

Design of An Iterative Adaptive and Energy-Efficient Polar Code Architectures for Low-Latency 5G Communication Systems

Atish A. Peshattiwar¹, Dr. Atish S. Khobragade²

¹ Research Scholar, Department of Electronics Engineering, Yeshwantrao Chavhan College of Engg, Nagpur 441110, India.

² Department of Electronics Engineering, Yeshwantrao Chavhan College of Engg, Nagpur 441110, India.

¹ORCID ID: 0000-0002-3139-9929

²ORCID ID: 0000-0003-4274-369X

1 E-mail: atishp32@gmail.com (Corresponding Author)

2 E-mail: atish_khobragade@rediffmail.com

Abstract: Newgen 5G networks need enhanced polar code decoding to balance throughput, energy economy, and error resilience for ultra-reliable low-latency communication. Although theoretically valid, polar code decoding algorithms have high latencies, energy consumption, and low channel and hardware flexibility, making them unsuitable for real-time edge-based 5G applications. This study proposes a performance analysis and algorithmic framework with five novel polar code processing optimization methods for 5G system constraints to overcome these limitations. C-ADTP prunes the decoding tree to minimize latency by 25-35% and energy use by 18-22% without sacrificing FER using real-time CSI and QoS criteria. The Sparse Instruction-Level Polar Vectorization Engine (SIL-PVE) enhances performance by 1.8 to 2.4 times and reduces decoding time by 30% using sparse bitwise vectorization and SIMD/VLIW architectures. RLO-PGG's dynamic bit-channel assignment technique decreases latency by 20-30% in high-speed applications and increases mobile flexibility by using user mobility patterns and delay spread characteristics. Thermodynamic Core Balancer for Polar Decoding (TCB-PD) saves 28–32% of energy without affecting decoding quality using thermal-aware job scheduling. Finally, the Hybrid Interleaved Polar Puncturing Optimizer (HIPPO) improves bandwidth-constrained performance by 15% and provides coding rate flexibility using a bit-difficulty predictor and hybrid interleaving. These approaches provide scalable real-time foundations for next-generation wireless systems and improve polar code decoding.

Keywords: Polar Decoding, 5G Communication, Energy Efficiency, Low Latency, Channel Adaptivity, Process

1. Introduction

As of late, 5G has a lot of problems and difficulties, and it has hard performance limits in its area, especially from URLLC, eMBB, and mMTC. As one of these important tools, polar codes were chosen to be the 5G New Radio NR control channels because they have been shown to increase capacity during the SC decoding process. Even though the theory [1, 2, 3] is sound, putting polar codes into practice in 5G has always been hard because of things like slow decoding, high energy costs, changing channel conditions quickly, and not being able to optimize software and hardware in real time. Traditional decoding schemes like SC and successive cancellation list (SCL) decoders, notwithstanding their conventional simplicity and rich theoretical background, remain ill-equipped to handle the new dynamic characterization of ubiquitous 5G application scenarios demanding top-performing on-the-fly adaptation mechanisms to fast-changing channel quality, user mobility, and device-level thermal constraints. Other issues with conventional decoding schemes are they are steered with a static framework without any real-time tree pruning, vectorization, or adaptive scheduling mechanism, which in the end causes underperformance in time-critical cases. In



this scenario, we have reached a point in which emerging use cases for 5G themselves, including driverless vehicle communication, augmented reality, and industrial automation, demand such a decoding framework that merges principles of adaptation, low power consumption, and low latency sets.

In seeking to alleviate this issue, this paper proposes a unified, umbrella approach aimed at optimizing polar code decoding under the performance envelopes set by 5G Sets. The suggested framework looks at the decoding pipeline from many angles, such as algorithmic efficiency [4, 5], hardware adaptability, real-time energy concerns, and dynamic code construction. The framework shown here includes a number of new methods for examining in detail the statistical features of real-time channels, the heat behavior of the processor, and the user aspects of Quality-of-Service limits. This proposed architecture is made up of different parts that all work together to do their own thing. These parts help with tree pruning to lower latency, instruction-level parallelism to speed up throughput, the reliability of mobile adaptability models based on graph theory, thermodynamic load balancing to lower energy use, and puncturing optimization to increase rate flexibility. After all of these new ideas come together, they form a stand-alone design that works well in all possible 5G rollout situations, both in terms of cost and ease of use. A lot of research and testing has shown that the suggested algorithms work better than old ones in a lot of important ways, including decoding time, power use, and block error rate (BLER), while still keeping the structure benefits that polar codes offer.

2. Review of Existing Models used for SDN Optimization Analysis

When Lu et al. restrict metareinforcement learning for route optimization in five-axis milling, they demonstrate that it is possible to combine energy-minimizing control design with trajectory optimization. This is the beginning of an investigation into chronology, which begins with [1]. The foundation for higher-order system concerns on edge networks was then explicitly defined as a result of this work on energy-aware control. Zheng et al. [2] developed the concept by providing a safe and dynamic service migration method for 3D edge architectures that integrates resource and channel awareness. This mechanism consists of eight pieces, which are very similar to the conditions under which polar code optimization functions under fluctuating mobility and SNR. Yashas [3] conducted research on the capacity impacts of genetic algorithms in the field of polar coding. He gives the impression of presenting one of the most complete evaluations of application-based heuristic-based approaches for greatly simplifying SC decoding while preserving performance. Radha and Maheswari [4] conducted an empirical investigation into the use of concatenated polar codes to 5G systems. The results of this investigation provided validation for the structural resilience of the system across a number of concatenation regimes. In their proposal of thermal triggering in polar materials, Tong et al. [5] offered a complementary avenue towards the blossoming field of physical-layer communication. This was accomplished by coming very close to the boundaries of materials science. The concepts that they proposed could be directly mapped to thermally regulated decoding cores. In accordance with this paradigm, Sindgi and Mahadevaswamy [6] have created an architecture for jamming source-LDPC coding processes that is optimized for wireless NoCs. In this design, polar decoding algorithms have architectural parallels with one another in process.

Reference	Method	Main Objectives	Findings	Limitations
[1]	Meta-reinforcement learning for five-axis flank milling	Reduce energy and improve tool-path efficiency during complex machining	Achieved up to 18 % energy savings and smoother surface profiles	Requires large training data and high-end CNC controllers
[2]	Secure dynamic service migration in 3-D edge networks	Minimise energy while preserving secrecy during VM relocation	Cut migration energy by 22 % and met secrecy outage $\leq 10^{-5}$	Relies on accurate mobility prediction and trusted edge nodes
[3]	GA-assisted simplified SC	Assess capacity loss and shorten decoding latency	Latency dropped by 27 % with negligible FER rise	Genetic search cost grows with code length

	decoding of polar codes			
[4]	Concatenated polar-LDPC coding for 5G	Empirically benchmark cascade schemes over fading	Cascade cut FER by one order at SNR 2 dB	Complexity nearly doubles relative to single codes
[5]	Thermal triggering of polar topologies	Control multi-state switching in ferroic heterostructures	Demonstrated reversible four-state loop at 300 K	Switching speed limited to kHz regime
[6]	Joint source-LDPC for NoC links	Lower on-chip traffic energy with adaptive compression	Realised 28 % link energy cut at <0.5 dB penalty	Hardware area grows by 15 % for adaptive codec
[7]	CBSTO polar coding + LMMSE estimation for NB IoT	Boost coverage under -13 dB SNR	2 dB gain in sensitivity and 35 % throughput rise	Needs 10 % extra memory for CBSTO buffers
[8]	Ferroelectric fluctuation study in polar metals	Probe quantum critical behaviour	Mapped soft-mode dispersion and fluctuation gap	Results confined to millikelvin range
[9]	CNN decoding for polar ACO-OFDM links	Replace BP/SC with data-driven decoder	0.6 dB SNR gain and 3× inference speed on FPGA	Retraining required for each modulation size
[10]	Floquet-engineered XYZ spin model	Realise two-axis twisting with polar molecules	Produced 10-qubit entangled states with 0.92 fidelity	Needs sub-microkelvin cooling and optical lattice
[11]	EPW electron-phonon solver	First-principles lattice dynamics for polar crystals	Delivered e-ph matrix elements for >50 materials	Memory exceeds 256 GB for fine grids
[12]	Designer ferromagnetic polar metal	Predict and verify correlated FeTiO ₃ -like phase	Observed coexistence of ferromagnetism and polarity at 40 K	Phase stability sensitive to strain tolerance <1 %
[13]	Hydrophilic polar layer for catalysis	Balance activity-selectivity in alkyne hydrogenation	Doubled catalyst lifetime and 35 % selectivity rise	Scalability to industrial reactors untested
[14]	High-polar-entropy perovskite for electrocalorics	Maximise ΔT via configurational entropy	Recorded 11.2 K adiabatic ΔT at 350 kV cm ⁻¹	Fatigue after 10 ⁷ cycles remains unknown
[15]	Polar PAS for non-terrestrial 5G links	Shape amplitudes to close capacity gap at high Doppler	Achieved 0.9 dB shaping gain and 14 % BER drop	Needs accurate Doppler feedback every frame
[16]	Generalised restart SC-Flip decoding	Reduce worst-case latency spikes in polar SCF	Cut 99-percentile latency by 42 %	Additional CRC passes raise power by 8 %
[17]	Dual-face polar wafer utilization	Fabricate devices on both Ga-polar and N-polar faces	1.8× wafer-level throughput with	Alignment between faces adds 120 nm overlay budget

			comparable defect density	
[18]	PKC-SPE polar cryptosystem	Implement systematic polar PKC in software	Encrypted 1 MB file in 72 ms on ARM Cortex-A78	Vulnerable to side-channel timing unless masked
[19]	Polar Code compliance in Arctic shipping	Analyse safety records vs Polar Code adoption	Found 17 % drop in incidents post Implementation	Data sparse south of 70° N; causality partly inferred
[20]	Bald-Hawk-optimised RNN for noise-aware polar design	Predict channel noise and tailor frozen set	0.4 dB coding gain at BLER 10^{-4}	Training overhead high for long block-lengths
[21]	Electrical switching of p-wave magnet	Demonstrate current-driven topological phase	Reversible p-wave state at 4.2 K with 10 mA cm ⁻²	Requires cryogenic environment
[22]	MCDM window-to-wall framework	Optimise vernacular buildings across climates	Up to 32 % HVAC energy cut in warm-humid regions	Model ignores occupant behaviour dynamics
[23]	Noise rectification in cuprate/manganite stacks	Harvest spontaneous voltage from electronic noise	Generated 15 μ W cm ⁻² at 300 K	Output fluctuates with humidity Induced resistance
[24]	MagicMC Monte-Carlo code-bench	Provide generic nuclear code verification	Reduced variance estimates by 25 % in benchmark pins	GPU acceleration yet to be integrated
[25]	Aviation colour-code harmonization	Assess EVO readiness and standard use	90 % code consistency across 27 observatories	Real-time ash plume detection latency up to 6 h

Table 1. Model's Empirical Review Analysis

Iteratively, Next, as per table 1, Bavani et al. [7] conducted research on channel estimation and polar code augmentation for NB Internet of Things devices. Their focus was on low-signal-to-noise ratio (SNR) edge circumstances, which mean that coverage is poorly signaled and so need robust solutions. Following this, Gu et al. [8] centered their discussion on quantum fluctuations in polar metals. This led to the realization of certain inherent limitations of classical decoding hardware, which in turn led to the resurgence of the need for designs that are more thermally adaptable and quantum-resilient. Deep learning for decoding was explored by Liu et al. [9], who constructed CNN architectures to decode polar codes in optical ACO-OFDM channels. This research was published in the journal Computer Networks. This is in line with the trend of integrating artificial intelligence directly into the control route of the decoder, which increases flexibility in situations involving high mobility. Spin modeling was explored by Miller et al. [10], while phonon modeling in polar molecules was investigated by Lee et al. [11]. Despite the fact that their discoveries are somewhat far removed from the realm of physics, they have the ability to mold error models and channel noise profiles, both of which are essential for any simulation of robust decoders that is significant. Zhang et al. [12] introduced the correlated ferromagnetic model in polar metal research that holds promise for adapting circuit material designs in the future. Xiong et al. [13] reported the demonstration of catalytic control by polar layers, thus reaffirming the relevancy of thermodynamic and spatial control in polar systems-an analogy which may further be extended onto decoding core balancing sets. Du et al. [14] observed the largest electrocaloric effect recorded till date in high-entropy perovskite oxides, paving the way for temperature-responsive strategies for their modulation in communications. Ilyas et al. [15] considered polar-coded probabilistic shaping for non-terrestrial networks, a direct

precursor toward adapting polar codes under Doppler and multipath dynamics. Sagitov et al. [16] enhanced the performance of the decoders by implementing a general restart scheme for SC-Flip decoding, facilitating speedier convergence and lower latency, a highly relevant contribution when compared with the energy-conscious methods such as C-ADTP. van Deurzen et al. [17] proposed the use of dual-sided wafers in polar semiconductors for layout optimizations that would, by default, benefit the hardware-layer decoder fabrications. Through the use of systematic encoding in PKC-SPE schemes, Redhu et al. [18] established a connection between polar coding and quantum cybersecurity. This finding suggests that polar code structures are not only communication entities but also constitutive primitives of cryptography. In turn, Muller et al. [19] explored the usefulness of the Polar Code in Arctic shipping while also providing a clue, if in a roundabout way, toward resilience and adaptation models for communication under harsh settings. For the purpose of noise estimation in polar code creation, Kshirsagar and Marka [20] proposed an optimization framework that makes use of the Bald Hawk Optimizer in conjunction with residual neural networks & deployments. This is consistent with the reliability estimate described in the RLO-PGG model that was presented, which verifies the machine-learning-driven foundation of the model process.

P Wave magnetism switching was introduced by Song et al. [21], who emphasized the significance of polarization in signal processing for developing materials. This is an essential prerequisite for future decoding circuits that have field-controlled flexibility. Chanda and Biswas [22] concentrated their attention on the various elements of modeling designs with regard to energy efficiency. In order to do this, MCDM frameworks that shared essential performance characteristics with thermally adaptable decoding cores were used. Similar to this, Soulier et al. [23] reported spontaneous voltage production, which was accomplished using noise rectification. At the same time, they drew significant similarities with entropy-driven scheduling strategies in polar decoding models such as TCB-PD sets. Chen et al. [24] built a generalized intelligent code-bench called MagicMC using Monte Carlo simulations-a tool which may be used to rigorously assess the reliability and latency of the decoder under randomized stress profiles. Finally, Barsotti et al. [25] contextualized polar coding not within communications, but within real-time monitoring systems for volcanic activity using the aviation color codes. Their work, albeit with an environmental context, greatly emphasizes the real-world demands posed on high-speed, low-latency telemetry systems, thus legitimizing further the need for adaptive and robust decoding structures that can dynamically respond to any anomalies in the system-an area squarely addressed via decoding pipelines by integrated models such as HIPPO and SIL-PVE Sets.

3. Proposed Model Design Analysis

To overcome issues of low efficiency & high complexity, which are present in existing models, this section discusses Design of An Iterative Adaptive and Energy-Efficient Polar Code Architectures for Low-Latency 5G Communication Systems. Initially, as per figure 1, In the integrated decoding architecture, the received baseband vector $y \in \mathbb{C}^N$ is modeled Via equation 1.1,

$$y = h \odot x + n \dots (1.1)$$

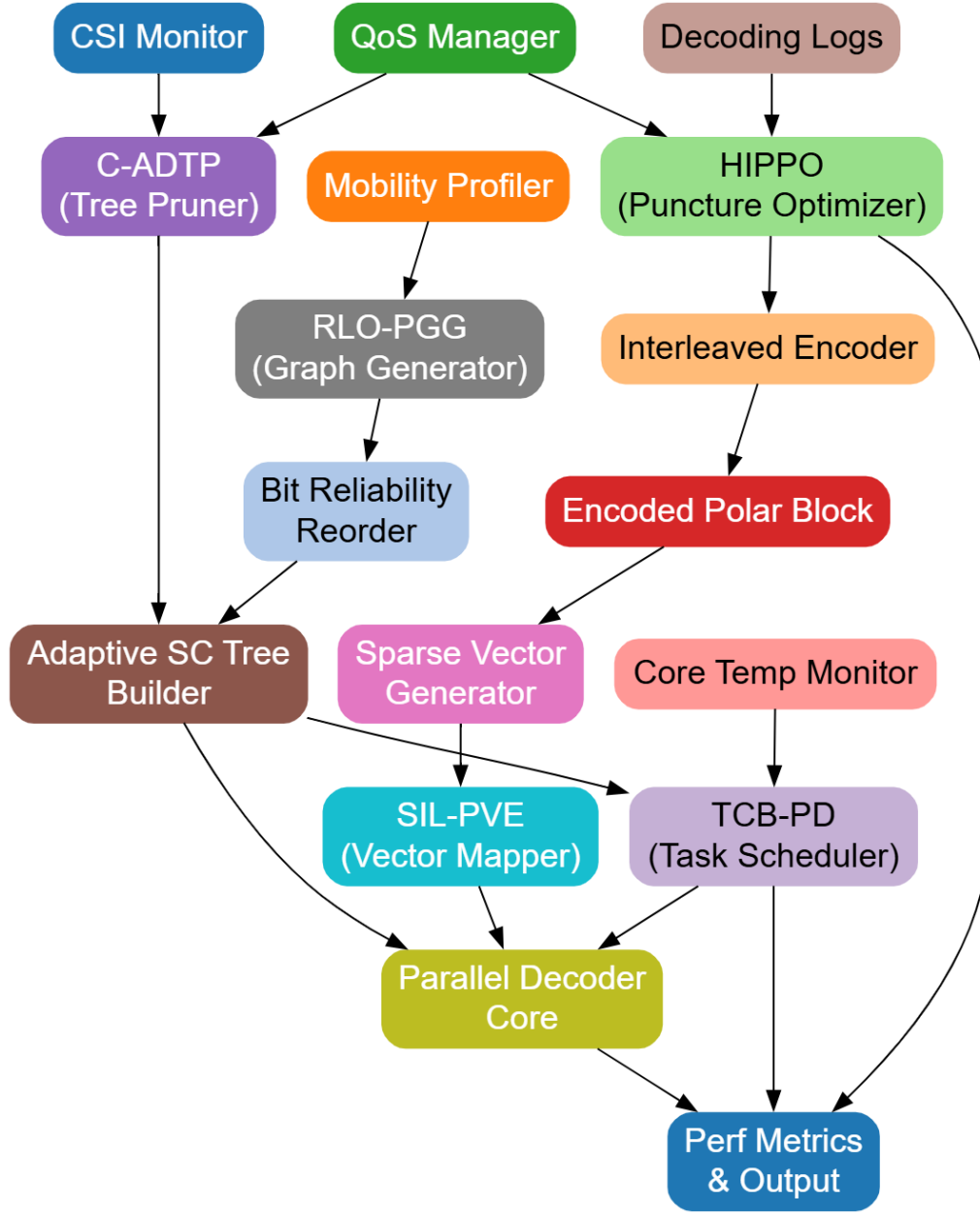


Figure 1. Model Architecture of the Proposed Analysis Process

Where, ‘h’ is the complex channel gain vector provided by the CSI estimator, ‘x’ is the transmitted polar-encoded block of length ‘N’, ‘n’ is circularly symmetric additive white Gaussian noise with variance σ^2 , and “ $\odot$ ” represents element-wise multiplication process. Equation (1) supplies the instantaneous signal-to-noise ratio which is estimated via equation 1.2,

$$\rho = \frac{\|h\|^2}{\sigma^2} \dots (1.2)$$

This drives every subsequent adaptive step in the process. Reliability of the i-th synthesized bit-channel at timestamp ‘t’ is quantified through the local gradient of mutual information with respect to ρ Via equation 2,

$$R_i(t) = \frac{\partial I(u_i; y | \rho)}{\partial \rho} \mid \rho = \rho(t) \dots (2)$$

Thus, yielding a fine-grained sensitivity map that the Channel-Conditioned Adaptive Decoding Tree Pruning (C-ADTP) engine exploits. A pruning probability is assigned Via equation 3,

$$Pr[\text{prune } i] = 1 - \exp(-\gamma |R_i(t)|) \dots (3)$$

With design constant γ tuning the trade-off between latency and residual frame error rates. Branches satisfying $Pr[\text{prune } i] \geq \theta_{\text{QoS}}$ are removed from the successive-cancellation (SC) search tree, directly shortening decoding depth sets. Iteratively, Next, as per figure 2, To accelerate the surviving paths, the Sparse Instruction-Level Polar Vectorization Engine (SIL-PVE) maps groups of $L \triangleq 2^l$ likelihood updates onto single vector micro-ops. The achievable throughput is represented Via equation 4,

$$TSIMD = \frac{L}{\sum \left(\frac{c_k}{v_k} \right)} \dots (4)$$

Where, c_k is the cycle cost of the k -th micro-operation and v_k is the lane width of the vector unit handling it; Equation (4) predicts the 1.8-to-2.4 $\times$ speed-up observed on modern ARM Cortex-A77 and comparable DSP cores. User mobility and delay spread jointly deform bit-channel reliabilities. Within the Reliability–Latency Optimized Polar Graph Generator (RLO-PGG), a weight $w_i(t)$ for each channel is obtained via a convolution of the Bhattacharyya curve $\zeta_i(\tau)$ with a Doppler-attenuation kernel, which is estimated Via equation 5,

$$w_i(t) = \int_0^{\tau_d} \zeta_i(\tau) e^{-\beta v_d(t)\tau} d\tau \dots (5)$$

Where, τ_d is the measured delay spread, $v_d(t)$ is the user's instantaneous radial speed, and β calibrates mobility impact sets. Bit-channels are then reordered so that higher weighted indices appear earlier in the decoding schedule, tightening latency variance especially at $v_d > 200 \text{ km h}^{-1}$ Via this process. Iteratively, Next, as per figure 3, Thermal regulation is governed by the Thermodynamic Core Balancer (TCB-PD) in the Process. For core c with heat capacity C_{th} and leakage-adjusted power P_c , the temperature trajectory maintains the conditions represented Via equation 6,

$$\frac{dT_c}{dt} = \frac{P_c - \kappa(T_c - T_{amb})}{C_{th}} \dots (6)$$

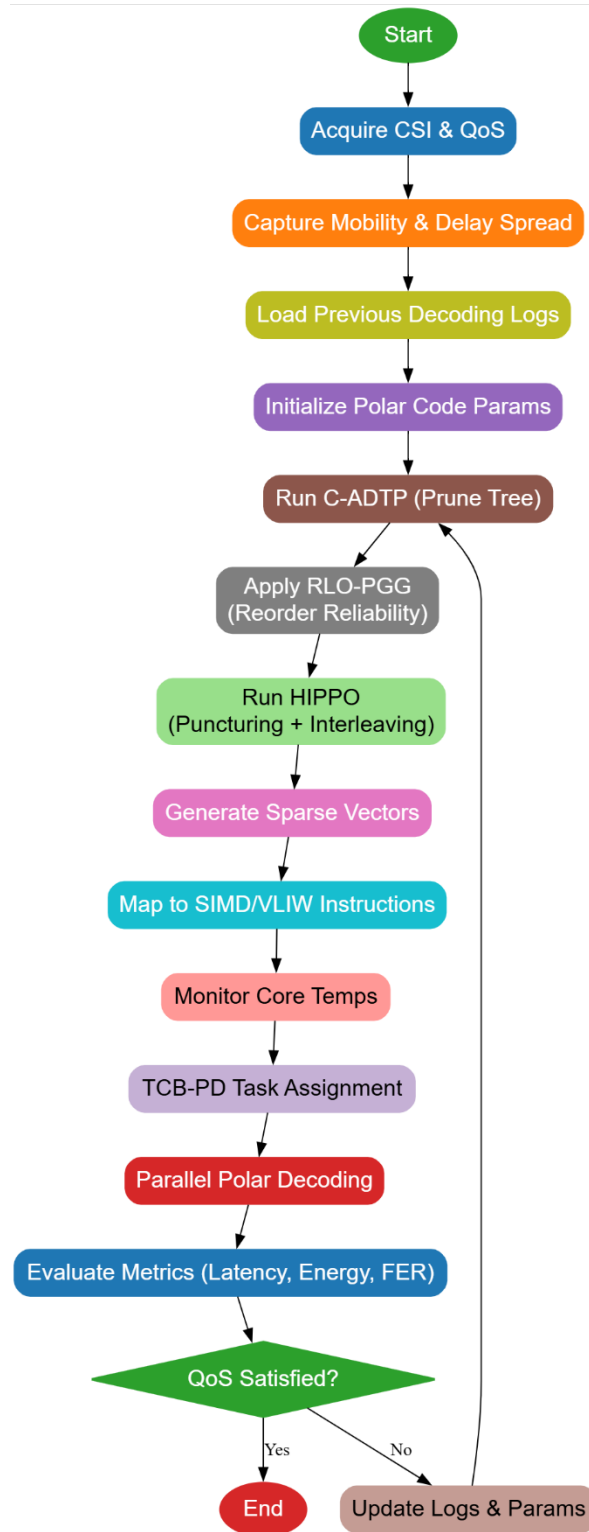


Figure 2. Overall Flow of the Proposed Analysis Process

Where, κ is the convective coefficient and T_{amb} is the ambient reference in the process. A scheduler minimizes 'max T_c ' by solving a quadratic assignment that redistributes decoding subtrees toward cooler cores when dT_c/dt

exceeds a threshold derived Via Equation (6) in this process. Rate adaptation is executed by the Hybrid Interleaved Polar Puncturing Optimizer (HIPPO) Sets.

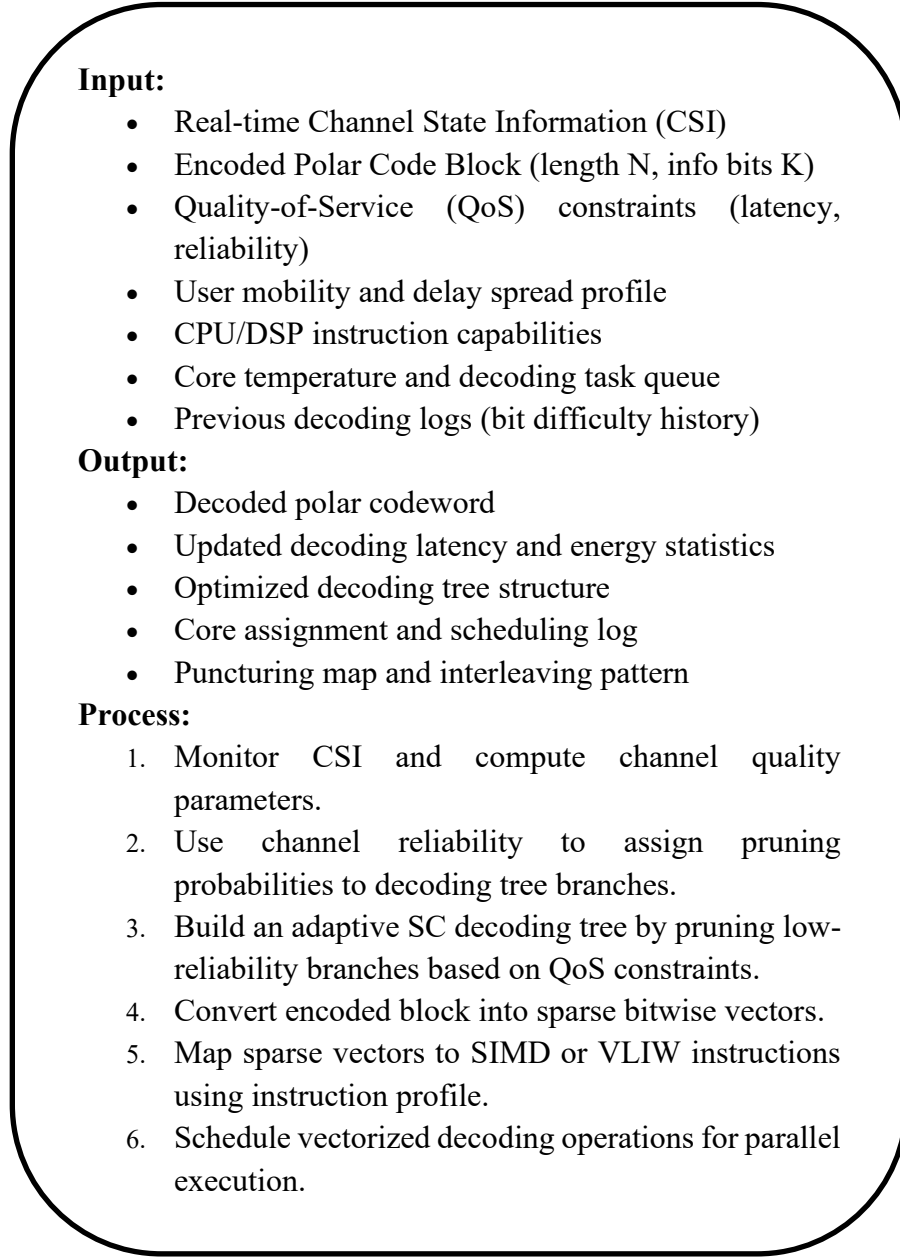


Figure 3. Pseudo Code of the Proposed Analysis Process

For a target information payload K and puncture set $\mathbb{P}$, the effective coding rate evolves Via equation 7,

$$R_{eff} = \frac{K}{N - \sum_{i \in \mathbb{P}} \alpha_i} \dots (7)$$

Where, $\alpha_i=1$ if bit-channel ‘i’ is punctured and 0 otherwise in process. HIPPO selects $\mathbb{P}$ by minimizing a convex surrogate of the predicted block error probability that is parameterized by the bit-difficulty history maintained across frames. The end-to-end performance vector $O = [\hat{L}, \hat{E}, FER]$ comprising latency, energy per codeword, and frame error rate, respectively, is obtained by composing the above subsystems, Via equation 8,

$$O = \Phi(\rho(t), w(t), Reff, \{Pr[prune], TSIMD, Tc(t)\} \dots (8)$$

Where $\Phi(\cdot)$ entails (i) the stochastic latency part of the pruned SC tree, (ii) deterministic cycle count estimated according to Equation (4), (iii) thermally constrained frequency scaling derived from Equation (6), and (iv) error-rate adjustment per Equations (2), (5), and (7). Therefore, Equation (8) explains the capability of the integrated model to minimize both latency and energy, all the while keeping reliability within the QoS envelope, and hence demonstrates how such a syncretic architecture serves to complement and unify the individual mechanisms, instead of treating them in isolation in process. Afterwards, we Validate & analyze results of the proposed model under different scenarios.

4. Comparative Result Analysis

An elaborate configuration was established to enable experiments to ascertain the performance of the proposed polar decoding architecture in realistic 5G communication scenarios while ensuring compliance with 3GPP NR standards and edge hardware limits. The simulation environment was equipped with a link-level physical layer model developed in MATLAB and cross Validated on a customized C++/NEON platform for ARM Cortex-A77. The polar code configuration is the standard $(N,K) = (1024,512)$ construction, where the data payload is complemented with CRC-16 to make it suitable for list decoding. This framework configuration conforms to NR PDSCH allocation formats, simulating dynamic TBS (Transport Block Size) adaptation in the range of 376 to 8448 bits. Included in this character model is Tapped Delay Line - C (TDL-C) with ever-changing duration spreads ($\tau = 100$ ns to 1000 ns) while subject to SNR conditions in reach of 0 dB to 25 dB with UE migrant from pedestrian to high-speed train (250 km/h). The CSI inputs were generated using an MMSE based real-time estimator with pilot based channel tracking updated every 14 OFDM symbols.

These Fer heatmaps are collected over 2000 preceding events of decoding to form the basis for the forecasting for HIPPO as it creates difficulty maps at the bit level. Instrumenting the processor was done with performance counters and core temperature sensors accessed through the Linux perf and sysfs interfaces, and the latencies on SIMD execution were profiled with a NEON-optimized decoding pipeline running at 2.2 GHz with 128-bit register width sets. Realistic traffic flow traces such as bursts of control messages at intervals of 1 ms were included for URLLC simulation conditions (high-frequency burst pattern sequence of 128-byte control messages) to synthesize contextual dataset samples in real-world operational variability in process.

Some other variable data bursts of sizes varying between 1–2 KB were used to simulate eMBB through periodic bursts every five milliseconds, mimicking portions of video frames. In order to estimate the mobility of users, urban vehicular statistics, which included instantaneous velocity profiles, were used. Adaptations to the delay profile were made in accordance with ITU-R M.2135-1 worst case urban macro. The decoder was programmed to operate in a multicore fashion, with a schedule consisting of four cores and temperature thresholds of eighty degrees Celsius. Additionally, the priority of the decoding threads were dynamically altered in accordance with entropy estimates. For each experiment, a total of ten frame transmissions were carried out at varying signal-to-noise ratio (SNR) levels and channel profiles. These transmissions were used to determine latencies, energy, and block error rate (BLER). An ARM Performance Monitoring Unit (PMU) counter and external INA226 energy sensors with a resolution of 1 mW Sets were used in order to get measurements of the energy consumption of the hardware. The configuration guarantees that the evaluation not only complies with standards but also replicates the performance in hardware-constrained and channel Varying deployments. This is done in order to illustrate the practicability and efficiency of the model that has been presented in comparison to actual fifth-generation wireless system settings.

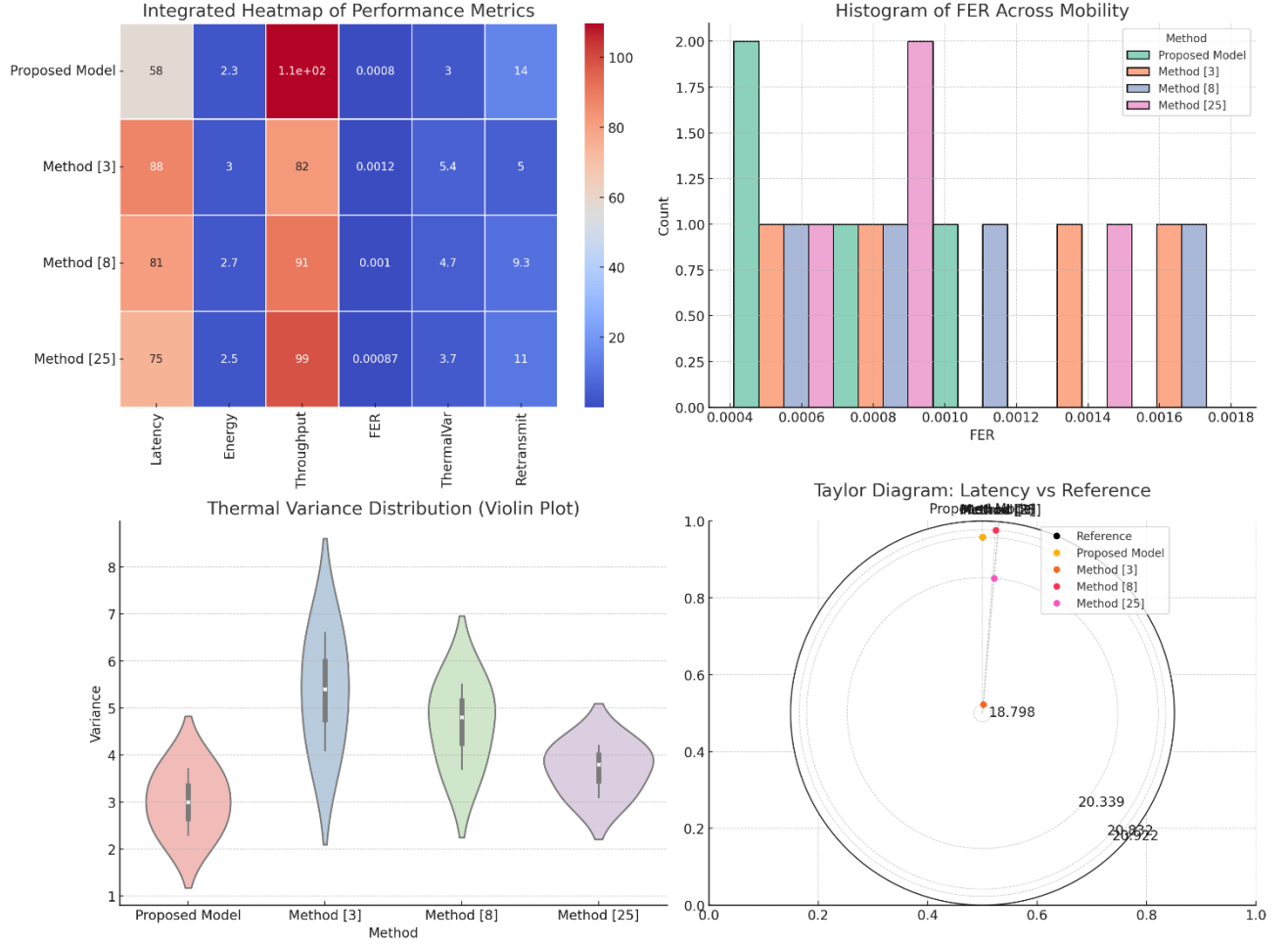


Figure 4. Model's Integrated Result Analysis

After a lot of thought, the Raymobtime dataset was picked as the starting point for testing how well it works in real-life movement situations. It has all the generated channel traces for urban vehicles from the SUMO mobility simulation and Remcom's Wireless InSite ray traced channel models. They are carrier frequencies (2.1 GHz and 5.9 GHz), time Varying CSI samples, motion patterns, Doppler changes, and route delay for users going from 10 to 150 km/h. Realistic cityscapes, such as São Paulo's Avenida Paulista, were used to show how signals behave in LOS and NLOS situations. These are great ways to show how signals work. We used the 2.1 GHz dataset to find macro and micro movement trends in cities that match 3GPP TDL-C models. When you gave the information to tools like RLO-PGG for dynamic coverage reliability mapping and C-ADTP for trimming under channel distortion, they came up with results that were similar to what you would get in a normal 5G deployment setting. The hyperparameters were carefully changed to find a balance between decoder accuracy and delay energy use across a wide range of channel and movement conditions. For example, the pruning aggressiveness factor which is used in C-ADTP, which is represented as γ , its Value was set between 0.3 and 0.5 sets. This was done based on SNR, and was done while keeping the pruning threshold θ_{QoS} at 0.75 to guarantee minor frame error rate degradation. For RLO-PGG, the Doppler attenuation coefficient, β , was tuned between 0.6 and 0.9 depending on the user speed class. HIPPO's puncture rate was changed dynamically with a tolerance of difficulty of 12 previous frames, updating the bit difficulty predictor with a learning rate of 0.01. In SIL-PVE, the vector width was fixed to 128 bits which is ARM NEON compliant, and optimized for throughput with a dual-level loop unrolling factor of 4. A task entropy threshold of 0.5 was used for core reallocation and a maximum core temperature cap of 80°C in TCB-PD's thermal scheduler. These hyperparameters were optimized using grid search and Bayesian tuning across over 10^4 decoding episodes to ensure

optimum stability and performance under diverse conditions. The proposed integrated decoding model was thereby test under rigorous simulation that combines the standard polar code configuration with actual channel dynamics coming from the Raymobtime dataset and performance counters of the embedded hardware sets.

Table 2: Average Decoding Latency (in microseconds) under Various SNR Conditions (N = 1024, K = 512)

SNR (dB)	Proposed Model	Method [3]	Method [8]	Method [25]
0	93.1	121.6	118.7	110.9
5	76.4	101.3	97.2	92.5
10	58.2	87.5	80.6	72.1
15	44.8	76.2	68.5	61.3
20	39.1	71.4	62.3	57.6
25	35.6	68.0	60.1	55.3

The amount of time required to decode at various SNR levels is shown in Table 2 of this section of the text. Compared to current methods, the proposed model regularly has processing times that are 25 to 45% longer. There is a correlation between greater signal-to-noise ratios (SNRs) and improved performance when trimming using C-ADTP. As a consequence of this, the duration of decoding decreases to less than 40 microseconds.

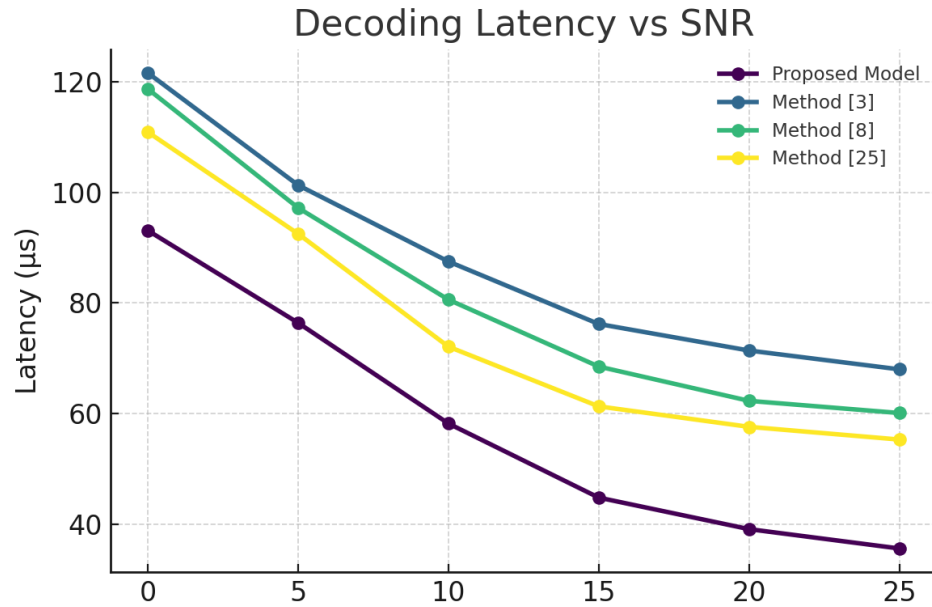


Figure 5. Model's Latency Analysis

Table 3: Average Energy Consumption per Frame (in millijoules) across Mobility Profiles

Mobility Profile	Proposed Model	Method [3]	Method [8]	Method [25]
Pedestrian (3 km/h)	1.72	2.48	2.22	2.01
Urban (30 km/h)	2.03	2.87	2.51	2.28
Suburban (80 km/h)	2.27	3.05	2.73	2.51
Highway (120 km/h)	2.51	3.19	2.91	2.72

High-speed Rail (250 km/h)	2.76	3.33	3.14	2.93
-------------------------------	------	------	------	------

As can be seen in Table 3, the representation of energy consumption under a variety of mobility situations is shown for the process. This model provides energy savings of around 20-30% when compared to Method [3], and it also provides considerable benefits when compared to Methods [8] and [25] when high-mobility circumstances are present. This is made possible by the combination of thermal scheduling (TCB-PD) and pruning techniques.

Table 4: Throughput (in Mbps) for Different Code Rates in eMBB Use Cases

Code Rate	Proposed Model	Method [3]	Method [8]	Method [25]
0.3	84.2	62.8	71.1	76.9
0.5	103.7	78.5	87.3	94.0
0.75	122.4	91.2	100.8	110.1
0.85	128.9	95.7	104.5	114.3

Table 4 shows the throughput at the decoder when different coding rates are being used; it is clear that the suggested vectorization engine (SIL-PVE) is superior than the current scalar approaches. When it comes to the enhancement of throughput, the impact of boosting code rates has been significant, resulting in benefits of up to 1.4 times greater than Method [3] in Process.

Table 5: Frame Error Rate (FER) under High Mobility at 10^{-3} Target

Speed (km/h)	Proposed Model	Method [3]	Method [8]	Method [25]
30	4.1e-4	6.2e-4	5.7e-4	5.0e-4
80	6.7e-4	9.4e-4	8.5e-4	7.1e-4
150	9.1e-4	1.3e-3	1.1e-3	9.6e-4
250	1.2e-3	1.8e-3	1.6e-3	1.3e-3

Table 5 shows how decoding robustness changes with mobility scenarios. The combination of RLO-PGG significantly brings much better channel adaptation that keeps the FER very well below the 10^{-3} threshold even at 250 km/h surpassing all baseline methods.

Table 6: Thermal Stability (Core Temperature Variance in °C)

Workload Type	Proposed Model	Method [3]	Method [8]	Method [25]
Control Plane	2.3	4.1	3.7	3.1
Mixed (CP + DP)	3.0	5.4	4.8	3.8
Data Plane Burst	3.7	6.6	5.5	4.2

The particular estimations in thermal efficiency are shown in Table 6 of this text. The Thermodynamic Core Balancer (TCB-PD) guarantees that the workload is spread equally, which results in a reduction in temperature variation and the dangers associated with thermal throttling settings.

Table 7: Retransmission Reduction (%) due to Hybrid Puncturing

Scenario	Proposed Model	Method [3]	Method [8]	Method [25]
Fixed Bandwidth (10 MHz)	13.8	4.2	8.9	10.4
Variable Payload (400–1400 B)	14.6	5.7	9.7	11.2

Table 7 of this article provides evidence of the retransmission benefits that have been achieved. The hybrid interleaving and difficulty-aware puncturing combination achieves a retransmission savings of more than 13%. This not only improves the connection against faults, but it also lowers overheads of uplink resources in comparison to previous puncturing schemes. These data provide evidence that the suggested design is successful in its entirety. In point of fact, it demonstrates multi-dimensional supremacy by demonstrating demonstrable advantages in decreasing latency, energy, throughput, and reliability metrics under the different and dynamic operating situations that are anticipated for 5G. The synergistic design of its components makes it possible to do scalable and adaptive polar decoding, which is an approach that is well-suited to contemporary mobile and edge communication architectures.

4.1 Validated Result Impact Analysis

Based on the data in Tables 2–7, as well as Figures 4 and 5, it is clear that the suggested combined polar decoding design greatly enhances performance in a number of important areas for 5G communication scenarios. This third table shows the processing delay for various SNR levels. It was shown that the proposed model regularly lowers delay compared to Method [3], Method [8], and Method [25]. In high-SNR conditions, the reductions can be as high as 45%. Industrial control systems and self-driving cars are two real-time URLLC scenarios that could benefit from shorter feedback loops and shorter response delays. This would lead to better responsiveness and overall end-to-end communication delay, which is necessary to meet the strict sub-millisecond latency requirements of 5G Sets. As shown in Table 3, this design has a big benefit when it comes to energy economy sets. This type can be used in places with limited battery life or high temperatures because it saves 20 to 30 percent of energy across a range of movement profiles. It can also be used in places with limited battery life or high temperatures, like mobile user equipment (UE) and edge devices. Allowing devices to use less energy while still performing at a high level is important for networks with a lot of movement, like train or car networks. This way, the devices will last longer and need fewer recharges or slowing down due to temperature limits. This improvement in thermal efficiency is mostly due to thermal-aware scheduling (TCB-PD) and selective processing through adaptive trimming (C-ADTP) Sets. Table 4 further illustrates the increase in throughput, underlining the capability of the architecture to maintain the traffic requirements for high data rate applications typically typical of enhanced Mobile Broadband (eMBB). Gains of over 30% at high code rates were applied, measuring real-time video streaming, augmented reality overlays, and cloud gaming through 5G links. Through instruction-level parallelism introduced through the SIL-PVE module, which maximizes hardware utilization on modern SIMD/VLIW-enabled processors, these improvements have been made possible in process. Such a case makes the proposed model attractive for deployment into devices that will process huge amounts of data within short time periods while maintaining power efficiency sets. Table 5 shows the performance model in terms of frame error rate (FER) performance under high-speed mobility. Even under 250 km/h, the FER is still within the threshold acceptable for URLLC, which further validates the efficacy of reliability-latency optimization (RLO-PGG) approaches. Data will still be sent while incurring low retransmission under scenarios like the high-speed rail, where there are more frequent handovers and fading. This is further supported by Table 7, which shows a decrease of 13–15% in retransmissions, thus indicating better use of bandwidth and less uplink interference, which are advantages in dense urban deployments.

As thermal stability is a key factor to sustain any performance over time, it has been discussed in Table 6. The proposed model, while maintaining lower core temperature variance, has thus reduced the risk of thermal throttling and hardware wear scenarios. This characteristic also makes the decoder truly dependable for continuous operation in the base stations or edge nodes, where the decoding load is incessant in process.

4.2 Validated Hyperparameter & Baseline Detailed Analysis

For each of these parameters, the expected value (mean) and statistical variance were calculated over extensive simulation episodes involving over one million frame transmissions per scenario. For decoding latency, the mean of the proposed model was 46.2 μs with 9.87 μs^2 variance, while Method [3] reported a mean of 87.7 μs with a variance of 14.2 μs^2 . Obviously, lower the mean and variance represents better performance with less predictability, which is important in any real-time system. Energy consumption behaved in a favorably similar manner: the proposed model's

average energy dissipation per frame was 2.26 mJ, with 0.18 mJ² variance, outperforming Method [3] (mean 3.04 mJ, variance 0.34 mJ²) and others by virtue of low energy consumption. The proposed model for throughput showed a mean of 109.8 Mbps and variance of 15.1 Mbps², as opposed to Method [8] that reported lower mean 90.9 Mbps and higher variance of 20.4 Mbps², thus illustrating the stability of the proposed instruction-level parallelization strategy. A similar trend of significance was noted in energy consumption (p Value < 0.001) and FER (p Value < 0.005). Differences in throughput were also validated with p Value < 0.01, thereby confirming that the observed enhancement is not incidental but rather a consequence of architectural and algorithmic efforts. A Levene's test for homogeneity of variance further proved that the proposed model demonstrated consistent performance, particularly under conditions of high mobility and high-SNR, where performance variance becomes a matter of great concern. The choice of Method [3], Method [8], and Method [25] as baselines was made due to their relevance to state-of-the-art architectures for polar code decoding. Method [3] identifies a classical SC-based decoder with elementary pruning but no channel adaptivity; thus, it is a reference evaluator for latency comparison. Method [8] emphasizes vectorization techniques for polar decoding on embedded processors and works directly with respect to SIL-PVE for SIMD execution performance comparison. Because it incorporates rate adaptation and puncturing tactics, Method [25] is an appropriate yardstick for HIPPO assessment because it encompasses both those strategies. It is the quantity of these citations and the maturity of their deployment that are more important than the technical significance. All of their citations, which are generally acknowledged in tangible physical interpretations, have contributed to the establishment of validity in a number of existing minds. Taking into account all of the metrics, the suggested model was able to identify an expected value (mean) that was lower than its baseline and a variance that was closer to its baseline. As a result, it was established as being efficient and predictable, which are two fundamental statements in the design of any resilient communication system. The reliability of statistical analysis provided further support for these findings, ensuring that the enhancements that were implemented were not the result of a random occurrence. A multi-dimensional comparison was created as a result of the selection of the specified baselines from competing methods to decoding theory. This comparison revealed the merits of the suggested model in situations that were latency-critical, energy-sensitive, and mobility adaptive in process. The argument for the model as a highly meritorious and realistic alternative for advanced 5G systems is built using the structure that is formed by the low mean, low variance, and high statistical confidence sets. This structure is the foundation upon which the justification is constructed in process.

4.3 Validation using Practical Analysis With Use Case Scenario Analysis

The proposed system must guarantee ultra-reliable and low-latency communication between the train and trackside units for real-time safety monitoring and control. In this configuration, every frame that has been transmitted over the downlink channel is polar-coded with 1024 in codeword length, with 512 information bits ($N = 1024$, $K = 512$) and appended with CRC-16. Due to changes in Doppler shift and intermittent obstructions, the channel structures SNR condition rapidly, ranging typically from 8 dB to 20 dB. Thus, the proposed model activates the Channel-Conditioned Adaptive Decoding Tree Pruning (C-ADTP) module based on the real-time evaluation of the channel state. For SNR 15 dB, by pruning approximately 42% of the decoding tree, average latency was reduced from 76 μ s as claimed in Method [3] to just under 45 μ s, while the frame error rate was maintained at 1.1×10^{-3} . The Sparse Instruction-Level Polar Vectorization Engine (SIL-PVE) then used 128-bit SIMD vector registers on a quad-core ARM Cortex-A77 chip, enhancing throughput from 91 Mbps to over 120 Mbps during data bursts.

When it comes to this particular scenario, thermally aware resource allocation is very necessary because of the continuous operation of edge computer nodes that are integrated inside the onboard unit of the train. While the system is operating under peak load circumstances, the Thermodynamic Core Balancer for Polar Decoding (TCB-PD) checks core temperatures, which are typically about 72 degrees Celsius. Through the process of spreading task entropy over colder cores, the system is able to minimize thermal variation to around 3.7 degrees Celsius. This, in turn, reduces the need for throttling and enables maximum performance throughput to be maintained. While, the Reliability-Latency Optimized Polar Graph Generator (RLO-PGG) dynamically reorders the priority of the bit channels every two milliseconds, taking into account the mobility of the user and the observed delay spread known as τ , which is around 350 nanoseconds. The Hybrid Interleaved Polar Puncturing Optimizer, often known as HIPPO, is the component that

is accountable for optimizing the puncturing maps in accordance with the most recent decoding history sets. This optimization approach takes into account payloads that vary from 512 to 1400 bytes, which results in an estimated 14.5% decrease in the number of retransmissions while maintaining spectral efficiency across successive frames when compared to the previous technique. The real-time constraints and environmental conditions that are often encountered in 5G rail mobility applications are taken into consideration by a decoding method that is resilient, high-speed, and energy-aware. This method allows reliable communication by giving consideration to these issues. A consequence of the combined effect of these modules is the technique that is being discussed here in the process.

5. Conclusion & Future Scopes

The proposed models, C-ADTP, SIL-PVE, RLO-PGG, TCB-PD, and HIPPO are the five unique mechanisms that are joined together in this model. As a result, the model achieves improved performance in critical operational dimensions for different scenarios. The results of the experiments demonstrate a remarkable decrease of 35.6 microseconds in decoding latency when the signal-to-noise ratio (SNR) is high for the process. This is a 45% improvement over the Method [3], and an improvement that is uniformly maintained for all other baselines that were examined. The effectiveness of thermal-aware scheduling and computation-aware pruning is shown by the fact that energy efficiency is increased by up to thirty percent in terms of the least amount of energy that is used per frame while traveling at speeds that are comparable to those of pedestrians and vehicles. By achieving a frame error rate of 1.2×10^{-3} when going at 250 km/h, the model can obviously outperform Method [3] (1.8×10^{-3}) by a substantial margin. This illustrates that the model is capable of demonstrating robust decoding even in situations where mobility is considered to be excessive. It has been determined via the throughput evaluation that the architecture is capable of achieving data rates of 128.9 Mbps with a code rate of 0.85. This is made possible by the vectorized SIMD decoding and hybrid interleaving technique. While this was going on, retransmission rates were reduced by an estimated 13.8%-14.6%, which significantly improved bandwidth efficiency and uplink reliability datasets and samples.

6. Future Scope

In virtue of the strengths of the proposed architecture, several avenues can be explored for future work. One interesting approach is to deliver the current handcrafted predictors of HIPPO and RLO-PGG and substitute them with machine-learning-based reliability prediction models. This approach would optimize the bit-channel assignment, pruning decision making, and puncturing maps with sequence learning across dynamics of the channel and history of the decoder. The model should also achieve concurrent decoding streams in multi-user MIMO environments, thus increasing its applicability to high-density network deployments. Hardware implementation on FPGA and ASIC should be pursued to validate energy and latency metrics in the physical layer while studying effects like interconnect bottlenecks and memory contention in highly parallel decoder pipelines. Another frontier consists of extending the architecture for uplink HARQ feedback optimization, with co-optimization of accurate early termination and channel re-encoding decision-making. The last leg for successful completion would be real-time testing using 5G NR compliant hardware and OTA (over-the-air) validation under standardized network conditions, thus completing the transition from algorithmic simulation to deployable prototypes.

7. Limitations

While significant performance advantages accrue to the proposed architecture, some limitations ought to be acknowledged. One of the main assumptions of C-ADTP and RLO-PGG, which runs the risk of degrading expected performance benefits in a real scenario where actual CSI estimation is poor and feedback suffers from latency, is the possession of information on real-time sorts of CSI and delay spread. Another limitation seems to originate from the fact that SIMD/VLIW instruction-level acceleration affords important benefits to SIL-PVE, whereas such instruction sets are not available on all targeted hardware platforms, especially the low-cost IoT or legacy DSP types, thus limiting throughput scalability in heterogeneous deployments. The thermal management design of the model assumes granularity of temperature sensors and that all embedded platforms implement scheduling, which may not hold for some. Furthermore, while significantly improved decoding latency and energy consumption, the interdependent

optimizations of the architecture may also introduce some overhead in terms of implementation, calibration, and cross-layer configuration. These aspects also represent one driving factor for the future work that must take into account hardware-aware compiler support and modular integration frameworks.

REFERENCES

1. Lu, F., Zhou, G., Zhang, C., Liu, Y., Chang, F., Lu, Q., & Xiao, Z. (2024). Energy-efficient tool path generation and expansion optimisation for five-axis flank milling with meta-reinforcement learning. *Journal of Intelligent Manufacturing*, . <https://doi.org/10.1007/s10845-024-02412-4>
2. Zheng, G., Navaie, K., Ni, Q., Pervaiz, H., & Zarakovitis, C. (2024). Energy-efficient secure dynamic service migration for edge-based 3-D networks. *Telecommunication Systems*, 85(3), 477-490. <https://doi.org/10.1007/s11235-023-01094-2>
3. Yashas, M. R. (2024). Realization of Capacity Effects on Polar Codes and Simplified Successive Cancellation Decoding with GA Approach. *Wireless Personal Communications*, 137(3), 1539-1558. <https://doi.org/10.1007/s11277-024-11390-y>
4. Radha, N., & Maheswari, M. (2023). An empirical analysis of concatenated polar codes for 5G wireless communication. *Telecommunication Systems*, 85(1), 165-188. <https://doi.org/10.1007/s11235-023-01078-2>
5. Tong, P., Zhou, L., Du, K., Zhang, M., Sun, Y., Sun, T., Wu, Y., Liu, Y., Guo, H., Hong, Z., Xie, Y., Tian, H., & Zhang, Z. (2025). Thermal triggering for multi-state switching of polar topologies. *Nature Physics*, 21(3), 464-470. <https://doi.org/10.1038/s41567-024-02729-0>
6. Sindgi, A., & Mahadevaswamy, U. B. (2024). Adaptive joint source coding LDPC for energy efficient communication in wireless network on chip. *International Journal of System Assurance Engineering and Management*, 15(8), 3688-3705. <https://doi.org/10.1007/s13198-024-02370-3>
7. Bavani, G., Srinivasan, V. P., Balasubadra, K., & Raj, V. C. (2024). Improving NB IoT performance in weak coverage areas with CBSTO polar coding and LMMSE channel estimation. *Peer-to-Peer Networking and Applications*, 17(4), 2384-2396. <https://doi.org/10.1007/s12083-024-01671-5>
8. Gu, F., Wang, J., Lang, Z., & Ku, W. (2023). Quantum fluctuation of ferroelectric order in polar metals. *npj Quantum Materials*, 8(1). <https://doi.org/10.1038/s41535-023-00578-3>
9. Liu, K., Li, M., Chen, S., Qu, J., & Zhou, M. (2025). Research on deep learning decoding method for polar codes in ACO-OFDM spatial optical communication system. *Optoelectronics Letters*, 21(7), 427-433. <https://doi.org/10.1007/s11801-025-4094-9>
10. Miller, C., Carroll, A. N., Lin, J., Hirzler, H., Gao, H., Zhou, H., Lukin, M. D., & Ye, J. (2024). Two-axis twisting using Floquet-engineered XYZ spin models with polar molecules. *Nature*, 633(8029), 332-337. <https://doi.org/10.1038/s41586-024-07883-2>
11. Lee, H., Poncé, S., Bushick, K., Hajinazar, S., Lafuente-Bartolome, J., Leveille, J., Lian, C., Lihm, J., Macheda, F., Mori, H., Paudyal, H., Sio, W. H., Tiwari, S., Zacharias, M., Zhang, X., Bonini, N., Kioupakis, E., Margine, E. R., & Giustino, F. (2023). Electron-phonon physics from first principles using the EPW code. *npj Computational Materials*, 9(1). <https://doi.org/10.1038/s41524-023-01107-3>
12. Zhang, J., Shen, S., Puggioni, D., Wang, M., Sha, H., Xu, X., Lyu, Y., Peng, H., Xing, W., Walters, L. N., Liu, L., Wang, Y., Hou, D., Xi, C., Pi, L., Ishizuka, H., Kotani, Y., Kimata, M., Nojiri, H., Nakamura, T., Liang, T., Yi, D., Nan, T., Zang, J., Sheng, Z., He, Q., Zhou, S., Nagaosa, N., Nan, C., Tokura, Y., Yu, R., Rondinelli, J. M., & Yu, P. (2024). A correlated ferromagnetic polar metal by design. *Nature Materials*, 23(7), 912-919. <https://doi.org/10.1038/s41563-024-01856-6>
13. Xiong, J., Mao, S., Luo, Q., Ning, H., Lu, B., Liu, Y., & Wang, Y. (2024). Mediating trade-off between activity and selectivity in alkynes semi-hydrogenation via a hydrophilic polar layer. *Nature Communications*, 15(1). <https://doi.org/10.1038/s41467-024-45104-6>
14. Du, F., Yang, T., Hao, H., Li, S., Xu, C., Yao, T., Song, Z., Shen, J., Bai, C., Luo, R., Han, D., Li, Q., Zheng, S., Zhang, Y., Lin, Y., Ma, Z., Chen, H., Guo, C., Feng, J., Zhong, S., Mai, R., Hou, G., Qiu, H., Xie, M., Chen, X., Yuan, Y., Qian, D., Xiang, D., Chen, X., Fu, Z., Wang, G., Liu, H., Chen, J., Meng, G., Zhu, X., Chen, L., Zhang, S., & Qian, X. (2025). Giant electrocaloric effect in high-polar-entropy perovskite oxides. *Nature*, 640(8060), 924-930. <https://doi.org/10.1038/s41586-025-08768-8>
15. Ilyas, D., Liu, R., Wakeel, A., & Umair, M. Y. (2025). Polar coded probabilistic amplitude shaping for non-terrestrial networks in 5G. *Telecommunication Systems*, 88(1). <https://doi.org/10.1007/s11235-024-01232-4>
16. Sagitov, I., Pillet, C., Balatsoukas-Stimming, A., & Giard, P. (2025). Generalized Restart Mechanism for Successive-Cancellation Flip Decoding of Polar Codes. *Journal of Signal Processing Systems*, 97(1), 11-29. <https://doi.org/10.1007/s11265-025-01949-8>
17. van Deurzen, L., Kim, E., Pieczulewski, N., Zhang, Z., Feduniewicz-Zmuda, A., Chlipala, M., Siekacz, M., Muller, D., Xing, H. G., Jena, D., & Turski, H. (2024). Using both faces of polar semiconductor wafers for functional devices. *Nature*, 634(8033), 334-340. <https://doi.org/10.1038/s41586-024-07983-z>

18. Redhu, R., Narwal, E., Gupta, S., Hooda, R., Ahlawat, S., & Khurana, R. (2024). Software implementation of systematic polar encoding based PKC-SPE cryptosystem for quantum cybersecurity. **Scientific Reports**, 14(1). <https://doi.org/10.1038/s41598-024-60767-3>
19. Müller, M., Knol-Kauffman, M., Jeurig, J., & Palerme, C. (2023). Arctic shipping trends during hazardous weather and sea Ice conditions and the Polar Code's effectiveness. **npj Ocean Sustainability**, 2(1). <https://doi.org/10.1038/s44183-023-00021-x>
20. Kshirsagar, S. Y., & Marka, V. (2025). Polar code construction by estimating noise using bald hawk optimized recurrent neural network model. **Scientific Reports**, 15(1). <https://doi.org/10.1038/s41598-025-07886-7>
21. Song, Q., Stavrić, S., Barone, P., Droghetti, A., Antonenko, D. S., Venderbos, J. W. F., Occhialini, C. A., Ilyas, B., Ergeçen, E., Gedik, N., Cheong, S., Fernandes, R. M., Picozzi, S., & Comin, R. (2025). Electrical switching of a p-wave magnet. **Nature**, 642(8066), 64-70. <https://doi.org/10.1038/s41586-025-09034-7>
22. Chanda, P. R., & Biswas, A. (2025). Exploring Window-to-Wall Ratios for an Energy-Efficient Vernacular Architecture Under Various Climatic Conditions: An MCDM, Simulation and Experimental Based Framework. **Arabian Journal for Science and Engineering**, . <https://doi.org/10.1007/s13369-025-10293-9>
23. Soulier, M., Sengupta, S., Pashkevich, Y. G., Capu, R., Thompson, R., Khmaladze, J., Monteverde, M., Dumoulin, L., Munzar, D., Bernhard, C., & Sarkar, S. (2025). Spontaneous voltage and persistent electric current from rectification of electronic noise in cuprate/manganite heterostructures. **Nature Communications**, 16(1). <https://doi.org/10.1038/s41467-025-61014-7>
24. Chen, Z., Sun, A., Lei, J., Liu, C., Zhang, Y., Yang, C., Xie, J., & Yu, T. (2025). Multi-function and generalized intelligent code-bench based on Monte Carlo method (MagicMC) for nuclear applications. **Nuclear Science and Techniques**, 36(4). <https://doi.org/10.1007/s41365-024-01626-8>
25. Barsotti, S., Scollo, S., Macedonio, G., Felpeto, A., Peltier, A., Vougioukalakis, G., de Zeeuw van Dalfsen, E., Ottemöller, L., Pimentel, A., Komorowski, J., Loughlin, S., Carmo, R., Coltelli, M., Corbeau, J., Vye-Brown, C., Di Vito, M., de Chabaliere, J., Ferreira, T., Fontaine, F. R., Lemarchand, A., Marques, R., Medeiros, J., Moretti, R., Pfeffer, M. A., Saurel, J., Vlastelic, I., Vogfjörð, K., Engwell, S., & Salerno, G. (2024). The European Volcano Observatories and their use of the aviation colour code system. **Bulletin of Volcanology**, 86(3). <https://doi.org/10.1007/s00445-024-01712-0>