

Design and Implementation of a Neural Network Based Predistorter for Enhanced Mobile Broadband

Chance Tarver*, Alexios Balatsoukas-Stimming^{†‡}, and Joseph R. Cavallaro*

*Department of Electrical and Computer Engineering, Rice University, Houston, TX, USA

[†]Department of Electrical Engineering, Ecole polytechnique fédérale de Lausanne Lausanne, Switzerland

[‡]Department of Electrical Engineering, Eindhoven University of Technology, Eindhoven, The Netherlands

Abstract—Digital predistortion is the process of using digital signal processing to correct nonlinearities caused by the analog RF front-end of a wireless transmitter. These nonlinearities contribute to adjacent channel leakage, degrade the error vector magnitude of transmitted signals, and often force the transmitter to reduce its transmission power into a more linear but less power-efficient region of the device. Most predistortion techniques are based on polynomial models with an indirect learning architecture which have been shown to be overly sensitive to noise. In this work, we use neural network based predistortion with a novel neural network training method that avoids the indirect learning architecture and that shows significant improvements in both the adjacent channel leakage ratio and error vector magnitude. Moreover, we show that, by using a neural network based predistorter, we are able to achieve a 42% reduction in latency and 9.6% increase in throughput on an FPGA accelerator with 15% fewer multiplications per sample when compared to a similarly performing memory-polynomial implementation.

Index Terms—Digital predistortion, neural networks, FPGA.

I. INTRODUCTION

Efficiently correcting nonlinearities in power amplifiers (PAs) through digital predistortion (DPD) is critical for enabling next-generation mobile broadband where there may be multiple radio frequency (RF) transmit (TX) chains arranged to form a massive multiple-input multiple-output (MIMO) system [1], as well as new waveforms with bandwidths on the order of 100 MHz in the case of mmWave communications [2]. Traditional DPDs use variations of the Volterra series [3], such as memory polynomials [4, 5]. These models consist of sums of various order polynomials and finite impulse response (FIR) filters to model the nonlinearities and the memory effects in a PA, respectively.

To learn the values of the parameters in a polynomial based model, an indirect learning architecture (ILA) is typically used in conjunction with some variation of a least squares (LS) fit of the data to the model [5]. In an ILA, a postinverse model of the predistorter is fitted based on the output of the PA [6, 7]. After learning the postinverter, the coefficients are copied to the predistorter. Although this simplifies the learning of DPD coefficients, it has been shown to converge to a biased

solution due to noise in the PA output [8, 9]. Moreover, the LS problem is often poorly conditioned [4]. In [10], a mobile graphics processing units (GPU) was used to implement the polynomial DPD with I/Q imbalance correction from [4]. This GPU implementation used floating-point and was able to avoid the challenges associated with the dynamic range requirements for memory polynomials. When implemented on an FPGA, a memory polynomial can be challenging due to the bit-widths that are necessary to perform the high-order exponentiation in fixed-point precision [11].

The overall DPD challenge has strong similarities to the problems encountered in in-band full-duplex (IBFD) communications [12–14], where a transceiver simultaneously transmits and receives on the same frequency, increasing the spectral efficiency of the communication system. However, this requires (among other techniques) digitally removing the significant self-interference from the received signal which not only consists of the intended transmission but also the nonlinearities added by the imperfections in the transmit chain including the PA. In [15], the author used neural networks (NNs) to perform the self-interference cancellation and found that it could achieve similar performance to polynomial based self-interference cancellation. This work was later extended to create both FPGA and ASIC implementations of the NN-based self-interference canceller [16]. It was found that, due to the regular structure of the NN and the lower bit-width requirements, it can be implemented to have both a higher throughput and lower resource utilization.

Inspired by the full-duplex NN work and the known problems of polynomial based predistortion with an ILAs, we recently proposed in [17] to use NNs for the forward DPD application. The NNs are a natural choice for such application as they are able to approximate any nonlinear function [18], making them a reasonable candidate for predistortion. The idea of using various NNs for predistortion has been explored in many works [19, 20]. However, the training method is unclear in [19], and their implementations require over ten thousand parameters. In [20], the training of the NN is done using an ILA which can subject the learned predistorter to the same problems seen with all ILAs.

Contribution: In our previous work [17], we avoided the standard ILA and we improved the overall performance by using a novel training algorithm where we first modeled the PA with a NN and then backpropagated through it to train a

The work of C. Tarver and J. R. Cavallaro was supported in part by the U.S. NSF under grants ECCS-1408370, CNS-1717218, and CNS-1827940, for the “PAWR Platform POWDER-RENEW: A Platform for Open Wireless Data-driven Experimental Research with Massive MIMO Capabilities.” The work of A. Balatsoukas-Stimming was supported by the Swiss NSF project PZ00P2_179686.

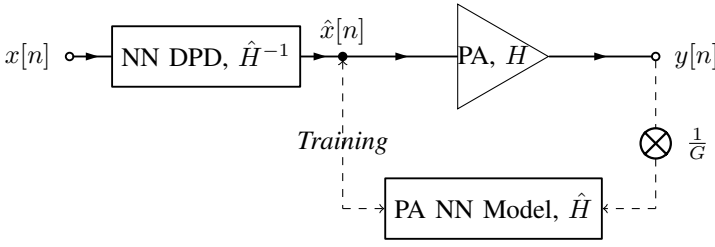


Figure 1. Architecture of the NN DPD system. The signal processing is done in the digital baseband and focuses on PA effects. The DAC, up/downconverters, and ADC are not shown in this figure, though their impairments are also captured.

DPD NN. We extend that work here to show that not only do we improve performance when compared to polynomial based DPD, but we do so with reduced implementation complexity. Furthermore, to realize the gains of the NN DPD, we design a custom FPGA accelerator for the task and compare it to our own polynomial DPD accelerator.

Outline: The rest of the paper is organized as follows. In Section II, we give an overview of our DPD architecture and methods. In Section III, we compare performance/complexity tradeoffs for the DPD NN to polynomial based predistorters. In Section IV, we compare FPGA implementations for memory polynomial and NN predistortion. Finally, in Section V we conclude the paper.

II. NEURAL NETWORK DPD ALGORITHM OVERVIEW

For the NN DPD system, we seek to place a NN based predistorter inline with the PA so that the cascade of the two is a linear system, as shown in Fig. 1. However, to train a NN, it is necessary to have training data, and in this scenario the ideal NN output is unknown; only the ideal PA output is known. To overcome this problem, we train a PA NN model to emulate the PA. We then backpropagate the mean squared error (MSE) through the PA NN model to update the parameters in the NN DPD [17].

A. Neural Network Architecture

We use a feed-forward NN that is fully-connected with K hidden layers, and N neurons per hidden layer. The nonlinear activation applied in hidden layers is chosen to be a rectified linear unit (ReLU), shown in (1), which can easily be implemented with a single multiplexer in hardware.

$$\text{ReLU}(x) = \max(0, x) \quad (1)$$

The input and output data to the predistorter is complex-valued, while NNs typically operate on real-valued data. To accommodate this, we split the real and imaginary parts of each time-domain input sample, $x(n)$, on to separate neurons.

Although PA-induced nonlinearities are present in the transmitted signal, the relationship between the input and output data is still mostly linear. Although in principle, a NN can learn this relationship given training data, this turns out to be difficult in practice [15]. As such, we implement a linear

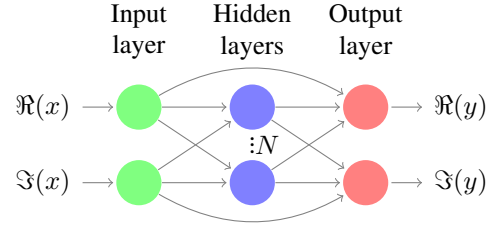


Figure 2. General structure of the DPD and PA neural networks. There are two input and output neurons for the real and imaginary parts of the signal, N neurons per hidden layer, and K hidden layers. The inputs are directly added to the output neurons so that the hidden layers concentrate on the nonlinear portion of the signal.

bypass in our NN that directly passes the inputs to the output neurons where they are added in with the output from the final hidden layer, as can be seen in Fig. 2. This way, the NN entirely focuses on the nonlinear portion of the signal.

B. Training

This work primarily focuses on the implementation and running complexity of the DPD application, which consists of inference on a pre-trained NN. The training is assumed to be able to run offline and, once the model is learned, significant updates will not be necessary and occasional offline re-training to account for long-term variations would be sufficient.

In [17], we first use input/output data of the PA to train a NN to model the PA behavior. We then connect a second DPD NN to the PA NN model. We treat the combined DPD NN and PA NN as one large NN. However, during the second training phase, we only update the weights corresponding to the DPD NN. We then connect the DPD NN to the real PA and use it to predistort for the actual device.

The process of predistorting can excite a different region of the PA than when predistortion is not used. To account for this, it is not uncommon in other DPD methods to have multiple training iterations. A similar idea is adopted in [17] and in this work. Once training of the PA and the DPD is performed, we then retransmit through the actual PA while using the DPD NN. Using the new batch of input/output data, we then can update the PA NN model and in turn refine the DPD NN. An example of the iterative training procedure is shown in Fig. 3, where the MSE training loss is shown for the PA NN model and the combined DPD-PA is shown for two training iterations.

III. COMPLEXITY COMPARISON

To evaluate the NN based predistortion, we present the formulation of both a memory polynomial and the NN. We then derive expressions for the number of multiplications as a function of the number of parameters in the models. In most implementations, multiplications are considered to be more expensive as they typically have higher latency and require more area and power. Additions typically have a minor impact on these metrics when compared to multiplications, so we omit them from this high-level analysis.

$$\hat{x}(n) = \sum_{\substack{p=1, \\ p \text{ odd}}}^P \sum_{m=0}^M \alpha_{p,m} x(n-m) |x(n-m)|^{p-1} + \sum_{\substack{q=1, \\ q \text{ odd}}}^Q \sum_{l=0}^L \beta_{q,l} x^*(n-l) |x^*(n-l)|^{q-1} + c \quad (2)$$

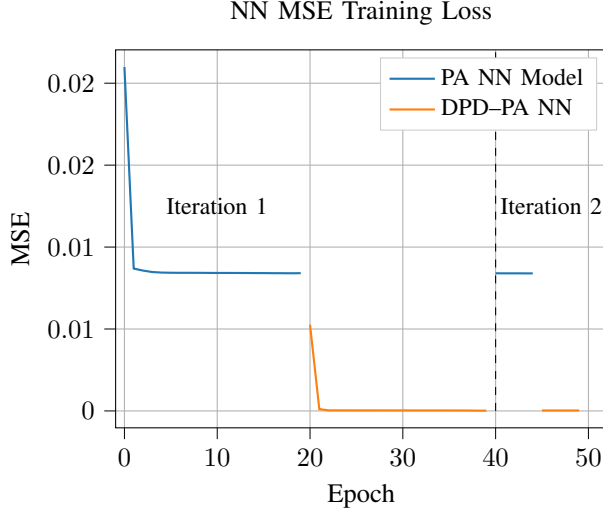


Figure 3. Example of iterative NN-DPD training for two training iterations, where 20 and 5 epochs are used in the first and second iteration, respectively.

A. Memory Polynomial Predistortion

An extension of a memory polynomial from [4] is shown in (2). This form of memory polynomial predistorts the complex baseband PA input $x(n)$ to be $\hat{x}(n)$ by computing nonlinearities of the form $x(n)|x(n)|^p$ and convolving them with an FIR filter for both $x(n)$ and its conjugate, $x^*(n)$. This conjugate processing gives the model the expressive power to combat PA nonlinearities and any IQ imbalance in the system. P and M are the highest nonlinearity order and memory depth in the main branch, while Q and L are the highest order and memory in the conjugate branch. The complex-valued coefficients $\alpha_{p,m}$ and $\beta_{q,l}$ represent the DPD coefficients that need to be learned for nonlinearity orders p and q and memory tap m and l . Finally, the DC term c accounts for any local oscillator leakage in the system.

The total number of complex-valued parameters in (2) is given as

$$n_{\text{PAR, poly}} = M \left(\frac{P+1}{2} \right) + L \left(\frac{Q+1}{2} \right) + 1. \quad (3)$$

Assuming three real multiplications per complex multiplication, we get the following number of multiplications in the system

$$n_{\text{MUL, poly}} = 3n_{\text{PAR, poly}} + \sum_{\substack{p=3, \\ p \text{ odd}}}^P \frac{1}{2} (p+5) + \sum_{\substack{q=3, \\ q \text{ odd}}}^Q \frac{1}{2} (q+5). \quad (4)$$

Here, each complex coefficient accounts for three multiplication. The expression, $x(n)|x(n)|^{p-1}$ is computed once for each n over a given p and delayed in the design to generate the appropriate value for each m . We note that $|x(n)|^{p-1}$ can always be simplified to $(\Re(x(n))^2 + \Im(x(n))^2)^{\frac{p-1}{2}}$ since p is odd. This accounts for $(\frac{p-1}{2} + 1)$ multiplications before being multiplied by the complex-valued $x(n)$ which adds 2 more multiplications. The same is true for the conjugate processing.

B. Neural Network Predistortion

The output of a densely connected NN is given by

$$\mathbf{h}_1(n) = f \left(\mathbf{W}_1 \begin{bmatrix} \Re(x(n)) \\ \Im(x(n)) \end{bmatrix} + \mathbf{b}_1 \right), \quad (5)$$

$$\mathbf{h}_i(n) = f(\mathbf{W}_i \mathbf{h}_{i-1}(n) + \mathbf{b}_i), \quad i = 2, \dots, K, \quad (6)$$

$$\mathbf{z}(n) = \mathbf{W}_{K+1} \mathbf{h}_K(n) + \mathbf{b}_{K+1} + \mathbf{W}_{\text{linear}} \begin{bmatrix} \Re(x(n)) \\ \Im(x(n)) \end{bmatrix}, \quad (7)$$

$$\hat{x}(n) = z_1(n) + 1j \cdot z_2(n), \quad (8)$$

where f is a nonlinear activation function (such as the ReLU from (1)), $\mathbf{W}_i$ and $\mathbf{b}_i$ are weight matrices and bias vectors corresponding to the i th layer in the NN, and j is the imaginary unit. The final output of the network after hidden layer K is given by (7) where the first element represents the real part of the signal, and the second element represents the imaginary part. In (7), $\mathbf{W}_{\text{linear}}$ is a 2×2 matrix of the weights corresponding to the linear bypass. In practice, we fix it to be the identity matrix, $\mathbf{I}_2$, to reduce complexity though these weights could also be learned in systems with significant IQ imbalance.

Assuming N neurons per hidden layer and K hidden layers, the number of multiplications is given by

$$n_{\text{MUL, NN}} = 4N + (K-1)N^2. \quad (9)$$

C. Results

The performance results for each predistorter as a function of the number of required multiplications are shown in Figs. 4–6. These results were obtained using the RFWebLab platform [21]. RFWebLab is a web-connected PA at Chalmers University. This system uses a Cree CGH40006-TB GaN PA with a peak output power of 6 W. The precision is 14 bits for the feedback on the ADC and 16 bits for the DAC. Using their MATLAB API, we test the NN predistorter using a 10 MHz OFDM signal. This signal has random data on 600 subcarriers spaced apart by 15 kHz and is similar to LTE signals commonly used in cellular deployments. It provides an interesting test scenario in that it has a sufficiently high

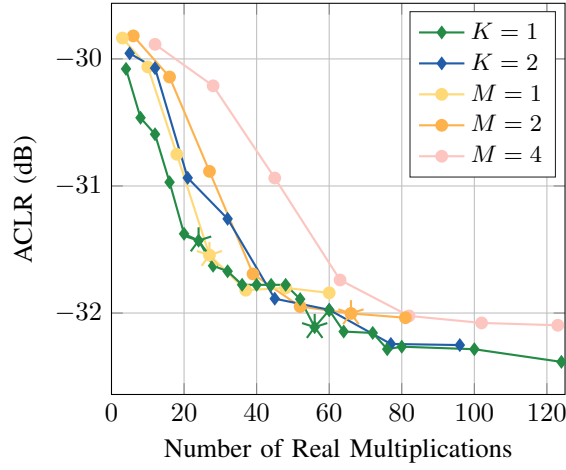


Figure 4. ACLR vs. number of multiplications for NN DPD (shown with diamonds) with up to $K = 2$ hidden layers and memory polynomial (shown with circles) with up to $M = 4$ memory taps. This represents the out-of-band performance of the predistorter. The stars represent design points that we implement in FPGA in the next section.

peak-to-average power ratio (PAPR) to make predistortion challenging. We train on 10 OFDM symbols then validate and present experimental results based on averaging over 10 different symbols. The Adam optimizer is used with an MSE loss function and batches of 32 samples. ReLU activation functions are used in the hidden layer neurons.

Specifically, we tested the following DPDs: (1) a NN DPD with $K = 1$ with $N = \{1, \dots, 20, 25, 31\}$ (dark green), (2) a NN DPD with $K = 2$ with $N = \{1, \dots, 8\}$ (light green), (3) a polynomial DPD without memory and with $P = 1$ to $P = 13$ (dark blue), (4) a polynomial DPD with $M = 2$ memory taps and with $P = 1$ to $P = 13$ (light blue), and (5) a polynomial DPD with $M = 4$ memory taps and with $P = 1$ to $P = 13$ (pink). For each of these polynomials, no conjugate processing was used hence $Q, L = 0$. A predistorter with $M = L = 4$ and $Q = P$ was also evaluated. However, the system did not have significant IQ imbalance, so the addition of the conjugate processing to the memory polynomial only had the effect of significantly increasing complexity. All DPDs were evaluated in terms of the adjacent channel leakage ratio (ACLR), the error vector magnitude (EVM), and the spectra of the post-PA pre-distorted signals.

1) *Out-of-band performance:* To measure the out-of-band performance, which is often the metric of most interest given by Federal Communications Commission (FCC) regulations and 3GPP standards, we compute the ACLR shown below as

$$\text{ACLR} = 10 \log_{10} \frac{P_{\text{adjacent}}}{P_{\text{channel}}}, \quad (10)$$

where P_{channel} is the signal power in the main channel, and P_{adjacent} is the signal power in the remainder of the band.

In Fig. 4, we observe that the NN DPD offers similar performance to the memoryless polynomial DPD for low numbers

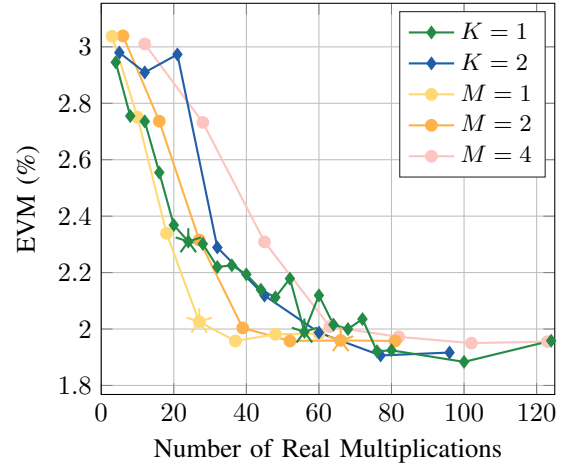


Figure 5. EVM vs. number of real multiplications for NN DPD (shown with diamonds) with up to $K = 2$ hidden layers and memory polynomial (shown with circles) with up to $M = 4$ memory taps. This represents the in-band performance of the predistorter. The stars represent design points that we implement in FPGA in the next section

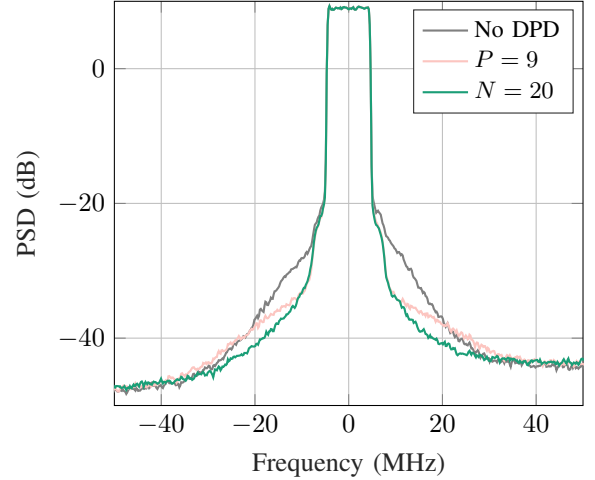


Figure 6. Example spectrum for the $M = 4$ polynomial and $K = 1$ NN. Each of these use around 80 multiplications per time-domain input sample to the DPD.

of multiplications and it is able to significantly outperform all polynomial DPDs as the number of multiplications increases.

2) *In-band performance:* Although the primary goal of predistortion is to reduce spectral regrowth around the main carrier, predistortion also reduces the EVM of the main signal. Reducing EVM can improve reception quality and is hence a desirable result. The EVM is computed as

$$\text{EVM} = \frac{\|\hat{\mathbf{s}} - \mathbf{s}\|}{\|\mathbf{s}\|} \times 100\%, \quad (11)$$

where $\mathbf{s}$ is the vector of all original symbols mapped onto complex constellations on OFDM subcarriers in the frequency domain, $\hat{\mathbf{s}}$ is the corresponding received vector after passing through the PA, and $\|\cdot\|$ represents the ℓ^2 norm.

In Fig. 5, we see the EVM versus the number of multiplications for each of the predistorters. As the number of multiplications increases, the EVM decreases, as expected. The memoryless polynomial DPD is able to achieve a low EVM for the smallest number of multiplications. However, the complexity is only slightly higher for the NN based DPD, which is able to achieve an overall better performance than all other examined polynomial DPDs.

3) *Spectrum Comparison*: The spectrum for both the memory polynomial and the NN DPDs are shown in Fig. 6. Here, both predistorters have the same running complexity of 80 multiplications per time-domain input sample. However, the NN is able to provide an additional 2.8 dB of suppression at ± 20 MHz.

IV. FPGA ARCHITECTURE OVERVIEW

In this section, we compare a NN DPD accelerator with a memory polynomial based implementation. We implement both designs in Xilinx System Generator and target for the Zynq UltraScale+ RFSoc ZCU1285 evaluation board. For the sake of this architecture comparison, we implement each to be fully parallelized and pipelined as to compare the highest throughput implementations of each. Based on the previous analysis, we implement both with 16-bit fixed point precision throughout.

We synthesize FPGA designs targeting two separate ACLRs. First, we target an ACLR of approximately -31.4 dB. This target is achieved with a NN with $N = 6$ neurons and $K = 1$ hidden layer and a 7th order memoryless polynomial. Second, we target a more aggressive ACLR below -32 dB. This is done with a NN with $N = 14$ neurons and $K = 1$ hidden layer. A memory polynomial with $M = 2$ and $P = 11$ is also used to achieve this.

A. Neural Network Accelerator

We implement the NN-DPD on FPGA with the goal of realizing high throughput via maximum parallelization and pipelining. The top-level overview of the design is shown in Fig. 7. Here, each wire corresponds to a 16-bit bus. The real and imaginary parts of the PA input signal stream in each clock cycle. Weights are stored in a RAM which can be written to from outside the FPGA design. After the RAM is loaded, the weights and biases are written to individual registers in the neuron processing elements (PEs) which cache them for fast access during inference. A chain of pipeline registers pass the inputs to the output to be added to the output of the final layer.

After the weights are loaded into RAM, the RAM controller loads each of the weights into a weights cache in each PE. To do this, a counter increments through each address in the RAM. The current address and the value at that address are broadcast to all neurons. Each address corresponds with a specific weight or bias. Whenever the weights cache in a neuron reads addresses corresponding to the weights and biases for its neuron, it saves the data into a register dedicated to that parameter. These registers output to the corresponding multiplier or adder.

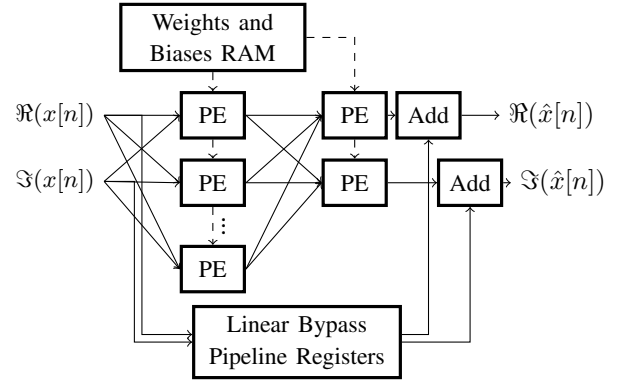


Figure 7. General structure of the NN FPGA implementation.

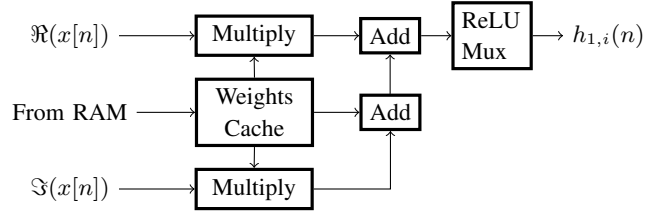


Figure 8. Example structure of a PE for the i th neuron in hidden layer 1.

An example neuron PE is shown in Fig. 8. Each PE is implemented with a sufficient number of multipliers for performing the multiplication of the weights by the inputs in parallel. The results from each multiplier are added together, along with the bias and passed to the ReLU activation function, which is implemented with a single multiplexer.

B. Polynomial Accelerator

The memory polynomial is also implemented using 16 bits throughout the design. We target the design for maximum throughput by fully parallelizing and pipelining it so that a new time-domain input sample can be streamed in each clock cycle. The main overall structure of the design is shown in Fig. 9. Each polynomial “branch” of the memory polynomial corresponding to nonlinear order p computes $x(n)|x(n)|^{p-1}$ and there is a branch for each p in the design. This computation from each branch is passed to an FIR filter with complex taps. Three multiplications are used for each complex multiplication in each filter. A RAM is implemented to interface with some outside controller for receiving updated weights. Once the coefficients α and β are loaded into the design, they can be moved from the RAM to registers near each multiply similarly to the cache implemented in the NN design.

C. Results

The Xilinx Vivado post-place-and-route utilization results are shown in Table I. Overall, the NN-based design offers numerous advantages over the memory polynomial. Specifically, for the target of an ACLR less than -32 dB, the NN

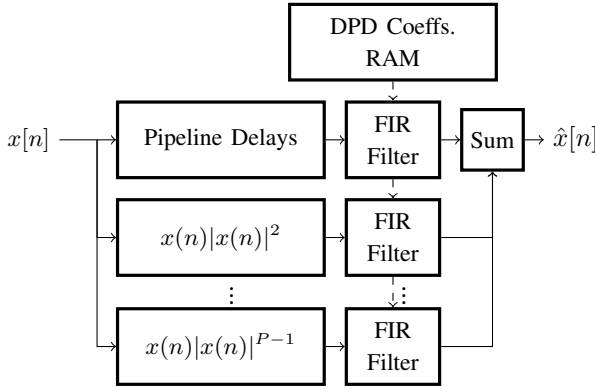


Figure 9. General structure of the high-throughput, low-latency, memory polynomial FPGA implementation.

Table I
COMPARISON OF PERFORMANCE AND FPGA UTILIZATION

Metric	ACLR: -31.4 dB		ACLR: -32	
	$N = 6$ $K = 1$	$P = 7$ $M = 1$	$N = 14$ $K = 1$	$P = 11$ $M = 2$
Num. of Params.	32	8	72	24
LUT	379	539	688	1424
LUTRAM	16	120	16	224
FF	538	991	1170	2730
DSP	24	27	56	66
Worst Neg. Slack (ns)	8.72	8.68	8.49	8.34
Max. Freq. (MHz)	783	756	661	603
Max. T/P (MS/s)	783	756	661	603
Latency (CC)	12	21	14	26

requires 48% of the lookup tables (LUTs), 42% of the flip-flops (FFs), and 15% reduction in the number of digital signal processors (DSPs). In terms of timing, there is a 9.6% increase in throughput with a 46% decrease in latency. These reductions in utilization occur while also seeing improved ACLR.

V. CONCLUSIONS

In this paper, we explored the complexity/performance tradeoffs for a novel, NN based DPD and found that the NN could outperform memory polynomials and offered overall unrivaled ACLR and EVM performance. Furthermore, we implemented each on an FPGA and found that the regular matrix multiply structure in the NN based predistorter led to a lower latency design with less hardware utilization when compared to a similarly performing polynomial-based DPD.

This work opens up many avenues for future work. This work can be extended to also compare performance/complexity tradeoffs for more devices with a wider variety of signals, including different bandwidths and multiple component carriers. It is also possible to include memory cells such as recurrent neural networks (RNNs) in the NN to account for memory effects. The NN is naturally well suited for a GPU implementation which would be interesting in software defined radio (SDR) systems. The NN complexity could also be further reduced with pruning, and the accuracy could

potentially be improved with retraining after quantization and pruning.

REFERENCES

- [1] E. G. Larsson, O. Edfors, F. Tufvesson, and T. L. Marzetta, "Massive MIMO for next generation wireless systems," *IEEE Commun. Mag.*, vol. 52, no. 2, pp. 186–195, Feb. 2014.
- [2] W. Roh *et al.*, "Millimeter-wave beamforming as an enabling technology for 5G cellular communications: Theoretical feasibility and prototype results," *IEEE Commun. Mag.*, vol. 52, no. 2, pp. 106–113, Feb. 2014.
- [3] A. Zhu, M. Wren, and T. J. Brazil, "An efficient Volterra-based behavioral model for wideband RF power amplifiers," in *IEEE MTT-S Int. Microw. Symp. Digest*, vol. 2, June 2003, pp. 787–790 vol.2.
- [4] L. Anttila, P. Handel, and M. Valkama, "Joint mitigation of power amplifier and I/Q modulator impairments in broadband direct-conversion transmitters," *IEEE Trans. Microw. Theory Techn.*, vol. 58, no. 4, pp. 730–739, Apr. 2010.
- [5] A. Katz, J. Wood, and D. Chokola, "The Evolution of PA Linearization: From Classic Feedforward and Feedback Through Analog and Digital Predistortion," *IEEE Microw. Mag.*, vol. 17, no. 2, pp. 32–40, Feb. 2016.
- [6] A. Balatsoukas-Stimming, A. C. M. Austin, P. Belanovic, and A. Burg., "Baseband and RF hardware impairments in full-duplex wireless systems: experimental characterisation and suppression," *EURASIP Journal on Wireless Commun. and Networking*, vol. 2015, no. 142, 2015.
- [7] D. Korpi, L. Anttila, and M. Valkama, "Nonlinear self-interference cancellation in MIMO full-duplex transceivers under crosstalk," *EURASIP Journal on Wireless Comm. and Netw.*, vol. 2017, no. 1, p. 24, Feb. 2017.
- [8] D. Zhou and V. E. DeBrunner, "Novel adaptive nonlinear predistorters based on the direct learning algorithm," *IEEE Trans. on Signal Processing*, vol. 55, no. 1, pp. 120–133, Jan. 2007.
- [9] R. N. Braithwaite, "A comparison of indirect learning and closed loop estimators used in digital predistortion of power amplifiers," in *IEEE MTT-S Int. Microw. Symp.*, May 2015, pp. 1–4.
- [10] K. Li *et al.*, "Mobile GPU accelerated digital predistortion on a software-defined mobile transmitter," in *IEEE Global Conf. on Signal and Inform. Process. (GlobalSIP)*, Dec. 2015, pp. 756–760.
- [11] M. Younes, O. Hammi, A. Kwan, and F. M. Ghannouchi, "An accurate complexity-reduced "PLUME" model for behavioral modeling and digital predistortion of RF power amplifiers," *IEEE Trans. on Ind. Electron.*, vol. 58, no. 4, pp. 1397–1405, Apr. 2011.
- [12] M. Jain *et al.*, "Practical, real-time, full duplex wireless," in *Proc. Int. Conf. on Mobile Comput. and Netw.* ACM, 2011, pp. 301–312.
- [13] M. Duarte, C. Dick, and A. Sabharwal, "Experiment-driven characterization of full-duplex wireless systems," *IEEE Trans. Wireless Commun.*, vol. 11, no. 12, pp. 4296–4307, Dec. 2012.
- [14] D. Bharadia, E. McMillin, and S. Katti, "Full duplex radios," in *ACM SIGCOMM*, 2013, pp. 375–386.
- [15] A. Balatsoukas-Stimming, "Non-linear digital self-interference cancellation for in-band full-duplex radios using neural networks," in *IEEE Int. Workshop on Signal Processing Advances in Wireless Commun. (SPAWC)*, June 2018, pp. 1–5.
- [16] Y. Kurzo, A. Burg, and A. Balatsoukas-Stimming, "Design and implementation of a neural network aided self-interference cancellation scheme for full-duplex radios," in *Asilomar Conf. on Signals, Systems, and Comput.*, Oct. 2018, pp. 589–593.
- [17] C. Tarver, L. Jiang, A. Sefidi, and J. Cavallaro, "Neural network DPD via backpropagation through a neural network model of the PA," in *Asilomar Conf. on Signals, Systems, and Comput.*, (to appear).
- [18] K. Hornik, "Approximation capabilities of multi-layer feedforward networks," *Neural Networks*, vol. 4, no. 2, pp. 251 – 257, 1991. [Online]. Available: <http://www.sciencedirect.com/science/article/pii/089360809190009T>
- [19] R. Hongyo, Y. Egashira, T. M. Hone, and K. Yamaguchi, "Deep neural network-based digital predistorter for doherty power amplifiers," *IEEE Microw. and Wireless Compon. Letters*, vol. 29, no. 2, pp. 146–148, Feb. 2019.
- [20] M. Rawat and F. M. Ghannouchi, "Distributed spatiotemporal neural network for nonlinear dynamic transmitter modeling and adaptive digital predistortion," *IEEE Trans. Instrum. Meas.*, vol. 61, no. 3, pp. 595–608, Mar. 2012.
- [21] "RF WebLab." [Online]. Available: <http://dpdcompetition.com/rfweblab/>