maximize its income with an $N = M - 1 = 24$ setting, which aligns with Corollary 3.1.

C. Application: Offline Parking Pricing

The proposed auction mechanism is suitable to provide an intelligent parking pricing model in a city that calculates with a spatial (s) parking demand $D(s)$. The demand model $D(s)$ can originate from historical time series. By setting a lower attitude towards walking, e.g., $\beta = 0.1$, and using the $D(s)$ demand model, we can simulate the auction methods for a known number of vehicles (M) and parking lots (N).

To demonstrate that the auction method can suggest efficient parking price settings, we created an abstract simulation. As there are hundreds of cities in the European Union that have imposed low emission zones (LEZs), various restrictions on vehicles that can enter specific areas [30], we also model parking prices at the perimeter of an LEZ. We used $N = 474$ parking spaces in groups of 6 spaces, 50 m apart from each other. Between the positions of 3000 m and 4000 m, there was an LEZ where there were no parking lots.

From a cross section view of a city, the parking demand of $M = 200$ vehicles followed the mixed probability distribution of $D_1(s) \sim \mathcal{U}(0, 5000)|_{100} + \mathcal{N}(1000, 100)|_{100}$, where 100 vehicles had a uniform demand across the whole a city, and 100 vehicles were headed towards the 1000 m point, following a normal distribution with $\sigma = 100$ m. For all j vehicles, the valuation was $p_{v_j} = 1000$ €, and the bid step was $\epsilon = 1.0$ € in each auction. In the initial phase, there was free parking in the city. After running the SIA method, we got the results

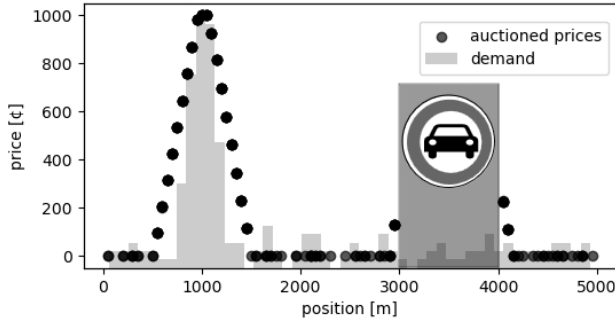


Fig. 6. SIA simulation results for a city containing a peak in the demand at the position of 1000 m, and a low emission zone between positions of 3000 m and 4000 m.

of Fig. 6., yielding that parking prices shall increase around the peak of demand, and at the perimeter of the LEZ, while parking could remain free in the rest of the city.

D. Application: Online SIA in a rural town's parking

Besides offline parking pricing optimization, SIAs might run in real time. For example, drivers might use their smartphones to navigate to a parking lot assigned to them by the SIA method, depending on their destination and attitude towards walking and parking prices. Hence, we have tested the SIA method in generated but realistic simulation of a typical European small rural town [31] having 10.000 inhabitants and commuters carrying out their daily activities. Consequently, this test is hardly a corner case for parking assignment, but there might be spatial and temporal over-demands in particular parts of the small town.

After 3 simulated days of bootstrapping, we simulated the morning traffic of the fourth day between 6:00 am and 10:00 am. During this simulation, we started SIAs for all empty parking spaces, and the simulated vehicles had to bid for them to reserve parking places, and to negotiate an hourly parking price. The starting prices for each auction were 50 ¢, and the bid step was $\epsilon = 5$ ¢. Each vehicle j had a valuation of $p_{v_j} = 1000$ ¢ = 10 €. Originally, vehicles aimed to use the curbside parking lot, provided on both sides of each road segment, but at the most visited sites, we placed 8 parking garages, see Fig. 7. Parking garages are considered *alternative* parking solutions and had a fixed price of 1000 ¢. Consequently, if a vehicle could not get a curbside parking space on auctions, it would be able to use the nearest parking garage.

We experimented with different β_j settings. We tried 3 scenarios in which all vehicles had the same attitude β_j of 0.1, 0.5 or 0.9. Furthermore, we simulated a MIX scenario with 90% of the vehicles having a $\beta = 0.1$ and 10% of the vehicles having a $\beta = 0.9$ attitude, representing that 90% of the drivers want to park near their destinations, but there is a minority, who prefers longer walking distances in favor of cheaper parking alternatives.

In addition to walking distances and auctioned hourly parking prices $p_{j,i}$, we also measured the duration of parking

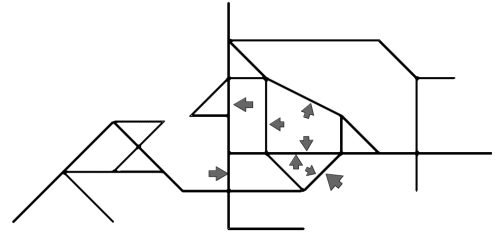


Fig. 7. Road network of the simulated town in Eclipse SUMO. Parking garages are indicated on their corresponding side of the roads.

TABLE I
ECLIPSE SUMO SIMULATION RESULTS (AVERAGE OF 5 RUNS)

β	total price [€]	distance [m]	auctioned price [€/h]	parking garage alternatives
0.1	4.86	13.81	51.97	3.28%
0.5	4.68	32.03	50.05	0.00%
0.9	4.68	32.02	50.00	0.00%
MIX	4.82	14.84	51.40	3.14%

$\tau_{j,i}$ (in seconds) to calculate the total parking price $P_{j,i}$ as:

$$P_{j,i} = p_{j,i} \cdot \left\lceil \frac{\tau_{j,i}}{60 \cdot 60} \right\rceil.$$

Moreover, we tracked the ratio of taking the parking garage alternative to the number of all simulation participation.

Table I summarizes the average results obtained from 5 distinct simulation runs of all β attitude settings. As we expected, when drivers prefer closer parking lots, it decreases walking distances and increases parking prices. In approximately 3% of the cases in this particular scenario, low β attitude setting creates such an over-demand that some vehicles have to opt for parking in parking garage instead of curbside parking lots.

IX. DISCUSSION

In this paper, we assume that parking fees and walking distances are the primary determinants of the choice of parking lot. In reality, there are additional influencing factors, for example, road tolls, traffic congestion, and the availability of public transportation. From an abstract point of view, this extra information could be integrated into (1), and with proper modifications, one could design attitude factors towards these new factors as well.

Moreover, the survey of [32] indicates some usual performance metrics to evaluate a dynamic parking pricing scheme. In the following, we describe how the SIA/LGB for parking assignment performs according to the factors relevant for dynamic parking pricing.

As the proposed method provides parking assignment, it reduces *environmental factors*, eliminates *parking searching time*; therefore, it improves *average speed* and *traffic flow*. By properly setting the attitude factor β_j , drivers can optimize its *travel time* and *utility* corresponding to their *ridership* and actual *demand*. Considering drivers' attitudes, the proposed SIA-based solution provides a more flexible approach of parking assignment compared to a posted pricing-based method, in which walking distance is defined as the social cost [17].

Following the offline results of SIA/LGB as in section VIII-C, municipalities can optimize parking prices or, following Corollary 3.1, the number of available parking lots to optimize their *revenue* by balancing the *supply*. However, drivers might face higher parking *costs* in periods of over-demand. Operators can use this extra income to cover their *operational costs*. Furthermore, Corollary 3.1 can also help parking lot operators maximize their incomes by optimizing the number of offered parking spaces on the auction rounds. If they have a demand forecast for parking, they can offer fewer than all free parking spaces in an area in specific periods. For example, a municipality might not offer all free parking spaces in the early morning to be able to serve the parking demand during later rush hours.

In terms of *computational complexity*, a fully parallel implementation of SIA auctions can run in linear time, see Theorem 1., while bidders shall run in a second-order polynomial time of the number of auctions (they shall regularly check the state of all auctions and send bids if necessary). In section VIII-D, we also presented that in a usual European rural town, around 3% of the *users get rejected*, and shall look for an alternative parking solution. In this study, we demonstrated the core parking assignment method, assuming all the drivers are cooperative, obey the rules, and occupy exactly those parking spaces reserved for them by the auction method. However, in the real world, not all drivers will or can use the implementation of the proposed system. To this end, we plan further analysis with different rates of participating and non-participating drivers to evaluate the system's properties in mixed traffic. Moreover, it would also be a further research question how strategic bidding of a group of drivers would impede the proposed method.

We have also assumed that the method has perfect information about parking lot occupancies. We ran simulations to relax this assumption and found that relatively accurate measurements are necessary (at least 93% of accuracy). Fortunately, state-of-the-art camera-based solutions can achieve this level of accuracy [33].

In this study, we assumed a predefined parking demand. In reality, various factors can influence where drivers would like to park. For example, one is likely to avoid standing in traffic congestion and prefer a park-and-ride (P+R) parking facility with good public transportation connections. We can model this attitude in the proposed solution as a slight preference for higher walking distances over parking costs. Moreover, there could be road tolls along the way from the suburbs towards the CBD of a city. One might offset the starting prices of the inner parking spaces with these fees to model the overall costs. Drivers with limited transportation budgets or who prefer cheaper parking lots might avoid these expensive alternatives by properly setting their β_j attitude factor.

Actual traffic situations, such as unusually slow traffic, also influence the choice of parking lot. By adjusting the β_j attitude factor, the proposed method could react to the changed situation. Consequently, properly setting the attitude factor might be a challenging task. By applying reinforcement learning, we might implement an automated approach to set the attitude factor for each driver appropriately. However, it

would require feedback from the drivers, which, for example, would be given by answering some questions (e.g., *Did you have to walk far?* or *Did you find the parking fees high?*).

X. CONCLUSION

In this paper we systematically analyzed how simultaneous online independent auctions with local greedy bidding strategy of [25] could be applied to the assignment of parking spaces. In our solution, similarly to [26], we used an attitude factor β_j that could be set by drivers to reflect their attitude towards closer and more expensive over more distant but cheaper parking spaces.

We theoretically proved that this method sooner-or-later terminates (Theorem 1). Additionally, we had propositions on how the prices are reaching each other and then increasing together during the auctions (Lemma 1), and what could be the price outcome of the auction method in non-overdemanded (Theorem 2, vehicles win parking spaces for the N th lowest starting prices) and over-demanded situations (Theorem 3, richest M vehicles obtains parking lots for the valuation of the M th richest drivers). Following those theorems, we provide a formula to calculate the optimal parking size ($M - 1$), if the incoming parking demand (M) is known (Corollary 3.1). Moreover, we showed that different attitude factors *have* an effect on the obtained parking lots (preferring cheaper parking lots results in more distant parking), see Theorem 2 and section VI.

By implementing SIAs in a PC environment, we demonstrated in section VIII-C how a municipality can use offline results of the auction method to implement demand-related parking prices, or how can an online deployment conduct the parking assignment in a simulated rural town, see section VIII-D.

With modern smartphones and mobile connection, there is no longer an obstacle to implement an auction-based parking guidance, assignment, and pricing system. We hope that our results add to the vision of a modern city's parking system that has the potential to make parking more efficient, and lead to urban redevelopment, leaving more space to parks, sustainable micromobility-based transportation, and local businesses, to make our habitat more livable. To encourage further experimentation and facilitate further studies, we provide our source codes here: https://github.com/alelevente/auction_theories.

APPENDIX A PROOF OF LEMMA 1.

Proof. We prove the statement by induction.

Without loss of generality, consider a simple scenario with two parking lots (A_1, A_2), two vehicles (B_1, B_2), and $p_{A_1}^{(0)} > p_{A_2}^{(0)}$. Following the LGB strategy, B_1 and B_2 bid for A_2 in the first $T = \left\lfloor \frac{p_{A_1}^{(0)} - p_{A_2}^{(0)}}{\epsilon} \right\rfloor$ steps. Let us assume that $p_{A_1}^{(T+1)} = p_{A_2}^{(T+1)}$ in the $(T + 1)$ st step, and B_2 placed the last bid on A_2 . Then B_1 can randomly select one of the following two options.

B_1 might bid for A_1 and wins it for $p_{A_1}^{(0)}$, and B_2 wins A_2 for $p_{A_1}^{(0)} - \epsilon$, since both B_1 and B_2 have placed bids on auctions and there will be no one who could overbid them.

Additionally, B_1 perhaps bids for A_2 and overbids B_2 . Hence, B_2 shall place a new bid. At this point, at $T + 2$, $p_{A_1}^{(T+2)} = p_{A_1}^{(0)}$ and $p_{A_2}^{(T+2)} = p_{A_1}^{(0)} + \epsilon$. In this case, B_2 will bid for A_1 as it is cheaper, and the auctions terminate: B_1 wins A_2 for $p_{A_1}^{(0)}$ and B_2 wins A_1 for $p_{A_1}^{(0)}$.

Let us assume that this statement is also true for $M - 1$ vehicles (and $N \geq M$ parking lots).

Now, we shall prove the statement for M vehicles. According to the induction assumption, after $T - \delta$ steps of the SIA, the first $M - 1$ vehicles are not overbid, and have bid for the cheapest $M - 1$ parking lots following the LGB strategy. Additionally, all of these $i \in \{1, 2, \dots, M - 1\}$ auctions have a current $p_{(M-1)-}^{(0)} \leq p_i^{(T-\delta)} \leq p_{(M-1)-}^{(0)} + \epsilon$ price. However, the M th vehicle is currently overbid in the $M - 1$ cheapest auctions in the $(T - \delta)$ th step. Supposing that $p_{M-}^{(0)} \geq p_{(M-1)-}^{(0)} + \epsilon$, the M th vehicle will bid for a parking space of which auction is inside the first $M - 1$ auctions as it is still cheaper than the M th parking space. Therefore, in the $(T - \delta + 1)$ st step, a vehicle that has not been overbid yet in the first $M - 1$ auctions, shall bid again on an auction making another vehicle bid in the $(T - \delta + 2)$ th step. This chain reaction ensures that in $\delta = \left\lfloor \frac{p_{M-}^{(0)} - p_{(M-1)-}^{(0)}}{\epsilon} \right\rfloor$ additional steps the first $M - 1$ auctions will reach the $p_{M-}^{(0)}$ price. Consequently, in the T th step, an overbid driver agent j will have two options:

Firstly, driver agent j might bid for parking lot with M th cheapest starting price $p_{M-}^{(T)} = p_{M-}^{(0)}$. As the $M - 1$ other vehicles will not be overbid this way, it will win this parking lot for $p_{M-}^{(0)}$, and the others will pay either $p_i^{(T)} = p_{M-}^{(0)}$ or $p_i^{(T)} = p_{M-}^{(0)} + \epsilon$ for their parking spaces. Otherwise, the M th parking lot would not be the cheapest for its $p_{M-}^{(0)}$ price.

Secondly, if j bids on any of the first $(M - 1)$ auctions for $p_{M-}^{(T)} = p_{M-}^{(0)}$ price, it will overbid another vehicle. It starts another chain reaction; however, at this time, the M th auction will be a relevant alternative for the M vehicles; hence, all the vehicles will win a parking lot for either a $p_{M-}^{(0)}$ or a $p_{M-}^{(0)} + \epsilon$ price. $\square$

APPENDIX B

PROOF OF THEOREM 3.

Proof. We prove the statement by induction.

Without loss of generality, consider a simple scenario with two parking lots (A_1, A_2), three vehicles (B_1, B_2, B_3), and $p_{v_{B_1}} > p_{v_{B_2}} > p_{v_{B_3}}$. Following the LGB strategy, B_1 , B_2 , and B_3 bid for A_1 and A_2 . According to Lemma 1, after T steps, we reach a $p_{A_1}^{(T+1)} = p_{A_2}^{(T+1)} = p_{v_{B_3}} + \epsilon$ state. In this state, B_3 is no longer capable of making additional bids for any of the auctions. Furthermore, if B_1 and B_2 sent the last bids for A_1 and B_2 , then they win A_1 and A_2 for a price of $p_{v_{B_3}}$. On the other hand, if B_3 has sent the last bid for for example A_1 for $p_{v_{B_3}}$, then, for example, B_1 is currently overbid. The B_1 vehicle has two options in the $(T + 1)$ th step.

It might bid for A_1 , it will win it for the current $p_{A_1}^{(T+1)} = p_{v_{B_3}} + \epsilon$ price. As B_2 is currently winning on A_2 (for $p_{v_{B_3}}$),

and B_3 cannot place further bids, the SIA terminates. On the other hand, if B_1 bids for A_2 , it will overbid B_2 for $p_{A_2}^{(T+1)} = p_{v_{B_3}} + \epsilon$. It makes B_2 bid for A_1 for $p_{A_1}^{(T+2)} = p_{v_{B_3}} + \epsilon$. As B_3 cannot make any further bids, the SIA will terminate.

Now, let us assume that the statement is also true for $M - N - 2$ vehicles (and $N < M - 2$ parking lots).

To prove the statement for $M - N - 1$ vehicles and $N < M - 1$ parking lots, it is easy to see that after some $T - \delta$ steps, $M - N - 2$ vehicles with the $(M - N - 2)$ nd lowest (or the $(N + 2)$ nd highest) valuations will not be able to bid further, as $\forall i : p_i^{(T-\delta)} > p_{v_{(M-N-2)-}}^{(0)}$. After additional δ steps, the current prices will reach the valuation of the $(M - N - 1)$ st lowest (or the $(N + 1)$ st highest) vehicle: $\forall i : p_i^{(T)} = p_{v_{(M-N-1)-}}^{(0)} + \epsilon (= p_{v_{(N+1)+}}^{(0)} + \epsilon)$. Regarding the vehicle with the $(N + 1)$ st highest valuation is overbid or not, there are two possibilities.

If it is overbid and it cannot place any further bids, then the first N vehicles with the highest N valuations will win the SIA for the price of $p_{v_{(N+1)+}}^{(0)}$.

If the vehicle is currently not overbid, then, in the $(T + 1)$ st step, a vehicle with higher valuation shall pose another bid, causing a chain reaction, after which each of the first N vehicles with the highest valuations will win the auctions for $p_{v_{(N+1)+}}^{(0)} + \epsilon$. $\square$

APPENDIX C

PROOF OF THEOREM 4.

Proof. We will prove the statement by Petri nets [34]. A general Petri net could cover all transitions and places of a restricted, i.e., deterministically timed Petri nets that has identical structure and entry points to the general Petri net. Hence, it would be easy to see that the fully parallel SIA is a generalization of the serialized implementation if we can define two Petri nets having similar structures and initial positions: one general Petri net modeling the former, and a deterministic timed Petri net modeling the latter algorithm. In the following, we demonstrate how such Petri nets could be constructed.

To construct corresponding Petri nets, we shall track the bids and the possibility of bidding of the bidder agents. Let us model the possibility of bidding, i.e., the *Overbid* state of a bidder in Fig. 3, as B_j places in the Petri net. Moreover, $A_{i,j}$ will be places in the Petri net to represent that bidder j is currently in *Winning* state on auction i . Naturally, there shall be no bids on the auctions at the beginning; hence, *starting transitions* $t_{s,j}$ put exactly one token to each B_j places. In addition to the *starting* $t_{s,j}$ transitions, the rest of the transitions are defined between the two sets of places mentioned above. There are exactly two transitions between any B_j and $A_{i,j}$ pair of places. *Bidding transitions* $t_{bi,j}$ originated from the B_j places and pointing to $A_{i,j}$ places represent that the bidder j bids at the i auction. *Gating transitions* $t_{gi,j}$ originate from $A_{i,j}$ places and terminate in B_j places, represent the *tell_overbid* transitions in Fig. 3. To ensure that a bidder always wins at most one auction, and to also ensure that auctions run the LGB strategy for a single item demand, each bidding transition $t_{bi,j}$ passes tokens to

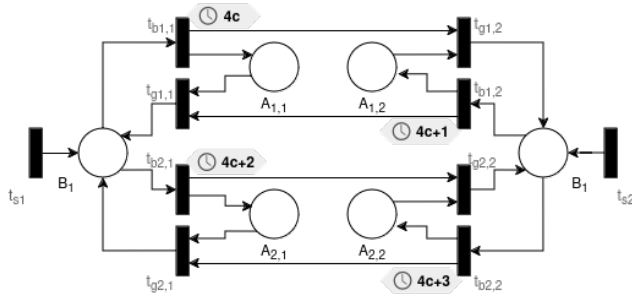
Attitude-driven Simultaneous Online Auctions
for Parking Spaces


Fig. 8. Structure of the Petri net models. With deterministic timing, it represents the serialized version. Without the timing, the Petri net corresponds to the fully parallel SIA execution with LGB strategy for 1-item demand.

each $t_{gi,k}$ gate transition, where $k \in \{1, 2, \dots, M\}$, $k \neq j$. The $t_{gi,k}$ gate transitions require at least two tokens on its inputs to release a single one. In this Petri net, if there is a token at the $A_{i,j}$ place, it means that the bidder j is currently winning on auction i . If there is a token at some B_j place, it means that the bidder j had been overbid, and it has not placed a newer bid since then. Fig. 8 demonstrates this Petri net structure for a simple $N = 2, M = 2$ case. This Petri net represents the SIA with an LGB strategy of [25].

Now, we shall focus on the necessary modifications of the untimed Petri net to represent the process defined in Algorithm 2. To model the iterations of line 6 and line 7 of Algorithm 2, we shall use a deterministic timed transitions Petri net. We assume that the transitions fire in an instant; however, they get enabled and disabled due to a periodic clock signal. Fundamentally, the outer iteration of line 6 prescribes that all transitions related to auction i shall precede the transitions of auction $i + 1$. Moreover, the inner iteration of line 7 requires that bidding transitions of bidder j shall precede the bidding transitions of bidder $j + 1$. Yielding bidding transitions $t_{bi,j}$ shall get enabled at exactly the $NMc + iM + j$, $c \in \{0, 1, \dots\}$ moments. As bidding transitions are required for firing the gating transitions, gating transitions can remain untimed ones. Fig. 8 shows an example of a simple $N = 2, M = 2$ case of the resulting timed Petri net.

As both Petri nets have a common structure and starting transitions, the timed Petri net is a restriction of the untimed version. Consequently, the untimed Petri net, corresponding to the method described in [25], can cover all transitions of the timed version. That yields that the parallelly running SIA with LGB strategy is a generalization of our serialized auction simulation method. $\square$

REFERENCES

- [1] D. Shoup, "Pricing curb parking," *Transportation Research Part A: Policy and Practice*, vol. 154, pp. 399–412, 2021, ISSN: 0965-8564. [DOI: 10.1016/j.tra.2021.04.012](#).
- [2] Z. Chen, Y. Yin, F. He, and J. L. Lin, "Parking reservation for managing downtown curbside parking," *Transportation Research Record*, vol. 2498, no. 1, pp. 12–18, 2015. [DOI: 10.3141/2498-02](#).
- [3] G. Pierce and D. S. and, "Getting the prices right," *Journal of the American Planning Association*, vol. 79, no. 1, pp. 67–81, 2013. [DOI: 10.1080/01944363.2013.787307](#).
- [4] J. Simićević, N. Milosavljević, G. Maletić, and S. Kaplanović, "Defining parking price based on users' attitudes," *Transport Policy*, vol. 23, pp. 70–78, 2012, ISSN: 0967-070X. [DOI: 10.1016/j.tranpol.2012.06.009](#).
- [5] F. Zong, P. Yu, J. Tang, and X. Sun, "Understanding parking decisions with structural equation modeling," *Physica A: Statistical Mechanics and its Applications*, vol. 523, pp. 408–417, 2019, ISSN: 0378-4371. [DOI: 10.1016/j.physa.2019.02.038](#).
- [6] P. A. Lopez, M. Behrisch, L. Bieker-Walz, et al., "Microscopic traffic simulation using SUMO," in *The 21st IEEE International Conference on Intelligent Transportation Systems*, IEEE, 2018.
- [7] E. F. and, "Zoning for parking as policy process: A historical review," *Transport Reviews*, vol. 24, no. 2, pp. 177–194, 2004. [DOI: 10.1080/0144164032000080485](#).
- [8] G. Mingardo, B. van Wee, and T. Rye, "Urban parking policy in europe: A conceptualization of past and possible future trends," *Transportation Research Part A: Policy and Practice*, vol. 74, pp. 268–281, 2015, ISSN: 0965-8564. [DOI: 10.1016/j.tra.2015.02.005](#).
- [9] W. Young, R. G. Thompson, and M. A. T. and, "A review of urban car parking models," *Transport Reviews*, vol. 11, no. 1, pp. 63–84, 1991. [DOI: 10.1080/01441649108716773](#).
- [10] P. L. Krapivsky and S. Redner, "Where should you park your car? The 1/2 rule," *Journal of Statistical Mechanics: Theory and Experiment*, vol. 2020, no. 7, p. 073 404, Jul. 2020. [DOI: 10.1088/1742-5468/ab96b7](#).
- [11] R. Arnott and J. Rowse, "Modeling parking," *Journal of Urban Economics*, vol. 45, no. 1, pp. 97–124, 1999, ISSN: 0094-1190. [DOI: 10.1006/juec.1998.2084](#).
- [12] E. Inci, "A review of the economics of parking," *Economics of Transportation*, vol. 4, no. 1, pp. 50–63, 2015, Special Issue on Collective Contributions in the Honor of Richard Arnott, ISSN: 2212-0122. [DOI: 10.1016/j.ecotra.2014.11.001](#).
- [13] C. Lei and Y. Ouyang, "Dynamic pricing and reservation for intelligent urban parking management," *Transportation Research Part C: Emerging Technologies*, vol. 77, pp. 226–244, 2017, ISSN: 0968-090X. [DOI: 10.1016/j.trc.2017.01.016](#).
- [14] A. Rahman, K. Rahman, and M. G. R. Alam, "Stable matching based parking lot allocation for autonomous vehicles," in *2021 2nd International Conference on Robotics, Electrical and Signal Processing Techniques (ICREST)*, 05-07 January, 2021, Dhaka, Bangladesh, 2021, pp. 372–376. [DOI: 10.1109/ICREST51555.2021.9331056](#).
- [15] T. Nakazato, Y. Fujimaki, and T. Namerikawa, "Parking lot allocation using rematching and dynamic parking fee design," *IEEE Transactions on Control of Network Systems*, vol. 9, no. 4, pp. 1692–1703, 2022. [DOI: 10.1109/TCNS.2022.3165015](#).
- [16] Y. Fujimaki, Y. Shibata, R. Namerikawa, and T. Namerikawa, "Decentralized optimal parking lot allocation by static and dynamic two-stage matching," *IFAC-PapersOnLine*, vol. 56, no. 2, pp. 689–694, 2023, 22nd IFAC World Congress, ISSN: 2405-8963. [DOI: 10.1016/j.ifacol.2023.10.1647](#).
- [17] I. R. Cohen, A. Eden, A. Fiat, and L. Jež, "Pricing online decisions: Beyond auctions," in *Proceedings of the 2015 Annual ACM-SIAM Symposium on Discrete Algorithms (SODA)*, January 4–6, 2015, San Diego, CA, USA, 2015, pp. 73–91. [DOI: 10.1137/1.9781611973730.7](#).
- [18] F. He, Y. Yin, Z. Chen, and J. Zhou, "Pricing of parking games with atomic players," *Transportation Research Part B: Methodological*, vol. 73, pp. 1–12, 2015, ISSN: 0191-2615. [DOI: 10.1016/j.trb.2014.12.003](#).
- [19] R. Arnott, "William vickrey: Contributions to public policy," *International Tax and Public Finance*, vol. 5, no. 1, pp. 93–113, 1998, ISSN: 1573-6970. [DOI: 10.1023/A:1008672627120](#).

- [20] H. Xiao, M. Xu, and Z. Gao, "Shared parking problem: A novel truthful double auction mechanism approach," *Transportation Research Part B: Methodological*, vol. 109, pp. 40–69, 2018, ISSN: 0191-2615. doi: 10.1016/j.trb.2018.01.008.
- [21] M. Cheng, E. Inci, S. X. Xu, and Y. Zhai, "A novel mechanism for private parking space sharing: The vickrey–clarke–groves auction with scale control," *Transportation Research Part C: Emerging Technologies*, vol. 150, p. 104 106, 2023, ISSN: 0968-090X. doi: 10.1016/j.trc.2023.104106.
- [22] S. Shao, S. X. Xu, H. Yang, and G. Q. Huang, "Parking reservation disturbances," *Transportation Research Part B: Methodological*, vol. 135, pp. 83–97, 2020, ISSN: 0191-2615. doi: 10.1016/j.trb.2020.03.005.
- [23] H. Yu, M. Huang, Y. Song, X. Wang, and X. Yue, "Making the most of your private parking slot: Strategy-proof double auctions-enabled staggered sharing schemes," *Transportation Research Part B: Methodological*, vol. 191, p. 103 102, 2025, ISSN: 0191-2615. doi: 10.1016/j.trb.2024.103102.
- [24] P. Wang, H. Guan, and P. Liu, "Modeling and solving the optimal allocation-pricing of public parking resources problem in urban-scale network," *Transportation Research Part B: Methodological*, vol. 137, pp. 74–98, 2020, Advances in Network Macroscopic Fundamental Diagram (NMFDD) Research, ISSN: 0191-2615. doi: 10.1016/j.trb.2019.03.003.
- [25] V. Bansal and R. Garg, "Simultaneous independent online auctions with discrete bid increments," *Electronic Commerce Research*, vol. 5, pp. 181–201, 2005. [Online]. Available: doi: 10.1007/s10660-005-6156-1.
- [26] S. R. Rizvi, S. Zehra, and S. Olariu, "MAPark: a multi-agent auction-based parking system in internet of things," *IEEE Intelligent Transportation Systems Magazine*, vol. 13, no. 4, pp. 104–115, 2021. doi: 10.1109/MTS.2019.2953524.
- [27] L. M. Ausubel and P. Cramton, "Vickrey auctions with reserve pricing," *Economic Theory*, vol. 23, no. 3, pp. 493–505, 2004, ISSN: 1432-0479. doi: 10.1007/s00199-003-0398-8.
- [28] A. Millard-Ball, R. R. Weinberger, and R. C. Hampshire, "The Shape of Cruising," *Findings*, Sep. 2021. doi: 10.32866/001c.28061.
- [29] R. Arnott and P. Williams, "Cruising for parking around a circle," *Transportation Research Part B: Methodological*, vol. 104, pp. 357–375, 2017, ISSN: 0191-2615. doi: 10.1016/j.trb.2017.07.009.
- [30] C. Holman, R. Harrison, and X. Querol, "Review of the efficacy of low emission zones to improve urban air quality in European cities," *Atmospheric Environment*, vol. 111, pp. 161–169, 2015, ISSN: 1352-2310. doi: 10.1016/j.atmosenv.2015.04.009.
- [31] L. Alekszejenkó and T. Dobrowiecki, "On vehicular data aggregation in federated learning. A case study of privacy with parking occupancy in Eclipse SUMO," in *SUMO User Conference Proceedings*, vol. Vol. 5 (2024). 2024, pp. 79–92. doi: 10.52825/scp.v5i.1100.
- [32] S. Saharan, S. Bawa, and N. Kumar, "Dynamic pricing techniques for intelligent transportation system in smart cities: A systematic review," *Computer Communications*, vol. 150, pp. 603–625, 2020, ISSN: 0140-3664. doi: 10.1016/j.comcom.2019.12.003.
- [33] L. Alekszejenkó and T. Dobrowiecki, "Effects of noisy occupancy data on an auction-based intelligent parking assignment," in *Proceedings of the 32nd Minisymposium of the Department of Artificial Intelligence and Systems Engineering Budapest University of Technology and Economics*, Feb. 3–4, 2025, Budapest, Hungary, 2025.
- [34] T. Murata, "Petri nets: Properties, analysis and applications," *Proceedings of the IEEE*, vol. 77, no. 4, pp. 541–580, 1989. doi: 10.1109/5.24143.



Levente Alekszejenkó was born in Fehérgyarmat, Hungary, in 1996. He received an M.Sc. degree in computer science engineering from the Budapest University of Technology and Economics, Budapest, Hungary, in 2021.

Since 2021, he has been a Ph.D. student at the Department of Artificial Intelligence and Systems Engineering, Budapest University of Technology and Economics, Budapest, Hungary. His research focuses on multi-agent intelligent systems in the transportation network.



Tadeusz Dobrowiecki was born in Warsaw, Poland, in 1952. He received the M.Sc. degree in electrical engineering from the Technical University of Budapest, Budapest, Hungary, in 1975, the Ph.D. (candidate of sciences) degree in 1981, and the DSc (Doctor of Sciences) degree in 2018, both from the Hungarian Academy of Sciences.

He joined the Department of Artificial Intelligence and Systems Engineering, Budapest University of Technology and Economics, Budapest, Hungary.

Currently he is a Professor Emeritus, teaching artificial intelligence and system identification.

Model-based mutation testing for Finite State Machine specifications with MTR

Gábor Árpád Németh

Abstract—In this article a model-based mutation testing technique has been introduced to a free and open-source model-based testing framework. The approach utilizes test suites that are generated with various test generation algorithms. It is investigated how efficient the resulting test suites are killing different types of mutations and their respective complexity. Guidelines are proposed to select the appropriate test generation methods for each mutation operator type independently. Using the ability to define a target score, the wide range of mutation generation and test generation options, one can create an appropriate trade off between fault coverage and test execution complexity.

Index Terms—model-based mutation testing, model-based testing, finite state machine, fault coverage, mutation operators

I. INTRODUCTION

In large software companies, while test execution is largely automated, the process of test design often remains a manual, thus resource-intensive and error-prone process. In Model-based testing (MBT) the product requirements are transformed into a formal specification model, from which test cases can be generated automatically according to some preset criteria [5]. This significantly reduces the cost and labor associated with traditional test design. However, selecting the right test generation algorithms that provide a proper trade off between the required fault detection capabilities and the complexity of the resulting test suite can be still a challenging task. To cope with this problem one can use model-based mutation testing (MBMT) approach [2], [3], [18], when different mutation operators are applied to the specification model itself to generate all possible mutated models which then can be used to select a test suite which has the desired fault coverage.

This article focuses on the MBMT of deterministic Finite State Machine (FSM) specifications which models have been extensively used in different problem domains such as telecommunication protocols and software [10], [11]. It is discussed, how *Model* $\gg$ *Test* $\gg$ *Relax*¹ (MTR) [20], a free and open source model-based tool can be used to discover the fault detection capabilities of different test generation algorithms for various types of mutation operators. It is also demonstrated how one can apply a MBMT technique in MTR to generate test suites that fulfill a given mutation target score for a given list of mutation operators. The main contribution of this article is that guidelines are proposed to select the appropriate test generation algorithms with their respective parameters for each mutation operator type separately that

result in the shortest test suite while providing a preset coverage of faults. Different strategies for the combinations of the above methods are also discussed.

The article is organized as follows. Section II discusses related terms regarding FSMs, MBT, MBMT, fault models and mutation operators. Section III overviews MBMT with the MTR framework. Section IV presents simulations to show the fault detection capabilities of different test generation approaches and to provide guidelines for selecting the most efficient ones for each mutation operator separately. The main results of the paper are concluded in Section V.

II. PRELIMINARIES

A. Finite State Machines

A Mealy Finite State Machine (FSM) M can be defined as $M = (I, O, S, T, s_0)$ where I , O , S and T are the finite and non-empty sets of *input symbols*, *output symbols*, *states* and *transitions* between states, respectively. s_0 denotes the *initial state*, where the machine starts from. Each transition $t \in T$ is a quadruple $t = (s_j, i, o, s_k)$, where $s_j \in S$ is the start state, $i \in I$ is an input symbol, $o \in O$ is an output symbol and $s_k \in S$ is the next state or end state.

FSM M is *deterministic*, if for each (s_j, i) state-input pair there exists at most one transition in T , otherwise it is *non-deterministic*. If there is at least one transition $t \in T$ for all state-input pairs, the machine is said to be *completely specified* (CS), otherwise it is *partially specified* (PS). In case of deterministic FSMs the output and the next state of a transition can be given as a function of the start state and the input of a transition, where $\lambda: S \times I \rightarrow O$ denotes the *output function* and $\delta: S \times I \rightarrow S$ denotes the *next state function*. Let us extend δ and λ from input symbols to finite *input sequences* I^* as follows: for a state s_1 , an input sequence $x = i_1, \dots, i_k$ takes the machine successively to states $s_{j+1} = \delta(s_j, i_j)$, $j = 1, \dots, k$ with the final state $\delta(s_1, x) = s_{k+1}$, and produces an *output sequence* $\lambda(s_1, x) = o_1, \dots, o_k$, where $o_j = \lambda(s_j, i_j)$, $j = 1, \dots, k$.

An FSM M is *strongly connected* iff for each pair of states (s_j, s_l) , there exists an input sequence which takes M from s_j to s_l . Two states, s_j and s_l of FSM M are *distinguishable*, iff there exists an $x \in I^*$ input sequence – called a *separating sequence* – that produces different output for these states, i.e.: $\lambda(s_j, x) \neq \lambda(s_l, x)$. Otherwise states s_j and s_l are *equivalent*. A machine is *reduced*, if no two states are equivalent.

An FSM M has a *reset message*, if there exists a special input symbol $r \in I$ that takes the machine from any state back to the s_0 initial state: $\exists r \in I : \forall s_j : \delta(s_j, r) = s_0$. The

Gábor Árpád Németh is with the Department of Computer Algebra, Faculty of Informatics, Eötvös Loránd University, Budapest, Hungary.
(e-mail: nga@inf.elte.hu)

¹ <https://modeltestrelax.org/>

reset is reliable if it is guaranteed to work properly in any implementation machine *Impl* of *M*.

B. Model-based testing

In FSM model-based testing (MBT) the requirements of the product are described as a specification FSM model *M*. The *test cases* – consisting of input sequences and their expected output sequences – are generated from *M* according to some preset criteria. The collection of test cases are called as *test suite*. To connect *M* to an actual System Under Test (SUT), a source code, called *adaptation code* is required that implements each transition of *M* as *keywords*. Utilizing the adaptation code, one can transform abstract test suites into executable ones to test the SUT². The SUT can be considered as a *black-box Impl* implementation machine of *M* with an unknown internal architecture, i.e. only its output responses to specific input sequences can be observed. *Conformance testing* verifies if the *observed output sequences* of *Impl* are equivalent to the *expected output sequences* derived from *M* – i.e. it determines if *Impl* conforms to *M*.

C. Mutation testing and model-based mutation testing

Mutation testing (MT) is a *white-box* testing technique, when various faulty system versions – called *mutants* – are created by applying simple modifications – called *mutation operators* – on the SUT itself that represents typical programming errors [17]. If a test can detect the modification from the original SUT, then it *kills* a given mutant, otherwise the given mutant remains *live*. A mutant can remain alive either if it is equivalent to the original SUT or if the test suite was unable to kill it. The efficiency of a test suite can be detected with a metric called *mutation score* that shows the ratio of the number of killed and the number of non-equivalent mutants. The MT technique relies on the *competent programmer hypothesis* [7] and the *coupling effect* [16]. The former one claims that experienced programmers tend to “create programs that are close to being correct”, i.e. it differs from the specification with relatively simple syntax and semantic faults [7]. The coupling effect assumes that those tests that identify all simple faults probably uncover more complex ones too [16].

In contrast to MT, in *Model-based mutation testing* (MBMT) the SUT is considered as a *black-box* i.e. its source code can not be mutated. Thus, a feasible approach is to introduce mutation operators in the formal system specification model itself [2], [3], [18]. In MBMT the following two different approaches can be used:

- In [3] a technique is presented, when a ts'_x test suite is generated from each M'_x mutated model and then executed on the SUT to be able to discover those faults that can not be identified by the *ts* test suite generated from the original, non-mutated *M* model. Note that as

²The time required for test execution is the function of the number of steps (i.e. the number of elements in the input/output sequence) in the test suite and the complexity of the adaptation code keywords. For example, in API (Application Programming Interface) testing, the adaptation code created for a given transition typically requires only a fraction of time to be executed compared to the one used in GUI (Graphical User Interface) testing.

test suites are generated from all mutant models independently and many possible non-equivalent mutations may exist for a given mutation operator type, the overall complexity of test generation and the resulting test suites can be enormous even in case of small specifications.

- Another method is when test suites are generated from the original model *M*, and then the given *ts* test suite is executed on all mutated models of M' . The fault detection capability of *ts* is determined by inspecting the amount of mutants it is able to kill, i.e. when *ts* provides an output executed on M'_x that differs from the one expected from *M* [2], [18]. The advantage of this approach compared to the first one is that the complexity of test generation and the resulting test suite is the fraction of the first one as the test suite is generated only once from model *M*. However, non-equivalent mutants may exist that can not be discovered from test suites generated just from *M*.

The current article focuses on the 2nd strategy.

D. FSM fault models and mutation operators

FSM *fault models* describe the assumptions of the test engineer about the FSM *Impl* (s)he is about to test. The following types of faults were introduced for CS, deterministic FSMs [6]:

- Output fault: for a given state-input pair, *Impl* produces an output that differs from the one specified in *M*.
- Transfer fault: for a given state-input pair, *Impl* goes into a state that differs from the one specified in *M*.
- Missing state: A state of *M* does not exist in *Impl*.
- Extra state: A non-specified state exists in *Impl*.

For PS, non-deterministic FSMs, these faults were extended with the following [4]:

- Missing transition: A transition of *M* does not exist in *Impl*.
- Extra transition: A non-specified transition exists in *Impl*.

These faults were extended later in other articles that dealt with FSM-based *mutation operators* [2], [12], [18]:

- Initial State Change: The state where the FSM starts from and reset leads to is changed to another state of *M*.
- Input change: The input of a transition is changed/deleted.

TABLE I
DIFFERENT TERMINOLOGIES FOR FSM FAULTS/MUTATION OPERATORS

[4], [6]	Used terminology in...		Notes
	[12]	[2], [18]	
-	Wrong-start-state	ISC (Initial State Changed)	-
-	Event-exchanged	COI (Change of Input)	May result in a non-det. FSM
-	Event-missing	MOI (Missing of Input)	Special case of COI. Contradicts to FSM formalism.
Output fault	Output-exchanged	COO (Change of Output)	-
-	Output-missing	MOO (Missing of Output)	Special case of COO
Transfer fault	Destination-exchanged	ESC (End State Changed)	May result in a non-strongly connected FSM
-	-	ROT (Reverse of Transition)	Combination of MOT and Extra Transition
Missing transition	Arc-missing	MOT (Missing of Transition)	May result in a non-strongly connected FSM
Extra transition	-	DOT (Redundant of Transition)	May result in a non-det. FSM
Missing state	State missing	MOS (Missing of State)	Associated transitions are also removed. May result in a non-strongly connected FSM
Extra state	State extra	-	Results in a non-strongly connected FSM
-	-	SSR (Start State Redundant)	Special case of extra state
-	-	ESR (End State Redundant)	Special case of extra state

Model-based mutation testing for Finite State Machine specifications with MTR

The used fault types of the referred papers and the differences in their terminologies are concluded in Table I.

Missing of Input (MOI) results in the changing of a state without receiving an actual input symbol which contradicts to the original FSM formalism. Change of Input (COI) may result in a non-deterministic FSM and in case of CS FSMs the set of input symbols must be extended to preserve determinism. Also COI can be created as a combination of Missing of Transition (MOT) and Extra Transition (ET).

An isolated, extra state (ES) can not be discovered with test suites generated from FSM M . Thus, the end state of an existing transition should be also modified to point to ES, and an ET that originates from ES should also be added.

Thus, the following mutation operators are considered: Initial State Changed (ISC), Change of Output (COO), Missing of Output (MOO), End State Changed (ESC), Missing of Transition (MOT), Extra Transition (ET), Missing of State (MOS), Extra State with 1-1 incoming/outgoing transitions (ES+).

III. MODEL-BASED MUTATION TESTING WITH MTR

Model $\gg$ *Test* $\gg$ *Relax* (MTR) offers a wide range of test generation algorithms for FSM-based specifications [20], this article focuses on the following:

- Transition Tour (TT) [14] for the *transition coverage* criteria, which generates the shortest test sequence that traverses all transitions of the specification, then it returns to the initial state.
- All-Transition-State (ATS) algorithm [19] for the ATS criteria [9]. This algorithm first traverses all transitions, then traverses all states again. It also creates alternative subsequences that try to reach all states in a way, that are as transition adjacent to each other as possible. In the current article ATS0, the non-iterative version of ATS is considered which provides 2 alternative subsequences.
- N-Switch Coverage (N-SC) algorithm [20] for the N-SC criteria [6], that covers all topologically possible, consecutive $N + 1$ transitions of the specification.
- Test generation using Harmonized State Identifiers (HSI) [13]. This algorithm creates a structured test suite that identifies all states of the machine, then all end states of transitions. The resulting test suite can be extended with sequences that guarantee to find a given θ number of extra states in *Impl* [8]. Note that in contrast with other algorithms listed above, HSI assumes that FSM M has reliable reset.

An overview of these algorithms which also compares their analytical and specific in MTR their practical complexities in detail (both for test generation and for the size of the resulting test suite) is presented in [20].

MBMT test generation can be selected in MTR, its pseudo code is described in Algorithm 1.

The inputs of the algorithm:

- The M specification FSM
- The Ω list of mutation operators, for which each ω_n element can either contain one (first order mutants) or more elements (higher order mutants)

Algorithm 1: Model-based mutation testing (MBMT)

```

input :  $M; \Omega = \{\omega_1, \dots, \omega_K\}; target\_score, ALG$ 
output:  $TS; mutation\_score$ 
1  $T := \{\}, TS := \{\}$ 
2 foreach  $\omega_n \in \Omega$  do
3    $M' := generateMutants(M, \omega_n)$ 
4   foreach  $ts \in T$  do
5      $execTestSuite(M', ts)$ 
6     if  $reachedTarget(M', ts, target\_score)$  and  $ts \notin TS$  then
7        $TS := TS \cup \{ts\}$ 
8   if  $\nexists ts \in T: reached\_target\_score$  for  $\omega_n$  then
9     while  $\neg reachedTarget(M', ts, target\_score)$ 
10       $\wedge \exists unused \in ALG$  do
11        $alg := First\ unused\ element\ of\ ALG$ 
12        $ts := generateTestSuite(M, alg)$ 
13        $T := T \cup \{ts\}$ 
14        $execTestSuite(M', ts)$ 
15       if  $reachedTarget(M', ts, target\_score)$  then
16          $TS := TS \cup \{ts\}$ 
17         break // goto line 2
18   if  $\nexists ts \in T: reached\_target\_score$  for  $\omega_n$  then
19      $cts := element\ of\ ALG\ with\ the\ highest\ score$ 
20      $TS := TS \cup \{cts\}$ 
21  $mutation\_score := calcMutScore(M, \Omega, TS)$ 
22 return  $TS, mutation\_score$ 
    
```

- The *target_score* the test engineer would like to achieve
- The *ALG* ordered list of applicable test generation algorithms with their parameters

The outputs of the algorithm:

- The *TS* set of test suites used to achieve *target_score*
- The actually achieved *mutation_score*

After initialization (line 1), the algorithm goes through the Ω list of mutation operators (line 2) and generates all possible M' mutant models for the given ω_n operator (line 3). Then it checks if there is a previously generated test suite ts in T (created for previous element(s) of Ω) that can be used to achieve *target_score* on all possible mutant models M' created for ω_n (lines 4-6). If such ts exists, ts is added to TS (line 7). If no such ts in T exist, then until *target_score* is reached and there exists any *alg* unused element of *ALG*, ts test suite is generated according to *alg* test generation algorithm and its parameters setting and ts is added to the T set of generated test suites (lines 8-12). Then it is checked if the test suite ts can be used to achieve *target_score*; if yes, ts is added to TS and process with next element of Ω (lines 13-16). If none of the test suites ts can be used to achieve *target_score*, then the algorithm simply adds the one that has the biggest score for ω_n (lines 17-19). If all elements of Ω list of mutation operators have been processed, the algorithm calculates the culminated mutation score (lines 20), returns the TS set of used test suites, the actually achieved *mutation_score* and terminates (line 21).

Note that by modifying the value of the *target_score*, the Ω list of mutation operators and the *ALG* ordered list of test generation algorithms (with their parameters), the test engineer has the ability to create a proper trade off between the desired fault coverage and the complexity of the generated test suites.

IV. SIMULATION RESULTS

Simulations were performed to find answers to the research questions related to MBMT in MTR below:

- Q1 What is the memory consumption of mutation generation in function of the size of the system and the type of the mutation operator?
- Q2 What is the fault coverage of the test suites of different test generation algorithms for different mutation types?
- Q3 What is the length of the test suites above?
- Q4 How different strategies for the *ALG* ordered list of applicable test generation algorithms perform best for different *target_scores* regarding the size of the resulting test sets?
- Q5 What are the answers to the Q2-Q4 questions for edge cases?
- Q6 How efficient are the different test suites to discover 2^{nd} order mutants?

TABLE II
INVESTIGATED SCENARIOS

ID	Number of states		size of step	I	O	section	Connecting to question
	min	max					
Scenario 1	5	50	5	5	10	IV-A1	Q1, Q2, Q3, Q4
Scenario 2	5	50	5	2	2	IV-A2	Q5
						IV-B	Q6
SIP UAC registration	4	-	-	11	3	IV-C	Q2

To answer the above questions, strongly connected, reduced, CS, deterministic random FSMs with reliable reset capability were generated with MTR in 2 scenarios and a small scale specification from the telecommunication domain was also investigated (see Table II). All possible ISC, COO, MOO, ESC, MOT, ET, MOS and ES+ mutant models were generated for each model, which keep the machine strongly connected and deterministic. The simulations were executed on servers running an Ubuntu 22.04.2 LTS operating system with 8 GB memory and one core of a shared AMD EPYC 7763 64-core CPU. Based on the results, guidelines are proposed in subsection IV-D to select the appropriate test generation methods for each mutation operator type independently.

A. First order mutations

1) *Scenario 1*: Consider Scenario 1, where the number of input and output symbols is 5 and 10, respectively.

Figure 1 shows the number of generated mutant models for each mutation operator. ES+ provided the most mutants as here the end state of all existing transitions can be changed to point to the extra state, and the input and the output of the new transition originating from the extra state can be selected from any combinations of the I/O sets of *M*. The memory consumption is shown in Table III. For ES+ mutants, the simulations could be performed only up to 15 states before the system ran out of the available 8 GB memory.

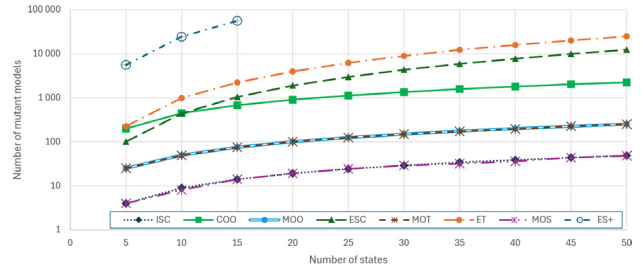


Figure 1. Scenario 1: Number of mutated models

TABLE III
SCENARIO 1: MEMORY CONSUMPTION OF MUTATED MODELS

Mutation operator	Memory consumption in GB for different number of states									
	5	10	15	20	25	30	35	40	45	50
ISC	0.00	0.00	0.00	0.01	0.01	0.01	0.01	0.01	0.01	0.01
COO	0.01	0.03	0.06	0.09	0.13	0.22	0.29	0.36	0.43	0.51
MOO	0.00	0.01	0.01	0.01	0.02	0.03	0.03	0.04	0.05	0.06
ESC	0.01	0.03	0.09	0.19	0.35	0.72	1.07	1.50	2.07	2.74
MOT	0.00	0.01	0.01	0.01	0.02	0.03	0.03	0.04	0.05	0.06
ET	0.01	0.06	0.19	0.41	0.72	1.48	2.23	3.14	4.26	5.60
MOS	0.00	0.00	0.00	0.01	0.01	0.01	0.01	0.01	0.01	0.01
ES+	0.18	1.36	4.90	-	-	-	-	-	-	-

TABLE IV
SCENARIO 1: MUTATION SCORE DISTRIBUTION FOR ALL MUTATED MODELS AND TEST GENERATION ALGORITHMS

Mutation operator	min	average	Mutation scores			max
			median	90th percentile		
ISC	1	1	1	1		1
COO	1	1	1	1		1
MOO	1	1	1	1		1
ESC	0.96	0.9971	1	0.9921		1
MOT	1	1	1	1		1
ET	0	0	0	0		0
MOS	1	1	1	1		1
ES+	0.9557	0.9937	0.9999	0.9773		1

TABLE V
SCENARIO 1: MUTATION SCORES FOR ESC

Test gen. alg.	Mutation scores for different number of states									
	5	10	15	20	25	30	35	40	45	50
TT	.96	.968	0.969	0.983	0.987	0.992	0.988	0.996	0.996	0.993
ATS0	1	.997	0.999	0.998	0.999	0.999	0.999	0.999	0.999	0.999
HSI	1	1	1	1	1	1	1	1	1	1
HSI ($\theta=1$)	1	1	1	1	1	1	1	1	1	1
1-SC	1	1	1	1	1	1	1	1	1	1
2-SC	1	1	1	1	1	1	1	1	1	1

TABLE VI
SCENARIO 1: MUTATION SCORES FOR ES+

Test generation algorithm	Mutation scores for different number of states		
	5	10	15
TT	0.9557	0.9773	0.9849
ATS0	0.9989	0.9995	0.9998
HSI	0.9813	0.9902	0.9994
HSI ($\theta=1$)	1	1	1
1-SC	1	1	1
2-SC	1	1	1

Test suites were generated with the TT, ATS0, 1-SC, 2-SC and HSI (with and without the extension to discover an extra state) test generation algorithms from the original, non-modified model *M*. The mutation detection capability of the generated test suites were investigated for all generated mutant models *M'* of each mutation operator type separately.

Table IV shows that all ISC, COO, MOO, MOT, MOS mutants were identified by all of the applied test generation algorithms, but none of the ET faults were discovered by any of these methods. The reason for the latter one is that by adding a new transition to a CS model, the *I* input symbol

Model-based mutation testing for Finite State Machine specifications with MTR

set of M should be extended to preserve determinism, but the new input symbol has not existed in the test suites generated from M , thus the new transition could not be selected.

The ESC and ES+ mutations were partially identified by the applied test suites, so the coverage of these faults were investigated for each test suite separately. Table V shows that HSI, 1-SC and 2-SC killed all ESC mutants, while ATSO nearly all of them and TT performed the worst. Table VI shows that 1-SC and 2-SC killed all of the ES+ mutants and ATSO nearly all of them. Interestingly, the test suite of the HSI-method has performed worse than ATSO to discover ES+ faults. As expected, when HSI was extended with the sequences designed to catch an extra state ($\theta=1$), it discovered all ES+ faults and TT performed the worst.

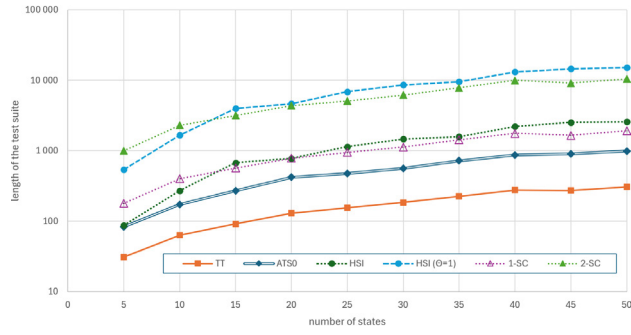


Figure 2. Scenario 1: Length of different test suites

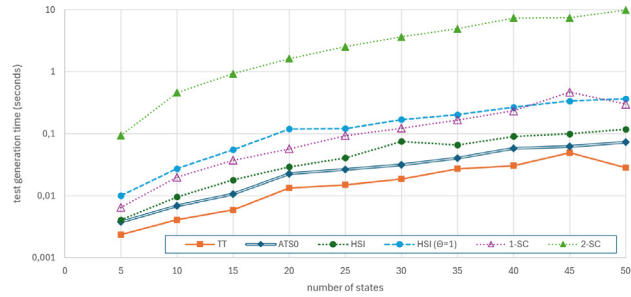


Figure 3. Scenario 1: Test generation time of different test suites

The length and the generation time of the different test suites are presented in Figures 2 and 3, respectively. As each test suite was generated from the original model M , its length and test generation time are independent from the type of the mutation operator, that was later applied on M .

TABLE VII

SCENARIO 1: LENGTH OF TEST SUITES USING DIFFERENT STRATEGIES TO KILL ALL, 99.8% AND 99% OF ESC MUTANTS

ALG / target score	Length of the test suite for different number of states									
	5	10	15	20	25	30	35	40	45	50
ATSO, HSI / 1	83	271	672	774	1139	1468	1583	2203	2523	2561
ATSO, 1-SC / 1	83	400	565	780	945	1124	1431	1763	1652	1907
ATSO, HSI / 0.998	83	271	271	420	474	561	719	864	905	989
ATSO, 1-SC / 0.998	83	400	271	420	474	561	719	864	905	989
ATSO, HSI / 0.99	83	173	271	420	474	561	719	864	905	989
ATSO, 1-SC / 0.99	83	173	271	420	474	561	719	864	905	989

The length of the resulting test suites were also investigated in the case when one was using different strategies for the ALG ordered list of applicable test generation algorithms (with their parameters) to kill all, 99.8% and 99% of ESC mutants. As Table VII shows, the $ALG=ATSO$, HSI and $ALG=ATSO$, 1-SC strategies result in almost the same length of test suites to kill all ESC mutants for Scenario 1. If $target_score$ was 0.998, the test suites of ATSO were included in the resulting TS set of test suites in all cases, but the one with 10 states. In case of $target_score = 0.99$, the test suites of ATSO were included in TS in all cases.

2) *Scenario 2*: An edge case was also investigated in Scenario 2, where both the number of input and output symbols were 2.

TABLE VIII

SCENARIO 2: MUTATION SCORE DISTRIBUTION FOR ALL MUTATED MODELS AND TEST GENERATION ALGORITHMS

Mutation operator	min	average	Mutation scores		
			median	90th percentile	max
ISC	0.7857	0.9765	1	0.9411	1
COO	1	1	1	1	1
MOO	1	1	1	1	1
ESC	0.8	0.9905	1	0.9743	1
MOT	1	1	1	1	1
ET	0	0	0	0	0
MOS	1	1	1	1	1
ES+	0.8142	0.9925	1	0.9855	1

One can observe in Table VIII, that due to less input and output symbols, the faults were harder to detect compared to Scenario 1. Again, all COO, MOO, MOT, MOS mutants were identified, but none of the ET faults were discovered by any of the test suites. In contrast to Scenario 1, not all ISC mutations are detected by all of the test suites. Also a higher rate of ESC faults remained undetected.

TABLE IX

SCENARIO 2: MUTATION SCORES FOR ISC

Test gen. algorithm	Mutation scores for different number of states									
	5	10	15	20	25	30	35	40	45	50
TT	1	1	0.7857	1	0.9583	1	0.9411	1	0.9772	1
ATSO	1	1	0.7857	1	0.9583	1	0.9411	1	0.9772	1
HSI	1	1	1	1	1	1	1	1	1	1
HSI ($\theta=1$)	1	1	1	1	1	1	1	1	1	1
1-SC	1	0.8888	0.8571	1	0.9583	1	1	0.9743	0.9545	1
2-SC	1	0.8888	0.8571	1	0.9583	1	1	0.9743	0.9545	1

The achieved mutation scores were investigated for each test suite separately for ISC, ESC and ES+ mutations. As shown in Table IX, only HSI was able to kill all ISC mutants and TT, ATSO, 1-SC and 2-SC were able to discover roughly the same amount of them. The reason for the former one is that HSI applied a reset symbol at the beginning of each test sequence and at the end of each test sequence the end state was verified. The reason for the latter one is that the test suites of TT, ATSO, 1-SC and 2-SC contain just one sequence, and the starting input subsequence x_0 of this sequence may enter the same state and produce the same output sequence for different initial states making it impossible to discover some ISC mutations.

The results for ESC are presented in Table X. All ESC mutations were killed by the HSI and 2-SC. The remaining test suites, ordered by efficiency, are: 1-SC, ATSO and TT. For ES+ mutations (Table XI) the HSI with the extension that is designed to catch 1 extra state and 2-SC were able

TABLE X
SCENARIO 2: MUTATION SCORES FOR ESC

Test gen. alg.	5	10	15	20	25	30	35	40	45	50
TT	0.8	0.978	0.899	0.973	0.957	0.969	0.968	0.978	0.974	0.983
ATS0	1	1	0.994	0.993	0.995	0.991	0.994	0.996	0.995	0.997
HSI	1	1	1	1	1	1	1	1	1	1
HSI ($\theta=1$)	1	1	1	1	1	1	1	1	1	1
1-SC	1	1	0.997	0.998	0.998	0.997	0.999	0.999	0.999	0.999
2-SC	1	1	1	1	1	1	1	1	1	1

TABLE XI
SCENARIO 2: MUTATION SCORES FOR ES+

Test gen. alg.	5	10	15	20	25	30	35	40	45	50
TT	.814	.966	0.931	0.984	0.975	0.987	0.985	0.988	0.990	0.991
ATS0	.971	.995	0.996	0.995	0.997	0.996	0.998	0.997	0.998	0.998
HSI	1	.990	1	1	1	1	1	1	1	1
HSI ($\theta=1$)	1	1	1	1	1	1	1	1	1	1
1-SC	1	1	1	1	1	1	1	0.999	1	1
2-SC	1	1	1	1	1	1	1	1	1	1

to discover all mutations. Even HSI without this extension and 1-SC discovered all ES+ faults in almost all cases. At and above 10 states ATS0 was able to catch at least 99.5% of ES+ faults. As expected, the TT was again the least effective to kill ES+ mutations.

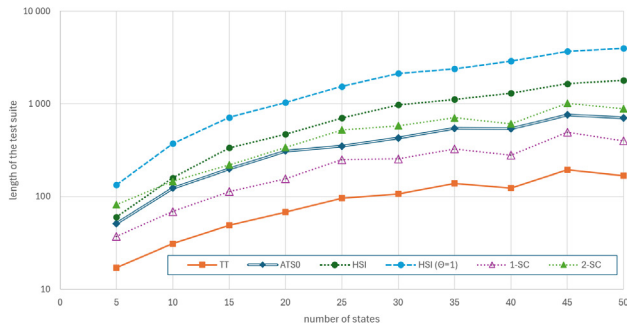


Figure 4. Scenario 2: Length of different test suites

The length of the different test suites is shown in Figure 4.

TABLE XII
SCENARIO 2: LENGTH OF TEST SUITES USING DIFFERENT STRATEGIES TO KILL ALL, 99.8% AND 99% OF ESC MUTANTS

ALG / target score	5	10	15	20	25	30	35	40	45	50
ATS0, HSI / 1	51	123	336	470	706	973	1116	1304	1650	1790
ATS0, 1-SC, 2-SC / 1	51	123	308	423	577	656	822	732	1146	1046
ATS0, HSI / 0.998	51	123	336	470	706	973	1116	1304	1650	1790
ATS0, 1-SC, 2-SC / 0.998	51	123	308	155	250	656	326	279	494	400
ATS0, HSI / 0.99	51	123	199	310	350	428	543	537	758	705

Different strategies were investigated for the ALG ordered list to kill all, 99.8% and 99% of ESC mutants in Table XII and to kill all and 99% of ES+ mutants in Table XIII.

B. Second order mutations

Second order mutation operators generate much more mutated models, than first order ones, thus simulations could be performed only for the subset of states even for Scenario 2. The mutation scores of different test suites for ES+, ISC and

TABLE XIII
SCENARIO 2: LENGTH OF TEST SUITES USING DIFFERENT STRATEGIES TO KILL ALL AND 99% OF ES+ MUTANTS

ALG / target score	5	10	15	20	25	30	35	40	45	50
ATS0, HSI, HSI ($\theta=1$) / 1	60	372	336	470	706	973	1116	1304	1650	1790
ATS0, 1-SC, 2-SC / 1	37	69	113	155	250	255	326	732	494	400
1-SC, HSI, HSI ($\theta=1$) / 1	37	69	113	155	250	255	326	1304	494	400
ATS0, HSI, HSI ($\theta=1$) / 0.99	60	123	199	310	350	428	543	537	758	705
1-SC, HSI, HSI ($\theta=1$) / 0.99	37	69	113	155	250	255	326	279	494	400

TABLE XIV
SCENARIO 2: MUTATION SCORES FOR ES+, ISC 2nd ORDER MUTANTS

Test generation algorithm	5	10	15	20	25
TT	0.9957	0.9996	0.9784	0.9999	0.9989
ATS0	0.9985	1	0.9992	1	0.9999
HSI	1	1	1	1	1
HSI ($\theta=1$)	1	1	1	1	1
1-SC	1	0.9996	0.9999	1	1
2-SC	1	1	1	1	1

TABLE XV
SCENARIO 2: MUTATION SCORES FOR ESC, ISC 2nd ORDER MUTANTS

Test generation algorithm	5	10	15	20	25	30	35
TT	1	1	0.9704	1	0.9982	0.9999	0.9971
ATS0	1	1	0.9989	1	0.9998	1	0.9997
HSI	1	1	1	1	1	1	1
HSI ($\theta=1$)	1	1	1	1	1	1	1
1-SC	1	1	0.9991	1	0.9999	1	1
2-SC	1	1	1	1	1	1	1

ESC, ISC second order mutants are presented in Table XIV and XV, respectively. HSI and 2-SC test suites were able to kill all of these mutants. The remaining test suites, ordered by efficiency, are: 1-SC, ATS0 and TT.

C. SIP UAC registration example

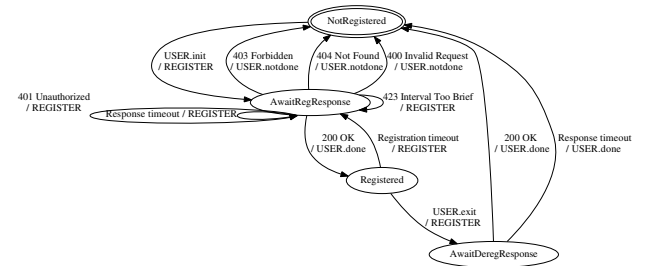


Figure 5. FSM for the registration process of the SIP user agent client

The mutation scores of the TT, ATS0, HSI, 1-SC and 2-SC test suites for different mutation operators for a specification FSM presented in Figure 5 were also observed. This FSM describes the registration process of the SIP (Session Initiation Protocol) User Agent Client (UAC) [1]³.

Before the results presented in Table XVI are discussed, a quick overview of ESC and ES+ faults are given. An ESC bug results to enter an invalid state. For example, the UAC assumes that its REGISTER request given for the 401 Unauthorized

³Here only the signaling level was considered; a step-by-step description of the construction of this FSM from call-flows can be found in [15].

Model-based mutation testing for Finite State Machine specifications with MTR

TABLE XVI
MUTATION SCORES FOR THE SIP UAC REGISTRATION EXAMPLE

Mutation operator	Mutation scores for PS / CS machines...					
	TT	ATSO	HSI	HSI ($\theta=1$)	1-SC	2-SC
ISC	1 / 1	1 / 1	1 / 1	1 / 1	1 / 1	1 / 1
COO	1 / 1	1 / 1	1 / 1	1 / 1	1 / 1	1 / 1
MOO	1 / 1	1 / 1	1 / 1	1 / 1	1 / 1	1 / 1
ESC	0.8928 / 0.9758	1 / 1	0.0357 / 1	1 / 1	1 / 1	1 / 1
MOT	1 / 1	1 / 1	1 / 1	1 / 1	1 / 1	1 / 1
ET	0 / 0	0 / 0	0 / 0	0 / 0	0 / 0	0 / 0
MOS	1 / 1	1 / 1	1 / 1	1 / 1	1 / 1	1 / 1
ES+	0.8939 / 0.9707	0.9939 / 0.9948	0.1 / 0.9951	0.9848 / 1	0.9939 / 1	0.9962 / 1

answer was successful without waiting for the 200 OK answer from the server, thus it enters the state *Registered* instead of *AwaitRegResponse*. ES+ mutations introduce a new functionality in the implementation, that was not specified. For instance, instead of returning to the *NotRegistered* state after an 400 Invalid Request has been received, the UAC goes to a distinct *Error* state.

Note that two different approaches were considered for test suite generation. In the first case, the test suites were generated directly from the original PS model presented in Figure 5. In the second case the former PS machine has been converted into a CS one by MTR by adding a loop transition without an output symbol for each undefined state-input symbol pair and test suites were generated from this CS FSM.

The actually achieved mutation scores for the PS and CS machines (see Table XVI) highly differ from each other for ESC and ES+ mutations in case of HSI test generation. The reason is the following: Most of the input symbols were defined only for 1 transition in case of the PS machine. Due to this, HSI was unable to find separating sequences⁴, resulting in that the structured test suite of HSI was unable to perform its state verification process at the end of each sequence. Thus, HSI could not identify ESC and ES+ mutations. However, this problem could be easily fixed by converting the PS machine into a CS one; in this case the test suite of HSI was able to apply state verifications at the end of each sequence. The test suite of HSI ($\theta=1$) contains longer subsequences which made it less sensitive to whether the FSM was partially or completely specified. As TT, ATSO, 1-SC and 2-SC provide just one sequence in their test suite, they were also affected much less by the completeness of the machine. But the fault detection capabilities of all methods were improved when the PS FSM was converted into a CS one.

D. Guidelines for selecting test generation algorithms

Based on the achieved *mutation_scores* and the length of the resulting test suites observed previously, our proposals for selecting the appropriate test generation algorithms for different types of mutation operators are summarized in Table XVII. Note that as COO, MOO, MOT, MOS mutations can be discovered by just traversing all transitions of FSM M , here the application of the shortest TT test suite is advised as it provides transition coverage. Due to its state verification part, HSI guarantees to find all ISC and ESC mutants in case of

TABLE XVII
GUIDELINES FOR TEST GENERATION ALGORITHMS FOR DIFFERENT MUTATION OPERATORS

Mutation operator	Proposed test generation algorithm for...		Notes
	complete coverage	good balance	
ISC	HSI	TT	HSI for CS FSMs with reliable reset or if a separating sequence exist for all state pairs ^{4,5} .
COO	TT	TT	-
MOO	TT	TT	-
ESC	HSI / (2-SC)	ATSO, 1-SC	HSI for CS FSMs with reliable reset or if a separating sequence exist for all state pairs ^{4,5} . 1-SC over ATSO is advised for sparse FSMs. No proof for 100% coverage of 2-SC.
MOT	TT	TT	-
ET	-	-	This type of fault can not be discovered.
MOS	TT	TT	-
ES+	HSI ($\theta=1$) / (2-SC)	ATSO, HSI / 1-SC	HSI ($\theta=1$) for CS FSMs with reliable reset or if a separating sequence exist for all state pairs ^{4,5} . No proof for 100% coverage of 2-SC. 1-SC & 2-SC over other options are advised for sparse FSMs. HSI over ATSO is advised for higher coverage.

CS machines⁵ or if a separating sequence exists for all pairs of states in case of PS machines⁴. HSI also assumes that the SUT has reliable reset. If these assumptions of HSI can not be fulfilled, for sparse models (where $|T| < 5 \cdot |S|$, that is $|I| < 5$ in case of CS FSMs) 2-SC can be used, as it found all ESC faults in our simulations, but there is no analytical proof for complete coverage. Although ATSO does not guarantee to find all ESCs, but can be a proper trade off between fault coverage and the length of the test suite as it discovers most of the ESC mutants with a fraction of the length of HSI and 2-SC test suites. For sparse models (where $|T| < 5 \cdot |S|$) 1-SC can be a suitable option over ATSO to find most of the ESCs. In edge cases, some ISCs may be undiscovered by TT, ATSO, 1-SC and 2-SC, but as they provide roughly the same fault detection capability, the shortest TT is proposed as a trade off between ISC coverage and the complexity of the test suite. For ES+, one can extend the HSI with extra state searching part ($\theta=1$) that kills all mutants in case of CS FSMs or if a separating sequence exists for all pairs of states in case of PS machines. For sparse models 2-SC can be also applied, as it also discovered all mutants in our simulations, but there is no analytical proof for complete coverage of ES+ mutations. The ATSO and HSI without the extra state extension can also be a suitable option to cover the most of ES+ faults, the latter one is proposed for a higher fault coverage at the cost of a longer test suite. Alternatively, one can use 1-SC for less dense models to find most of the ES+ mutants. As discussed previously, ET mutations can not be discovered with test suites which are generated from the original, non-mutated specification M .

If one uses MBMT with MTR, for the ALG ordered list of test generation algorithms the following strategies are proposed to discover different mutation types. For COO, MOO, MOT, MOS mutation operators: $ALG=TT$. For ESC faults: $ALG=1-SC$, HSI for sparse FSMs and $ALG=ATSO$, HSI for more dense FSMs. If the FSM does not have reliable reset, then $ALG=1-SC$, 2-SC strategy can also be applied for ESCs in case of sparse models. For ISC mutations $ALG=TT$, HSI is advised (or if HSI not applicable, just $ALG=TT$). For ES+ mutations $ALG=1-SC$, 2-SC, HSI ($\theta=1$) for sparse FSMs and $ALG=ATSO$, HSI ($\theta=1$) for more dense FSMs are proposed.

Note that the processing order for the Ω list of mutation operators are not optimized in line 2 of Algorithm 1. Thus,

⁴Note that the HSI test generation implemented in MTR throws a warning if separating sequence does not exist for a given state pair.

⁵Note that PS machines can be easily converted into CS ones.

if one would like to generate a set of test suites that is able to catch different mutation types this should be taken into account. One should apply those mutations first in the Ω ordered list of mutation operators, that are harder to detect to avoid adding a weaker test suite to the test set before adding a stronger one that would make the former one unnecessary. According to presented results, in case of first order mutants, the following order is suggested in Ω : (1) ES+, (2) ESC, (3) ISC, (4) MOT/MOS/COO/MOO.

Using the results discussed previously, one can set a *target_score* that is to be achieved for a given Ω list of mutation operators and the proper *ALG* ordered list of applicable test generation algorithms to fulfill the desired coverage with the lowest possible length in the resulting test set of test suites.

V. CONCLUSION

In the current paper it is presented how model-based mutation testing can be applied for finite state machine specifications in the free and open source *Model $\gg$ Test $\gg$ Relax* model-based testing framework. The test engineer can set the list of different types of first or higher order mutants (s)he interested in, the ordered list of test generation algorithms (with their parameters) to be applied from a wide list of options and a target mutation score which is to be achieved.

The memory consumption of mutation generation, the fault coverage of the different test generation strategies for different mutation operators and the length of the resulted test suites were investigated by simulations including first and second order mutants. Guidelines were also given for selecting the appropriate test generation algorithm and the ordered list of these methods with their respective parameters for each mutation operator separately. A processing order for different types of mutation operators for efficient test suite generation is also proposed if one would like to cover multiple types of mutations. Using the above guidelines with the wide range of setting possibilities one can create an appropriate trade off between fault coverage and the size of the resulting test set of test suites.

ACKNOWLEDGEMENTS

The author would like to thank the students who took part in implementation of the following parts of the framework: Máté István Lugosi for the TT and ATS test generation, for the simulation script, and for the basic change injection functionality, Tódor Dávid Nyeste for the HSI-method, Bálint Miksa Rumpler for the N-SC test generation, the extension of change injection and for the mutation testing logic above all test generation algorithms.

REFERENCES

- [1] RFC 3261: SIP: Session Initiation Protocol, 2002. <https://tools.ietf.org/html/rfc3261> Accessed: 2025-07-10.
- [2] Danial Nikbin Azmoudeh and Yvan Labiche. Analysis of mutation operators for FSM testing. In *2023 IEEE International Conference on Software Testing, Verification and Validation Workshops (ICSTW)*, pages 300–307, 2023. **doi:** 10.1109/ICSTW58534.2023.00060.
- [3] Fevzi Belli, Christof J. Budnik, Axel Hollmann, Tugkan Tuglular, and W. Eric Wong. Model-based mutation testing—Approach and case studies. *Science of Computer Programming*, 120:25–48, 2016. **doi:** 10.1016/j.scico.2016.01.003.
- [4] Gregor von Bochmann, Anindya Das, Rachida Dssouli, Martin Dubuc, Abderrazak Ghedamsi, and Gang Luo. Fault Models in Testing. In *Proceedings of the IFIP TC6/WG6.1 Fourth International Workshop on Protocol Test Systems IV*, pages 17–30, Amsterdam, The Netherlands, 1991. North-Holland Publishing Co.
- [5] Manfred Broy, Bengt Jonsson, Joost-Pieter Katoen, Martin Leucker, and Alexander Pretschner (Eds.). *Model-Based Testing of Reactive Systems*. Springer, 2005. **doi:** 10.1007/b137241.
- [6] T. Chow. Testing software design modelled by finite-state machines. *IEEE Transactions on Software Engineering*, 4(3):178–187, May 1978. **doi:** 10.1109/TSE.1978.231496.
- [7] R. A. DeMillo, R. J. Lipton, and F. G. Sayward. Hints on Test Data Selection: Help for the Practicing Programmer. *Computer*, 11(4):34–41, 1978. **doi:** 10.1109/C-M.1978.218136.
- [8] Rita Dorofeeva, Khaled El-Fakih, and Nina Yevtushenko. An Improved Conformance Testing Method. In Farn Wang, editor, *Formal Techniques for Networked and Distributed Systems – FORTE 2005*, volume 3731 of *Lecture Notes in Computer Science*, pages 204–218. Springer, Berlin, Heidelberg, 2005. **doi:** 10.1007/11562436_16.
- [9] István Forgács and Attila Kovács. *Practical Test Design*. BCS, The Chartered Institute for IT, 2019.
- [10] Drago Hercog. Protocol Specification and Design. In *Communication Protocols*. Springer, Cham, 2020. **doi:** 10.1007/978-3-030-50405-2_2.
- [11] Gerard J. Holzmann. *Design and Validation of Protocols*. Prentice-Hall, 1990.
- [12] Yue Jia and Mark Harman. An Analysis and Survey of the Development of Mutation Testing. *IEEE Transactions on Software Engineering*, 37(5):649–678, 2011. **doi:** 10.1109/TSE.2010.62.
- [13] Gang Luo, Alexandre Petrenko, and Gregor V. Bochmann. Selecting Test Sequences for Partially-Specified Nondeterministic Finite State Machines. In *Proceedings of the IFIP WG6.1 7th International Workshop on Protocol Test systems VI*, pages 91–106. Springer, 1995. **doi:** 10.1007/978-0-387-34883-4_6.
- [14] S. Naito and M. Tsunoyama. Fault detection for sequential machines by transition-tours. In *Proceedings of the 11th IEEE Fault-Tolerant Computing Conference (FTCS 1981)*, pages 238–243. IEEE Computer Society Press, 1981.
- [15] Gábor Árpád Németh and Péter Sótér. Teaching performance testing. *Teaching Mathematics and Computer Science*, 19(1):17–33, 2021. **doi:** 10.5485/TMCS.2021.0518.
- [16] A. Jefferson Offutt. Investigations of the software testing coupling effect. *ACM Trans. Softw. Eng. Methodol.*, 1(1):5–20, January 1992. **doi:** 10.1145/125489.125473.
- [17] Mike Papadakis, Marinos Kintis, Jie Zhang, Yue Jia, Yves Le Traon, and Mark Harman. Chapter six - mutation testing advances: An analysis and survey. volume 112 of *Advances in Computers*, pages 275–378. Elsevier, 2019. **doi:** 10.1016/bs.adcom.2018.03.015.
- [18] S. C. Pinto Ferraz Fabbri, M. E. Delamaro, J. C. Maldonado, and P. C. Masiero. Mutation analysis testing for finite state machines. In *Proceedings of 1994 IEEE International Symposium on Software Reliability Engineering*, pages 220–229, 1994. **doi:** 10.1109/ISSRE.1994.341378.
- [19] Gábor Árpád Németh and Máté István Lugosi. Test generation algorithm for the All-Transition-State criteria of Finite State Machines. *Infocommunications Journal*, 13(3):56–65, 2021. **doi:** 10.36244/ICJ.2021.3.6.
- [20] Gábor Árpád Németh and Máté István Lugosi. MTR Model-Based Testing Framework. *Infocommunications Journal*, 16(2):11–18, 2024. **doi:** 10.36244/ICJ.2024.2.2.



Gábor Árpád Németh obtained his MSc in Electrical Engineering and his PhD in Computer Science at the Budapest University of Technology and Economics (BME), Department of Telecommunication and Media Informatics (TMIT) in 2007 and 2015, respectively. He worked at Ericsson between 2011 and 2018 on a performance testing tool used in the telecommunication industry. Currently, he works at the Eötvös Loránd University (ELTE) on topics related to requirement engineering and software testing.

Guidelines for our Authors

Format of the manuscripts

Original manuscripts and final versions of papers should be submitted in IEEE format according to the formatting instructions available on

<https://journals.ieeeauthorcenter.ieee.org/>
Then click: "IEEE Author Tools for Journals"
- "Article Templates"
- "Templates for Transactions".

Length of the manuscripts

The length of papers in the aforementioned format should be 6-8 journal pages.

Wherever appropriate, include 1-2 figures or tables per journal page.

Paper structure

Papers should follow the standard structure, consisting of *Introduction* (the part of paper numbered by "1"), and *Conclusion* (the last numbered part) and several *Sections* in between.

The Introduction should introduce the topic, tell why the subject of the paper is important, summarize the state of the art with references to existing works and underline the main innovative results of the paper. The Introduction should conclude with outlining the structure of the paper.

Accompanying parts

Papers should be accompanied by an *Abstract* and a few *Index Terms (Keywords)*. For the final version of accepted papers, please send the short cvs and *photos* of the authors as well.

Authors

In the title of the paper, authors are listed in the order given in the submitted manuscript. Their full affiliations and e-mail addresses will be given in a footnote on the first page as shown in the template. No degrees or other titles of the authors are given. Memberships of IEEE, HTE and other professional societies will be indicated so please supply this information. When submitting the manuscript, one of the authors should be indicated as corresponding author providing his/her postal address, fax number and telephone number for eventual correspondence and communication with the Editorial Board.

References

References should be listed at the end of the paper in the IEEE format, see below:

- a) Last name of author or authors and first name or initials, or name of organization
- b) Title of article in quotation marks
- c) Title of periodical in full and set in italics
- d) Volume, number, and, if available, part
- e) First and last pages of article
- f) Date of issue
- g) Document Object Identifier (DOI)

[11] Boggs, S.A. and Fujimoto, N., "Techniques and instrumentation for measurement of transients in gas-insulated switchgear," *IEEE Transactions on Electrical Installation*, vol. ET-19, no. 2, pp.87–92, April 1984. DOI: 10.1109/TEI.1984.298778

Format of a book reference:

[26] Peck, R.B., Hanson, W.E., and Thornburn, T.H., *Foundation Engineering*, 2nd ed. New York: McGraw-Hill, 1972, pp.230–292.

All references should be referred by the corresponding numbers in the text.

Figures

Figures should be black-and-white, clear, and drawn by the authors. Do not use figures or pictures downloaded from the Internet. Figures and pictures should be submitted also as separate files. Captions are obligatory. Within the text, references should be made by figure numbers, e.g. "see Fig. 2."

When using figures from other printed materials, exact references and note on copyright should be included. Obtaining the copyright is the responsibility of authors.

Contact address

Authors are requested to submit their papers electronically via the following portal address:

https://www.ojs.hte.hu/infocommunications_journal/about/submissions

If you have any question about the journal or the submission process, please do not hesitate to contact us via e-mail:

Editor-in-Chief: Pál Varga – pvarga@tmit.bme.hu

Associate Editor-in-Chief:

József Bíró – biro@tmit.bme.hu

László Bacsárdi – bacsardi@hit.bme.hu



29th Conference on Innovation in Clouds, Internet and Networks

Disruptive Intelligent Technologies for Next-Generation Clouds and Networks

30 March – 02 April, 2026
Athens, Greece



Call For Papers

Digital transformation enhances both economic development and the quality of life for citizens. This evolution is based on the deployment of a network, computing and service infrastructure that is highly sophisticated in terms of its size, the diversity of its components and the range of specialized uses it allows. Such an infrastructure is made of a large set of heterogeneous hardware/software components resulting in a system of systems whose availability, reliability, performance, interoperability and energy efficiency are becoming major challenges. Disruptive intelligent technologies are key enablers for this evolution, driving automation, enabling early detection of faults and cyberattacks, and optimizing resource usage to meet stringent service performance and resource requirements.

A special focus will be set on the conference theme “Disruptive Intelligent Technologies for Next-Generation Clouds and Networks”.

Areas of interest and topics include, but are not limited to:

Intelligent network service management

- Intelligent end-to-end service configuration
- Intelligent and cooperative management models across heterogeneous network segments from IoT to Cloud and back
- Intelligent, resilient and robust schemes in the IoT-edge-cloud continuum service management
- Enablers for cross-network intelligence and services trustworthiness
- Intelligent and flexible networks and services for high adaptability
- Digital twin for network service management
- Distributed intelligence in B5G/6G networks

IoT-edge-cloud continuum

- Intelligent edge for vertical industries
- Digital twin challenges for the IoT-edge-cloud continuum and life cycle automation
- Vehicular cloud services
- Industrial networks and services for an intelligent automation
- Functional decomposition and orchestration, service chaining
- Communication protocols for ultra-low latency and ultra-high reliability
- Fog, edge and multi-access computing and networking
- Security, trust and privacy for the IoT-edge-cloud continuum

Enablers and services for intelligent wireless and mobile networks

- Network and service architecture, protocols, application-network layer interworking
- Scalability and multi-tenancy in beyond wireless, cellular and emerging networks
- Green networking and efficient resource usage and service delivery
- Security, trust and privacy management in large systems and smart environments

Network automation and orchestration architectures and protocols

- Zero-touch network and service management and orchestration
- Artificial Intelligence for networks and networks for Artificial Intelligence
- Orchestration of mobile and ephemeral resources, automated infrastructure discovery
- Network data analytics, anomaly detection and feature extraction
- Autonomic and cognitive networking
- Automated data, control, and management planes
- Intent-based network management
- Self-driving networks
- Network programmability
- Automated QoS/QoE management
- Orchestration of distributed cloud, edge and network resources
- Security, societal, and legal aspects of network automation
- Federated learning-based applications
- Federated Networks

Technologies for innovation in clouds, edge, networks and IoT

- SDN, NFV, service function chaining, function placement and network embedding
- Serverless computing and FaaS
- Application-aware networking
- Blockchain and Distributed Ledger Technologies in networking and network services
- Advanced multimedia and real-time communications
- QoE/QoS Assurance
- Lightweight virtualization technologies, software and hardware acceleration
- Energy-efficient software-defined infrastructures
- Measurement, monitoring and telemetry
- Time-sensitive and deterministic networking
- Open-source projects, testbeds and open standards

Important dates

Paper Submission Due
October 31, 2025

Notification of Acceptance
December 20, 2025

Camera-Ready Papers due
January 8, 2026

Conference Date
March 30 - April 2, 2026

General co-Chairs:

Molka Gharbaoui
(Scuola Sant'Anna Pisa, Italy)

Christos Douligeris
(University of Piraeus, Greece)

TCP co-Chairs:

Eliftherios Mamatas
(University of Macedonia, Greece)

Giovanni Schembra
(University of Catania, Italy)

<https://www.icin-conference.org/index.php/call-for-papers/>



Who we are

Founded in 1949, the Scientific Association for Infocommunications (formerly known as Scientific Society for Telecommunications) is a voluntary and autonomous professional society of engineers and economists, researchers and businessmen, managers and educational, regulatory and other professionals working in the fields of telecommunications, broadcasting, electronics, information and media technologies in Hungary.

Besides its 1000 individual members, the Scientific Association for Infocommunications (in Hungarian: HÍRKÖZLÉSI ÉS INFORMATIKAI TUDOMÁNYOS EGYESÜLET, HTE) has more than 60 corporate members as well. Among them there are large companies and small-and-medium enterprises with industrial, trade, service-providing, research and development activities, as well as educational institutions and research centers.

HTE is a Sister Society of the Institute of Electrical and Electronics Engineers, Inc. (IEEE) and the IEEE Communications Society.

What we do

HTE has a broad range of activities that aim to promote the convergence of information and communication technologies and the deployment of synergic applications and services, to broaden the knowledge and skills of our members, to facilitate the exchange of ideas and experiences, as well as to integrate and

harmonize the professional opinions and standpoints derived from various group interests and market dynamics.

To achieve these goals, we...

- contribute to the analysis of technical, economic, and social questions related to our field of competence, and forward the synthesized opinion of our experts to scientific, legislative, industrial and educational organizations and institutions;
- follow the national and international trends and results related to our field of competence, foster the professional and business relations between foreign and Hungarian companies and institutes;
- organize an extensive range of lectures, seminars, debates, conferences, exhibitions, company presentations, and club events in order to transfer and deploy scientific, technical and economic knowledge and skills;
- promote professional secondary and higher education and take active part in the development of professional education, teaching and training;
- establish and maintain relations with other domestic and foreign fellow associations, IEEE sister societies;
- award prizes for outstanding scientific, educational, managerial, commercial and/or societal activities and achievements in the fields of infocommunication.

Contact information

President: **FERENC VÁGUJHELYI** • elnok@hte.hu

Secretary-General: **GÁBOR KOLLÁTH** • kollath.gabor@hte.hu

Operations Director: **PÉTER NAGY** • nagy.peter@hte.hu

Address: H-1051 Budapest, Bajcsy-Zsilinszky str. 12, HUNGARY, Room: 502

Phone: +36 1 353 1027

E-mail: info@hte.hu, Web: www.hte.hu