

Pipelined Diagonal Matrix Code Encoder/Decoder for 5G/6G Systems

¹Dunna kiran kumar, ²Dodla Sriram, ³TankalaVijayaLaxmi, ⁴Kadumula Rama, ⁵Pandiri Hemlatha

¹²³⁴⁵Dept. of Electronics & Communication Engineering Avanthi'S St.Theressa Institute of Engineering & Technology Vizianagaram, India

Abstract- The increasing data rates and stringent reliability requirements of 5G and emerging 6G communication systems demand high-throughput and low-latency error control coding techniques. Diagonal Matrix Codes (DMC) have gained attention due to their strong error detection and correction capabilities with moderate hardware complexity. This work presents a pipelined VLSI architecture for Diagonal Matrix Code encoder and decoder tailored for 5G/6G systems. The proposed design employs parallel and pipelined processing stages to enhance throughput while minimizing critical path delay. Both encoding and decoding operations are implemented using synthesizable hardware modules, enabling efficient integration into high-speed baseband processors. Performance evaluation demonstrates that the pipelined architecture achieves improved latency and higher operating frequency compared to non-pipelined implementations, making it suitable for next-generation wireless communication systems requiring reliable and real-time data transmission. Further the designs are modified with data aware adaptive activation with bypass correction logic and parity reuse for efficient hardware utilization.

Keywords- Diagonal Matrix Code (DMC), Error Control Coding, 5G Communication Systems, 6G Wireless Networks, High-Throughput Hardware, Low-Latency Systems.

1. INTRODUCTION

One of the most significant technological changes in contemporary digital systems is the quick development of wireless communication systems from the earliest cellular generations to the fifth-generation (5G) infrastructure and the expected sixth-generation (6G) paradigm. The shift to 5G and 6G redefines the architectural requirements of digital hardware in terms of ultra-reliability, ultra-low latency, extreme connectivity density, and real-time flexibility, in contrast to previous generational shifts that mostly concentrated on raising peak data rates. Massive machine-type communication (Mmtc), ultra-reliable low-latency communication (URLLC), and enhanced mobile broadband (Embb) are all expected to be supported by modern wireless networks while also meeting strict scalability and energy efficiency standards.

The target latency for upcoming 6G systems is anticipated to drop below 100 microseconds for crucial applications like autonomous systems, remote surgery, augmented reality, and intelligent edge computing. In 5G

networks, latency requirements have already been lowered to the order of milliseconds. The design of on-chip inter connect topologies, integrated memory, and baseband processors is greatly impacted by these stringent latency restrictions. Data continuously passes through intricate signal processing chains in these systems, which include equalizers, precoding units, channel estimation blocks, Fast Fourier Transform (FFT) modules, modulators, demodulators, and forward error correction engines. Because each of these phases depends on intermediate buffering and temporary data storage, internal memory dependability is essential to system performance as a whole.

Aggressive growth of semiconductor technology has also brought significant reliability issues. Transistor geometries grow more compact as CMOS devices decrease into the sub-micron and nano-meter regimes, while supply voltages are lowered to control dynamic power usage. The critical charge needed to change a stored logic state is decreased by these scaling patterns. Digital circuits are therefore more susceptible to transitory disruptions such electromagnetic interference, voltage noise, temperature fluctuations, and radiation-induced Single Event Upsets (SEUs). Moreover, charge-sharing effects raise the likelihood of Multiple Bit Upsets (MBUs), in which neighbouring memory cells flip at the same time, in high-density memory arrays.

Clustered and neighbouring bit errors have become more common in contemporary integrated circuits, especially in FPGA and ASIC-based communication processors. Clustered mistakes necessitate more complex correction algorithms that can recognize and fix several incorrect bits inside a single data word, in contrast to isolated single-bit errors that can be fixed using traditional Hamming codes. Such errors have the potential to disrupt channel estimation matrices, corrupt modulation symbols, spread throughout pipeline stages, and ultimately raise the bit error rate (BER) at the system level if they are not fixed.

To improve digital reliability, error detection and correction (EDAC) techniques have been used for a long time. Single-error correction with little redundancy and straightforward implementation is offered by traditional Hamming codes. Although SEC-DED methods may detect double faults, their correction strength is still restricted. Although algebraic codes like Reed-Solomon and BCH provide strong burst error correction, they necessitate intricate finite-field arithmetic operations that raise latency, power

consumption, and hardware complexity. Widely used for channel coding, iterative decoding schemes like Low-Density Parity-Check (LDPC) codes approach Shannon capacity. However, because of their variable latency and high switching activity, they are not as well suited for real-time systems that require lightweight internal memory protection.

An internal error correcting system in 5G and 6G baseband processors needs to meet several requirements at once. To avoid pipeline disturbance, the correction latency must first be deterministic and constrained. Second, high-frequency operation possibly in the hundreds of megahertz to several gigahertz range must be compatible with throughput. Third, in order to address clustering and neighbouring bit upsets, repair capacity needs to go beyond single-bit errors. Fourth, to guarantee space and power efficiency, hardware implementation must rely on basic logic operations. Lastly, in order to accommodate expanding data pathways and changing communication standards, the architecture needs to be scalable.

A potential class of structured codes that strike a balance between hardware efficiency and corrective capabilities are matrix-based error correction codes. Matrix codes improve error localization without adding complicated arithmetic by organizing data bits into two-dimensional matrices and implementing parity restrictions across several spatial dimensions. Both vertical and horizontal parity relationships are used in traditional matrix codes. These methods might, however, not be able to adequately adjust for diagonal or clustered error patterns.

This idea is expanded upon by diagonal-based matrix coding, which adds parity requirements along the data matrix's diagonal routes. Each data bit thus takes part in several independent parity equations in various spatial orientations. This multi-dimensional constraint structure boosts robustness against neighbouring multi-bit mistakes and enhances error distinguishability. A 32-bit data word arranged in a 4x8 matrix could be safeguarded by 24 diagonal parity bits, allowing for the correction of up to eight incorrect bits inside the word, as previously shown by the Diagonal Hamming Code (DHC) architecture. When compared traditional single-error correction systems, this is a significant improvement.

Additionally, pipelining methods have been used in diagonal coding schemes to reduce combinational delay. Higher operating frequencies are made possible by reducing the critical route length by adding registers in between the encoding and decoding processes. Even while pipelining enhances timing performance, more architectural adjustments are required to satisfy the demands of contemporary wireless systems for energy efficiency and streaming compatibility.

Because of external channel coding methods, the likelihood of mistake in internal memory or pipeline registers is typically low in high-speed communication contexts. Consequently, there is needless switching activity and an increase in dynamic power consumption when full correction logic is continuously activated for each data word. Reducing needless toggling is crucial for energy-efficient operation since dynamic power is proportional to switching activity, capacitance, supply voltage, and frequency. Furthermore, simple duplication of parity computation circuitry causes a linear increase in hardware resources as data lengths in next-generation communication systems increase. Scalability is restricted in the absence of logic sharing or hardware reuse, especially in FPGA systems where slice LUTs and routing resources are limited.

These findings highlight the need for an improved diagonal matrix coding architecture that incorporates adaptive activation, balanced pipeline segmentation, hardware reuse, and streaming-friendly design principles while maintaining robust multi-bit correction capacity. Thus, an Adaptive Pipelined Diagonal Matrix Code architecture designed especially for 5G and 6G communication systems is proposed in this work. The suggested solution uses parity reuse techniques to maximize hardware efficiency, selective corrective activation to decrease switching activity, bypass logic to lower effective latency in common-case scenarios, and hierarchical syndrome detection to discover error-free states early. The suggested architecture seeks to increase reliability, throughput, dynamic power consumption, and scalability for next-generation wireless infrastructures by integrating coding theory with realistic VLSI design considerations.

The suggested solution tackles the core issues of performance, energy efficiency, and dependability in contemporary communication processors by combining structured diagonal coding with adaptive pipelined architecture. In developing 5G and 6G networks, the resulting framework provides a strong basis for fault-tolerant, high-speed, real-time operation.

2. LITERATURE SURVEY

[1] High packing densities in sophisticated FPGAs make semiconductor memories more susceptible to temperature induced faults, which calls for effective error-detecting and correcting (EDAC) techniques for sensitive data. A pipelined matrix-based EDAC algorithm based on Hamming codes is shown in this study. It can correct up to 6-7 random bit mistakes and 8-bit burst errors in 8-bit data. The suggested design uses pipelining and optimized memory representation to increase code pace while decreasing area, latency, and power. Verilog-implemented and verified on a Zynq-7000 FPGA with Vivado 2023.2, the design surpasses current methods in terms of code and

technology.[2] In order to decrease the execution time of instructions, this study proposes a variable-length pipeline microcontroller design based on the RISC-V RV32IM ISA. Instructions without data dependencies can finish without needless waiting thanks to the implementation of an in-order dispatch with out-of-order writeback mechanism. With behavior handling of multiplication and division operations to behavior execution stages, the suggested architecture, which is based on the Hummingbird E200, allows for variable pipeline lengths for different instructions. The design reaches a 120 MHz operational frequency when implemented in TSMC 0.18 μm technology, which is a major improvement over the 50 MHz Hummingbird E200 in the same process.[3] The design and FPGA implementation of a multi-stage pipelined RISC-V processor utilizing System Verilog HDL are presented in this work. The processor incorporates essential modules memory, hazard unit, and pipeline registers for effective execution, utilizing the open-source and adaptable instruction set of RISC-V. Test benches and ModelSim–Intel FPGA Edition simulations are used to validate the design, which is then synthesized on an Altera DE1-SoC board. The suggested approach is appropriate for undergraduate VLSI design teaching and is consistent with contemporary processor design techniques.[4] For real-time applications, this study suggests a pipelined VLSI architecture based on a single MAC for the quick computation of the one-dimensional discrete wavelet transform (DWT). Using only high-pass and low-pass filters, the design efficiently shares resources to reduce computational complexity, and pipelining speeds up processing. Comparing the architecture to traditional parallel MAC designs, it achieves notable area and power reductions.

It is implemented in Verilog and realized on a Xilinx Virtex5 FPGA. According to experimental results, the working frequency is 276 MHz with a power consumption of 200 Mw, and the area is reduced by 44.8% and the power savings are 31.97%. [5] A machine-learning-based design framework that adaptively modifies clock frequency based on the propagation latency of individual instructions is proposed in this research. The TigerMIPS processor incorporates a random forest model as an extra pipeline step that uses execution history, operands, and operation type as features to classify instruction delays in real time. In 45 nm CMOS technology, gate-level evaluation shows a 70% speedup with 30% energy reduction by coarse-grained categorization. Further performance gains of up to 89% speedup with 15.5% energy savings are achieved using finer-grained delay classification.[6] This paper describes a hardware accelerator for LU decomposition that uses two linear arrays of processing engines mapped across different SLR regions and is implemented on an FPGA. Block LU decomposition calculations across the PE arrays are simplified and effectively scheduled by the suggested design. The design, when implemented on an Alveo U50 FPGA, yields an average throughput of 95 GFLOPS and a peak

performance of 128 GFLOPS. Compared to an Intel MKL implementation on a 4-core Xeon CPU, the accelerator delivers approximately $15\times$ latency speedup.[7] Crosspoint resistive memory array-based analogue matrix computing (AMC) circuits and the theoretical underpinnings of the AMC idea are presented in this study. To carry out essential matrix operations, such as multiplication, inversion, pseudoinverse, and eigenvector calculation, in a single step, four types of AMC circuits are shown. The use of external input, feedback settings, and methods for dealing with negative-valued matrices are among the fundamental design concepts covered in the study. Transfer function-based techniques are used to analyze circuit stability and temporal complexity.[8] A near memory floating-point computing architecture is presented in this study for effective sparse matrix–vector multiplication (SpMV) in AI systems with limited resources. Utilizing data locality and memory-level parallelism, the architecture lowers latency and boosts performance by integrating computing at the edge of memory arrays. The architecture provides high-performance floating-point operations and is designed for processing sparse data in parallel. When combined with a 4 KB SRAM array, the solution operates at 1 GHz and produces up to 370 MFLOPS with just a 17% area overhead.[9] In this research, a low-latency parallel approach based on the Jacobi technique that uses CORDIC rotations to efficiently compute the eigenvalues and eigenvectors of real symmetric matrices is proposed. Reducing the per-rotation computation time is the main contribution, which greatly lowers the parallel Jacobi method's total latency. This approach is used to construct a highly accurate and flexible VLSI architecture that supports real skew-symmetric, complex Hermitian, and complex skew-Hermitian matrices. In comparison to state-of-the-art alternatives, the design, when implemented on a Zynq Ultra Scale+ FPGA, delivers 172.75 MHz operation with 43.75% reduced latency and 89.4% increased accuracy.[10] The universal VLSI architecture for calculating matrix inverses and pseudoinverses utilizing QR decomposition and back substitution is shown in this work. The suggested architecture does away with expensive division and square-foot operations by using the Newton-Raphson and CORDIC algorithms. The findings demonstrate that the design can accommodate different matrix dimensions while still satisfying timing, resource, and real-time performance requirements. The architecture is appropriate for high-performance signal processing and control applications because it provides near-floating-point precision.[11] In this work, Gate Diffusion Input (GDI) logic is used to develop and simulate Hamming code encoder and decoder circuits for low-power digital communication systems. GDI is an effective VLSI design method that behavior size, latency, and power consumption while drastically lowering circuit complexity. Micro wind technology is used to model and simulate the suggested designs, and the results are contrasted with traditional CMOS-based implementations. According to simulation data, GDI is a better low-power design technique since it reduces power by over

50% while also improving delay performance and transistor count.[12] The difficulty of creating effective error detection and repair codes for System-on-Chip (SoC) systems with constrained memory and processing power is discussed in this study. In order to lower the bit error rate in the waterfall area, an improved turbo coding strategy is suggested by linking the information and parity bits of neighboring code blocks. In order to behavior link length and regulate hardware complexity, the design restricts coupling to nearby blocks, potentially lowering electromagnetic interference. The suggested method works better than current turbo coding strategies, according to a performance study based on error rate vs repetition count. [13] This study describes the construction of circuits for low power digital communication that use CNTFET technology in conjunction with inverter-based pass transistor logic to encode and decode Hamming codes. Compared to traditional logic approaches, the suggested method maintains minimal circuit complexity while providing a better trade-off between power, delay, and area. The suggested method is used to develop encoder and decoder architectures, which are then simulated in 32 nm technology using the HSPICE tool. The advantages of the suggested design over conventional CMOS-based implementations are shown by comparative findings.[14] Convolutional encoder-decoder machine learning models that operate in the frequency domain are used in this study to provide a low-power hardware architecture for on-device audio denoising. A quantized network is created to maintain enhanced quality while drastically lowering memory and computational complexity by as much as 75%. A processing- element activation routing strategy is implemented to further increase efficiency, lowering on-chip memory access energy by 3.24 Mj/frame at 0.65 V and 18.5 MHz. When compared to previous efforts, the suggested chip provides the greatest PESQ audio quality score and operates in real-time.[15] By investigating the memory effect in atomically thin two-dimensional (2D) nanomaterials, this work advances our understanding of non-volatile memory devices and their engineering applications. The paper focuses on innovative advancements in monolayer memory devices (atomristors), where a virtual conductive point model explains how switching ehavior results from atomic vacancies and metal adsorption. Few-layer 2D materials, on the other hand, behave similarly to traditional conductive-bridge memory. Promising uses for these 2D atomristors include memristive devices for neuromorphic computing, non-volatile RF switches for 5G/6G systems, and zero-power electronics.

3. EXISTING DESIGN

The Diagonal Hamming Code (DHC) architecture, which was created to enhance multi-bit error correcting capability in embedded memory systems, is the foundation of the current design. The DHC approach's core idea is to create parity constraints along diagonal channels by

rearranging a linear data word into a structured two-dimensional matrix. In contrast to traditional horizontal or vertical parity methods, the diagonal coding structure improves error localization capacity by taking advantage of spatial correlations among bits in the matrix representation. With the use of XOR-based logic, the design seeks to deliver better multi-bit implementation.

A 32-bit data word is converted into a 4×8 matrix arrangement in this architecture.

The matrix representation includes geometric positioning that allows structured grouping of bits across diagonal directions, as opposed to presenting the data as a straightforward sequential bit stream. In order to maintain spatial adjacency relationships across various orientations, each diagonal group is made up of four data bits that are chosen from various rows and columns. Modulo-2 addition operations are used to generate three parity bits for each diagonal group. The 4×8 matrix contains eight of these diagonals, therefore for every 32 data bits, a total of 24 superfluous bits are created. Consequently, the encoded codeword has 56 bits, resulting in a code rate of roughly 0.57. For similar correction strength, this redundancy is still far lower than that of more intricate algebraic codes like BCH or Reed-Solomon, even though it is larger than that of basic SEC-DED schemes.

The DHC can be understood as a structured linear block code defined over the binary Galois field from a coding-theoretic perspective. The code's parity-check matrix is built using diagonal parity equations; each row of the matrix represents a distinct diagonal constraint. The requirement that the product of the codeword vector and the parity-check matrix equal zero is satisfied by a valid codeword. The diagonal parity bits are recalculated and compared with the stored parity bits to determine the syndrome vector during decoding. The word is deemed error-free if the syndrome vector is zero. Bit inversion allows for repair if the syndrome is non-zero since its pattern identifies the locations of incorrect bits.

The multi-dimensional constraint enforcement of the diagonal coding system is its strongest point. The overlap of constraint coverage is increased since each data bit takes part in several diagonal parity equations. Multiple neighbouring faults impact multiple parity equations at the same time. Even in cases when errors are clustered, this overlapping activation results in unique syndrome patterns that enable precise error localization. Up to eight incorrect bits inside the 32-bit word might be fixed by the design. Compared to traditional Hamming-based SEC-DED methods, which detect double faults and fix only one error, this correction strength is noticeably stronger. Even if the correcting capability has improved, there are still certain difficulties with the DHC's hardware implementation. Combinational XOR networks, which produce diagonal parity bits from chosen data bits, make up the majority of

the encoder. The combinational delay of the encoder is determined by the depth of the XOR trees, despite the fact that XOR logic is comparatively straightforward and well-supported in FPGA and ASIC implementations. A parity equation's equivalent XOR tree deepens with an increase in the number of inputs, which lengthens the propagation delay. Additionally, the connectivity overhead caused by FPGA routing limitations may increase this latency. More complexity is added by the decoder architecture. All diagonal parity bits must be recalculated, a multi-bit syndrome vector must be created, the incorrect bit positions must be determined by interpreting the syndrome, and the relevant bits must be inverted to remedy the problem. These operations are cascaded within a single combinational path in a non-pipelined implementation. As a result, the decoder's critical path lengthens considerably, lowering its maximum working frequency. Such lengthy combinational pathways limit throughput and can result in timing violations at higher clock frequencies, which makes them unattractive for high-speed communication systems.

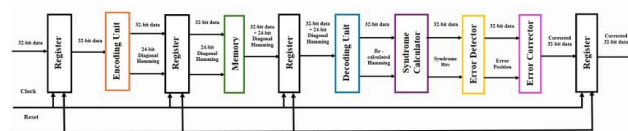


Fig1: Pipelined Diagonal Hamming Code (PDHC).

A Pipelined Diagonal Hamming Code (PDHC) architecture was presented in order to alleviate the timing constraints of the non-pipelined DHC. In this version, pipeline registers were added to the encoder and decoder before and after significant logic blocks. The effective combinational delay per stage was decreased by splitting the combinational logic into many stages that were separated by registers. In FPGA implementations, this segmentation allowed for better timing closure and higher clock frequency operation. Comparing the experimental synthesis results to the original non-pipelined architecture, it was shown that the critical path delay and logic power consumption were reduced. Nevertheless, pipelining did not significantly change the corrective mechanism's operating behavior, even though it increased frequency scalability. Whether or not there was a mistake, the correction logic continued to operate for each processed word. Because of external channel coding, the internal bit error probability is usually low in real-world communication systems. Consequently, there is needless switching activity when the whole corrective network is continuously activated. This static activation pattern helps prevent unnecessary energy loss because dynamic power consumption in CMOS circuits is proportional to switching activity, capacitance, supply voltage, and clock frequency.

Furthermore, each diagonal group's parity generating networks were constructed separately in the current design.

Despite having a similar structure, these networks were not designed with hardware reuse in mind. Consequently, redundant XOR operations led to an increase in gate count in ASIC realizations and slice LUT Consumption in FPGA implementations. Scalability is restricted by this redundancy, especially when expanding the design to accommodate wider data lines that are necessary for future communication standards.

The current architecture's focus on embedded memory protection rather on streaming data environments is another drawback. The architecture was not initially designed to integrate into high-throughput baseband pipelines, where deterministic one-word-per-cycle processing is crucial, even if it is appropriate for memory read-write protection applications.

In conclusion, the pipelined version of the current Diagonal Hamming Code and its robust multi-bit correction capability offer enhanced timing performance and modest redundancy. Nevertheless, there are still restrictions on streaming compatibility, adaptive corrective control, hardware reuse, and energy economy. The creation of an improved architecture that maintains corrective strength while meeting the performance and power needs of upcoming 5G and 6G communication systems is driven by these limitations.

4. PROPOSED DESIGN

While maintaining their robust multi-bit correction capacity, the suggested Adaptive Pipelined Diagonal Matrix Code architecture is designed to overcome the energy and performance constraints seen in the traditional Diagonal Hamming Code and its pipelined counterpart. The primary goal of the suggested architecture is to convert diagonal matrix coding from a memory-focused defence mechanism into a high-throughput, power-conscious, and streaming-compatible architecture that can be incorporated into 5G and soon-to-be 6G communication systems. The likelihood of internal data corruption is typically low in high-speed wireless communication contexts since external channel coding methods like LDPC or polar codes are present. However, the resulting clustered bit flips need to be repaired deterministically within stringent temporal limits when internal errors arise as a result of radiation-induced disturbances, voltage fluctuations, or process variances. In order to improve operating frequency, lower dynamic power consumption, and increase hardware scalability, the suggested architecture is made to retain the capacity to correct up to eight incorrect bits inside a 32-bit data word.

In the suggested architecture, the encoder restructures the computation of diagonal parity into pipeline phases that are balanced. The XOR networks that generate parity are dispersed among several pipeline levels rather than calculating every diagonal parity bit in a single

combinational block. The maximum possible clock frequency is raised by carefully balancing the logic depth of each step in order to reduce the maximum combinational delay per stage. The critical path is not dominated by any one diagonal parity path thanks to this balanced pipeline segmentation. Thus, the architecture is able to work at high frequencies, which are characteristic of contemporary FPGA-based baseband processors and ASIC implementations in wireless infrastructure. Hierarchical syndrome evaluation is introduced by the architecture at the decoder side to increase energy efficiency. A lightweight initial detection step determines whether there is a parity mismatch in the received word instead of immediately triggering full correction logic for each data word. The correction circuitry stays dormant and the data is delivered straight to the output pipeline register when the calculated syndrome vector equals zero, which indicates that there is no error. This bypass situation predominates during regular operation because the likelihood of internal errors occurring in well-designed communication systems is statistically low. As a result, the correction network's switching activity is much decreased, which lowers the dynamic power consumption.

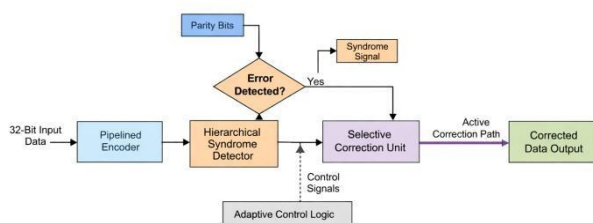


Fig2: Proposed Adaptive Pipelined Diagonal Matrix Code Architecture

The architecture selectively engages correction logic based on the diagonal group determined by the syndrome pattern when a non-zero syndrome is found. Only the rectification channel corresponding to the impacted diagonal is engaged, rather than inversion networks across all data bits. In unaffected logic regions, this focused activation prevents needless toggling and significantly lowers switching activity. In CMOS circuits, dynamic power is related to operating frequency, capacitance, supply voltage, and switching activity. Therefore, reducing needless toggling directly increases energy efficiency without sacrificing throughput.

The optimization of parity computation through hardware reuse is another noteworthy improvement in the suggested architecture. Conventional methods result in duplicate XOR operations since parity formulas for various diagonal groups are calculated independently. However, a thorough examination shows that diagonal equations share a number of intermediate XOR terms. The total number of XOR gates is decreased by reorganizing the parity generating network to reuse common sub-expressions. This results in less routing congestion and lower slice LUT consumption in FPGA implementations.

Reduced connection capacitance and gate count in ASIC implementations enhance area efficiency and may result in lower leakage power.

By guaranteeing that one encoded or decoded 32-bit word is processed per clock cycle following initial pipeline filling, the suggested design preserves deterministic performance. The possible throughput is $32 \cdot f$ bits per second if the operational frequency is represented by f . In 5G and 6G systems, where fluctuating decoding latency is undesired, this stable and predictable throughput characteristic is especially crucial for ultra-reliable low-latency communication scenarios. In contrast to iterative decoding systems that create temporal uncertainty and necessitate numerous cycles per codeword, the suggested diagonal matrix code architecture corrects in a single pipeline pass.

Reliability-wise, the design maintains the capacity to fix up to eight incorrect bits in a 32-bit word. The likelihood of correctable error patterns, given a bit error probability p , encompasses all possible combinations with up to eight incorrect bits. In contrast to SEC-DED techniques that simply fix one issue, the reliability coverage is significantly improved. Additionally, the system's overall power consumption is reduced by lowering switching activity through adaptive activation and bypass logic. Long-term dependability is improved by lower soft error rates, which are indirectly caused by lower junction temperatures brought on by lower dynamic power. Another important factor in the suggested design is scalability. By increasing matrix dimensions while maintaining diagonal grouping principles, the diagonal matrix structure can be proportionately increased as communication systems move toward wider data routes in 6G architectures. To maintain controlled combinational depth, the balanced pipeline methodology can be expanded appropriately. Deeper pipelining adds more registers, but the increased frequency scalability and energy-delay optimization make the trade-off worthwhile. Future extensions will be feasible because hardware complexity increases roughly linearly with word size.

All things considered, the suggested Adaptive Pipelined Diagonal Matrix Code design optimizes correction capability, operating frequency, power efficiency, and hardware utilization in a balanced manner. In order to meet the performance and reliability needs of next-generation wireless communication systems, the design incorporates structured diagonal coding with adaptive logic control and pipeline balancing. For real-time 5G and 6G infrastructures, it converts diagonal matrix coding from a static memory protection method into a dynamic, high-performance, and energy-conscious error control solution.

5. RESULTS AND DISCUSSION

The pipelined designs are developed in Xilinx Vivado tool for the hardware of choice is 28nm CMOS technology based Artix-7 FPGA board (xc7z020clg484-1). The system requirements include windows 10 OS, 4 GB RAM installed with Xilinx Vivado 2023.2. The simulation result of proposed design is as shown in fig. 3.

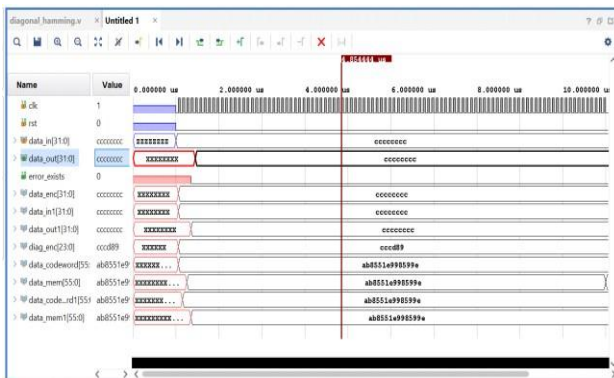


Fig3: Proposed Simulation Waveform

The suggested Diagonal Matrix Code (DMC) encoder and decoder architecture's proper hardware implementation is confirmed by the technical diagrams produced following synthesis. The diagram shows how input ports, combinational logic blocks, and output ports used for error correction and parity generation are connected. The synthesized design makes effective use of hardware with 76 logic cells, 67 I/O ports, and 143 nets. The pipelined processing method, which lowers critical route time and increases throughput, is validated by the structured logic structure. These findings show that the suggested architecture is appropriate for high-speed 5G/6G communication systems that need error control coding that is dependable and low-latency.

The fig.4(a),4(b) and 4(c) shows technology schematic of the existing and proposed designs.

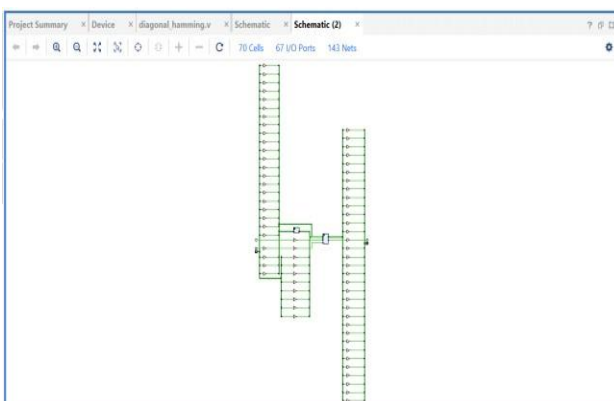


Fig a) Existing Technology Schematic Design

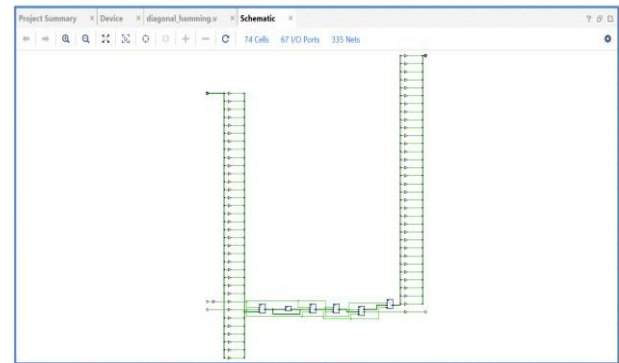


Fig b) Pipelined Technology Schematic Design

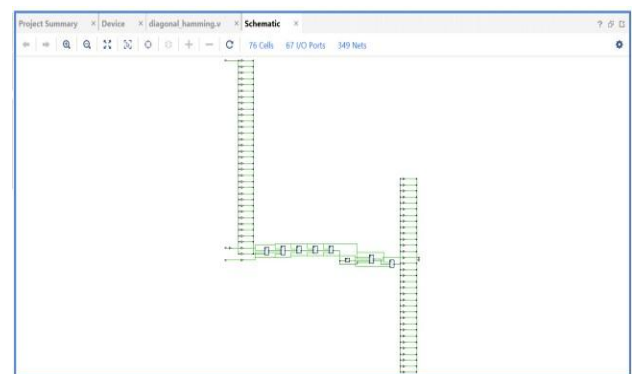


Fig I Proposed Technology Schematic Design

The performance of the pipelined, suggested, and current Diagonal Matrix Code designs is contrasted in Table X. With critical path delays of 35.898 ns and 35.388 ns, respectively, the pipelined and current implementations show comparable time characteristics. The suggested architecture, on the other hand, dramatically lowers the delay to 5.791 ns, demonstrating better timing performance. Additionally, power analysis reveals that the suggested design lowers overall on-chip power usage to 10.018 W as opposed to 23.46 W and 23.36 W for the pipelined and current systems. Adaptive activation techniques, parity

The suggested design is better suited for high-speed 5G/6G communication systems since it exhibits lower latency, lower power consumption, and increased hardware efficiency.

Table1: Comparison Results Among Existing & Proposed Designs

Parameter	Existing Design	Pipelined Design	Proposed Design
Maximum Setup	35.898	35.388	5.791
Logic Delay (ns)	8.149	8.271	3.122
Net Delay (ns)	27.749	27.117	2.669
Maximum Hold	0.676	0.257	0.301
Total On-Chip Power (W)	23.46	23.36	10.018
Dynamic Power (W)	22.421	22.324	8.979
Static Power (W)	1.039	1.039	1.039
Signal Power (W)	3.759	3.774	0.870
Logic Power (W)	3.878	3.714	0.162
I/O Power (W)	14.785	14.835	7.948
Junction Temperature (°C)	125	125	125

6. CONCLUSION

An effective pipelined Diagonal Matrix Code (DMC) encoder and decoder architecture for fast 5G/6G communication systems was presented in this paper. To increase hardware efficiency and processing performance, the suggested architecture uses parity reuse, data-aware adaptive activation, parallel processing, and bypass correction logic. Compared to the current and pipelined designs, the suggested architecture dramatically decreases the critical path delay to 5.791 ns and the overall on-chip power consumption to 10.018 W, according to implementation and synthesis data. The design is appropriate for next-generation wireless communication systems because of these enhancements, which show increased throughput, decreased latency, and optimal hardware use. In order to offer scalable and energy-efficient solutions for ultra-high data rate 6G communication networks, future work will concentrate on further refining the design using sophisticated low-power approaches and optimized FPGA/ASIC implementations.

7. REFERENCES

1. Kavya, C. H., & Varadarajan, S. (2025). Pipelined diagonal matrix codes for error correction in embedded memories. *Engineering, Technology & Applied Science*

Research, 15(5), 28201–28207. <https://doi.org/10.48084/etasr.12069>.

2. Yen, M.-H., Tsou, C.-H., Lin, T.-F., Lin, Y.-H., Ku, Y.- F., & Kao, C.-T. (2023). VLSI implementation of RISC-V MCU with a variable stage pipeline. In *Proceedings of the 2023 IEEE 6th International Conference on Knowledge Innovation and Invention (ICKII)* (pp. 161–165). IEEE. <https://doi.org/10.1109/ICKII58656.2023.10332786>.

3. Loh, S. H., Tan, G. Y., & Sim, J. J. (2025). VLSI design course with FPGA implementation of pipelined RISC-V processor. In *Proceedings of the 2025 IEEE 15th International Conference on Control System, Computing and Engineering (ICCSCE)* (pp. 36–41). IEEE. <https://doi.org/10.1109/ICCSCE65566.2025.11182670>.

4. Vaithyanathan, D., James, B. P., & Mariammal, K. (2024). A high-speed computational pipeline single MAC-based VLSI architecture for real-time signal and image processing. In *Proceedings of the 2024 Fourth International Conference on Advances in Electrical, Computing, Communication and Sustainable Technologies(ICAECT)*(pp.16).IEEE. <https://doi.org/10.1109/ICAECT60202.2024.10469571>.

5. Ajirlou, A. F., & Partin-Vaisband, I. (2022). A machine learning pipeline stage for adaptive frequency adjustment. *IEEE Transactions on Computers*, 71(3), 587–598. <https://doi.org/10.1109/TC.2021.3057764>

6. Kumar, M. K., Choudry, Z., & Purini, S. (2023). Accelerating LU decomposition of arbitrarily sized matrices on FPGAs. In *Proceedings of the 2023 International VLSI Symposium on Technology, Systems and Applications (VLSI-TSA/VLSI-DAT)* (pp. 1–4). IEEE. <https://doi.org/10.1109/VLSI-TSA/VLSI-DAT57221.2023.10134072>

7. Sun, Z. (2022). Analog matrix computing with resistive memory: Circuits and theory. In *Proceedings of the 2022 International Symposium on VLSI Technology, Systems and Applications(VLSI-TSA)*(p.12).IEEE. <https://doi.org/10.1109/VLSI-TSA54299.2022.9770998>

8. Eggermann, G., Rios, M., Ansaloni, G., Nassif, S., & Atienza, D. (2023). A 16-bit floating-point near-SRAM architecture for low-power sparse matrix-vector multiplication. In *Proceedings of the 2023 IFIP/IEEE31st International Conference on Very Large Scale Integration (VLSI-SoC)* (pp. 1–6).IEEE. <https://doi.org/10.1109/VLSISoC57769.2023.10321838>

9. Sharma, R., Shrestha, R., & Sharma, S. K. (2022). Low-latency and reconfigurable VLSI architectures for computing eigenvalues and eigenvectors using CORDIC-

based parallel Jacobi method. IEEE Transactions on Very Large Scale Integration (VLSI) Systems,30(8),1020–1033. <https://doi.org/10.1109/TVL SI.2 022.3170526>

10. P. S., V., & Mula, S. (2024). A parameterized VLSI architecture for real-time inversion of square matrices and Moore–Penrose pseudoinversion of rectangular matrices. In Proceedings of the 2024 31st IEEE International Conference on Electronics, Circuits and Systems (ICECS) (pp. 1–4). IEEE. <https://doi.org/10.1109/ICECS61496.2024.10848994>

11. Mythry, S. V., HarshaVardhini, P. A., Bandaru, T., Gunamgari, S. R., Bandari, K., & Balagoni, N. (2024). The low power implementation of Hamming-code encoder and decoder using GDI logic for error-free transmission and reception in digital data communication. In Proceedings of the 2024 2nd International Conference on Computer, Communication and Control (IC4) (pp. 1–6).IEEE.<https://doi.org/10.1109/IC457434.2024.10486366>

12. Wamankar, R., & Raghuwanshi, A. (2024). Design of recursive decoding-based coupled turbo encoder- decoder for VLSI applications. In Proceedings of the 2024 International Conference on Inventive Computation Technologies (ICICT) (pp. 1591–1597). IEEE. <https://doi.org/10.1109/ICICT60155.2024.105 44543>

13. Lokesh, P., Kumar, M. M., & Rambabu, S. (2025). High-efficiency CNTFET-based inverter and pass transistor design for Hamming code encoder and decoder for fault-tolerant data transmission. In Proceedings of the 2025 Fourth International Conference on Smart Technologies, Communication and Robotics (STCR) (pp. 1–6). IEEE. <https://doi.org/10.1109/STCR62650.2025.11020446>

14. Kochar, D. V., Ashok, M., & Chandrakasan, A. P. (2025). A 0.75 mm² 407 μ W real-time speech audio denoiser with quantized cascaded redundant convolutional encoder-decoder for wearable IoT devices. In Proceedings of the 2025 38th International Conference on VLSI Design and 23rd International Conference on Embedded Systems (VLSID)(pp.180–185).IEEE.<https://doi.org/10.1109/VLSID64188.2025.00044>

15. Materials for memory, computing and 6G applications. (2022). In Proceedings of the 2022 International Symposium on VLSI Technology, Systems and Applications (VLSI-TSA) (p. i). IEEE. <https://doi.org/10.1109/VLSI-TSA54299.2022.9770976>