

Uncertainty Quantification with Invertible Neural Networks for Signal Integrity Applications

Osama Waqar Bhatti, Oluwaseyi Akinwande, Madhavan Swaminathan

School of Electrical and Computer Engineering, School of Material Science and Engineering
3D Systems Packaging Research Center (PRC), Georgia Institute of Technology, Atlanta, GA

Abstract—We present a machine learning based tool to quantify uncertainty for prediction problems regarding signal integrity. Harnessing invertible neural networks, we convert the inverse posterior distribution given by the network to address uncertainty in frequency responses as a function of design space parameters. As an example, we consider a differential plated-through-hole via in package core and predict S-parameters from its geometrical properties. Results show 3.3% normalized mean squared error when compared with responses from a fullwave EM simulator.

Index Terms—Invertible neural networks, jacobian, uncertainty quantification, differential via-pair, signal integrity

I. INTRODUCTION

Modern electronic systems comprise high dimensional design space parameters that have to be tuned in order to obtain desired circuit response. Before fabrication, these systems undergo iterative design cycles using simulation software. Traditional EM solvers, while accurate, solve complex partial differential equations iteratively to obtain frequency response from the geometrical properties of the electromagnetic structure. Such design schemes consume time and computational resources. Recently, machine learning (ML) techniques have proved quite promising for design optimization and uncertainty quantification for signal and power integrity problems [1]. Generally, the design space of an electromagnetic structure is parameterized and fed as an input to the ML framework which outputs the frequency response. Given a dataset $D = \{X, Y\}$ where X is the design space parameter set and Y is the response space, we can write:

$$Y = T(X) + \epsilon \quad (1)$$

where $T(\cdot)$ is the forward transformation and ϵ is standard Gaussian noise inherently present in the system. However, current ML schemes have deterministic outputs that have no information about the reliability of their predictions in the frequency space. One way to quantify uncertainty is by using probabilistic modeling. We assume that X and Y are random variables sampled from their prior distributions $p(X)$ and $p(Y)$ respectively. To achieve confidence bounds, we need to determine the conditional distribution $p(Y|X)$ and $p(X|Y)$. The latter distribution is the goal of inverse design. The problem of inverse design is to determine the set of input combinations that lead to the desired response, that is, to find the inverse surrogate model $G(\cdot) = T^{-1}(\cdot)$. Such a problem is inherently ill-posed, the reasons for which are twofold: (1)

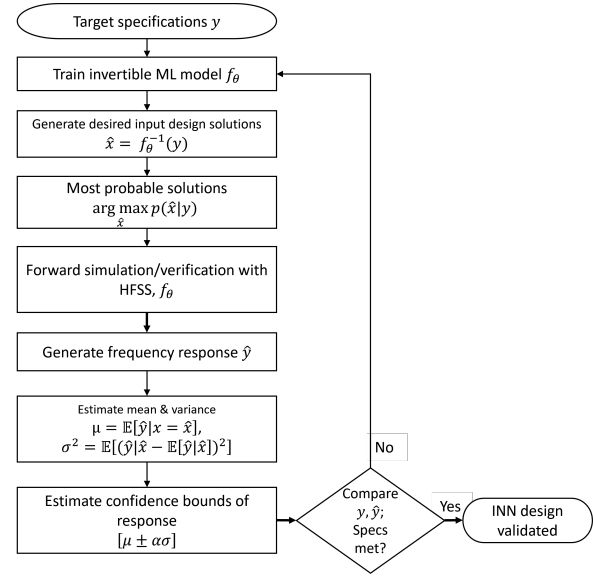


Fig. 1. Flow of inverse design and its uncertainty quantification.

Does the inverse transformation $G(\cdot)$ exist? (2) If the inverse exists, is it unique?

Several architectures have been proposed to address the problem of invertibility. In recent years, artificial neural networks (ANNs) have been deployed as an effective tool for microwave design and modeling problems [2]. State-of-the-art generative models like the generative adversarial network (GAN) [3], variational auto-encoder (VAE) [4] and invertible neural network (INN) have been developed to output posterior conditional distributions instead of a deterministic design solution. Not only does this enable the designer to have multiple candidate choices but also give an evaluation of the reliability of the model [5]. In contrast to other inverse methods, the INN used in our proposed approach provides mode stability and tractability.

In this paper, we propose to utilize INNs to achieve inverse posterior distribution $p(X|Y)$. We then pick the most probable sample points from this distribution and undergo a forward pass through the INN. This gives us the uncertainty bounds for the frequency response Y . This approach is illustrated in Fig. 1. As an example, we consider a differential via pair and parameterize its design space. The S-parameters of the resultant structure are the output response.

II. INVERTIBLE NEURAL NETWORKS

INNs are neural networks comprising stacked invertible blocks that can learn bijective transformations between function spaces.

A. Invertibility by Construction

The architecture of the INN addresses the existence, uniqueness and stability of inverse solutions. Given a sample x from design space X and the corresponding y from the response space Y that are related through the transformation $y = f(x)$, we can form a relationship between their probability densities through the change-of-variables technique [6], [7]:

$$p_Y(y|\theta) = p_X(f_\theta^{-1}(y)) \cdot \left| \left(\frac{df_\theta^{-1}}{dx} \right) \right|, \quad (2)$$

where we define all the composition of the INN architecture in a single function f_θ , and θ is the set of all network parameters. The INNs are made of stacks of reversible blocks with inputs x and outputs y , and they can be trained in both directions simultaneously, as shown in Fig. 2. In addition to the outputs y of the system, a set of latent variables z can be defined which encode the lost information in the forward direction. Variables z can be sampled from a standard normal distribution, which, when passed through the trained network in the reverse direction, conditioned on an output y , result in the conditional posterior distributions $p(x|y)$. Each reversible block is shown in Fig. 3. It is called affine coupling block and it ensures easy invertibility and a tractable Jacobian. The block's input vector is halved into $[x_1, x_2]$, and they are transformed by an affine function with coefficients e^s and t [7] [6]:

$$y_1 = x_1, \quad y_2 = x_2 \circ e^{s(x_1)} + t(x_1). \quad (3)$$

Given the block's output $[y_1, y_2]$, these expressions are invertible through:

$$x_1 = y_1, \quad x_2 = (y_2 - t(y_1)) \circ e^{-s(y_1)}. \quad (4)$$

(3) represents the forward mapping while (4) represents the inverse mapping (see Fig. 3 for a graphical illustration). The use of element-wise additive (+) and multiplicative ($\circ$) operations allows the inverse of the transformation to be easily computed without requiring the scale $s(\cdot)$ and shift $t(\cdot)$ networks to be inverted, which could be arbitrarily complex. The bijectivity of the INN model allows for bi-directional operation and training, and therefore both forward and inverse processes can be well learned [8]. We accumulate losses when the network is trained in the forward and reverse directions which are backpropagated to the network and the weights are automatically adjusted as part of the learning process.

The losses include the mean square error (MSE) and the maximum mean discrepancy (MMD). The INN is trained in both forward and reverse directions.

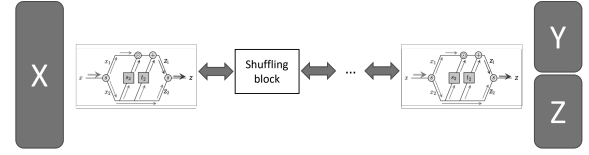


Fig. 2. Architecture of INN (x : input, y : output, z : latent variable)[9].

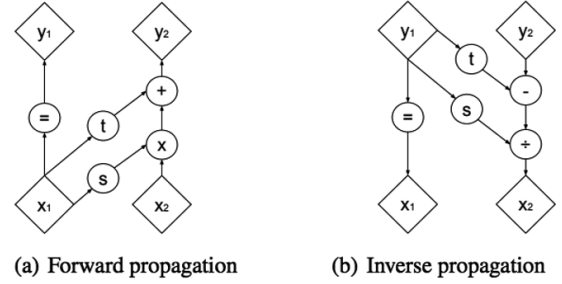


Fig. 3. RealNVP block enabling forward and backward propagations [7].

B. Addressing Uncertainty

After training the INN, inference is performed. We sample z coming from a known distribution $p(z)$, generally assumed to be a standard Gaussian. We append the sampled z with the target y_{target} and go through backward pass of the network to obtain the inverse posterior distribution $p(x|y_{target})$. The expected value of this distribution gives us the mean, and the variance of the distribution shows the sharpness of the input design tuple. The goal, here, is to obtain the variance and mean of the forward posterior distribution $p(y|x)$. To achieve this, we consider a range of the most probable input design tuples provided by INN and undergo a forward pass to obtain upper and lower confidence bounds for our frequency responses.

Algorithm 1: INN training

Input: training data: $\{X, Y\}$, $\#epochs$, learning rate: α , $p(z) = \mathcal{N}(0, \mathbb{I}_{D_z})$

Output: training losses: $\mathcal{L}$, trained model

```

1 while  $i \leq \#epochs$  do
2   for  $x_{batch}, y_{batch} \in \{X, Y\}$ , do
3      $[y_{pred}, z_{pred}] = f_\theta(x_{batch})$ 
4      $\mathcal{L}_y = \text{MSE}(y_{pred}, y_{batch})$ 
5      $\mathcal{L}_z = \text{MMD}(q(y, z), p(y)p(z))$ 
6     sample  $z \sim p(z)$ 
7      $x_{pred} = f_\theta^{-1}([y_{batch}, z])$ 
8      $\mathcal{L}_x = \text{MMD}(g(x), p(x))$ 
9      $\mathcal{L}_{total} = w_x \mathcal{L}_x + w_y \mathcal{L}_y + w_z \mathcal{L}_z$ 
10     $p \leftarrow p - \alpha \nabla(\mathcal{L}_{total})$ 
```

TABLE I
 CONTROL PARAMETERS OF THE PTH STRUCTURE

Parameter		Unit	Min	Max
μ -via Diameter	$d_{\mu\text{-via}}$	μm	30	70
μ -via Pad Diameter	$d_{\text{pad},\mu\text{-via}}$	μm	31	140
BU Layer Thickness	h_{BU}	μm	20	35
μ -via Top Antipad Radius	$r_{a,\text{BU},\text{TOP}}$	μm	100	500
μ -via Bot. Antipad Radius	$r_{a,\text{BU},\text{BOT}}$	μm	100	500
PTH Pitch	v_P	μm	300	1200
Core Thickness	h_{core}	μm	100	1200
BU Cu Thickness	$t_{c,\text{BU}}$	μm	10	20
Core Cu Thickness	$t_{c,\text{Core}}$	μm	11	40
PTH Diameter	d_{PTH}	μm	100	250
PTH Pad Diameter	$d_{\text{pad},\text{PTH}}$	μm	110	500
PTH Top Antipad Radius	$r_{a,\text{PTH},\text{TOP}}$	μm	50	500
PTH Bot. Antipad Radius	$r_{a,\text{PTH},\text{BOT}}$	μm	50	500

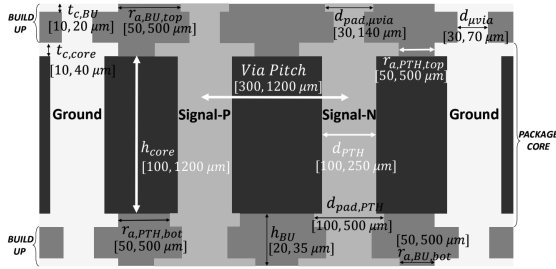


Fig. 4. Parameters of the differential PTH in package core [10].

III. EXAMPLE: DIFFERENTIAL PTH PAIR IN PACKAGE CORE

We consider an application of modeling a differential plated-through-hole (PTH) pair in package core along with the microvias that connect to build-up layers. Such structures are common since they enable vertical interconnection for signals. As such, the signal integrity of such differential vias is crucial for high-speed interfaces. We achieve an inverse surrogate model for the structure shown in Fig. 4. Obtaining the inverse posterior enables us to quantify uncertainty in the S-parameters of the shown structure.

A. Model Setup

The design space is parameterized as a 13- D input design tuple. The minimum and maximum values are shown in Table I. Each input combination in the design space has a corresponding four-port scattering (S) parameter matrix from 0.1-100 GHz with steps of 100 MHz. The objective is to determine an invertible mapping from the design space X and frequency response Y . Since the structure is partially reciprocal and symmetric, we only consider $S_{11}, S_{12}, S_{13}, S_{14}, S_{33}$ and S_{34} . We take the magnitude of the S-parameters, resulting in an output dimension of 6000. We draw 682 samples using Latin Hypercube Sampling (LHS) and obtain S-parameters using Ansys HFSS. The data is split into train and test sets

for the INN model. We use 500 samples for training and the remaining for evaluation of the model.

B. Results

We train the INN for 50 epochs with 100 iterations per epoch optimizing the model with an initial learning rate of $\alpha = 0.01$ using Adam optimizer. We train with an adaptive exponentially decreasing learning rate until the model converges. On random, we choose a desired response y_{target} from the test set. Next, we sample $z \sim p(z)$ for 5,000 times to obtain $x = f_{\theta}^{-1}(y, z)$. For this application, we choose the dimensionality of z to be 1000. The inverse distributions for each dimension of x is shown in Fig. 5. We also plot the prior distributions before conditioning on y_{target} . Starting with a uniform prior, we see that the posterior inverse distribution becomes dense around a certain range of design tuples that the model suggests are most likely to produce the target distribution. For each dimension, we choose the design tuple for which the model is most confident. These input combinations are fed back into the INN to obtain a set of frequency responses. This range determines the lower and upper bounds for the target frequency responses.

We simulate the chosen input ranges into a forward simulator to obtain confidence intervals. In Fig. 6, we plot the y_{target} from the test set coming from the 3D EM solver. We compare it with the output from the INN. We find that the mean of the predicted frequency responses from the INN closely matches the test set values. Specifically, we use the normalized mean-squared error as loss metric over each frequency response in the test set:

$$NMSE = \frac{1}{N_d D_y} \times \sum_{d=1}^{D_y} \sum_{n=1}^{N_d} \times \left(\frac{\sum_{m=1}^N (S_{n,d}[m] - \hat{S}_{n,d}[m])^2}{\sum_{m=1}^N (S_{n,d}[m] - \frac{1}{N} \sum_{m=1}^N \hat{S}_{n,d}[m])^2} \right) \quad (5)$$

where N_d are the number of evaluation designs for the model and $D_y = 6$ represents the magnitude of the learnt S-parameters. The NMSE value for the proposed approach is 3.3%.

IV. CONCLUSION

We propose a method to perform uncertainty quantification of frequency response as a function of design space parameters using invertible neural networks for signal integrity applications. Specifically, we illustrate a differential plated-through-hole pair in package core as an example. We provide lower and upper confidence bounds for output 4-port S-parameters. We achieve a normalized mean-squared error of 3.3% on the test set.

ACKNOWLEDGMENT

This material is based upon work supported by the National Science Foundation under Grant No. CNS 16-24810-Center for Advanced Electronics through Machine Learning (CAEML). The authors would also like to pay their gratitude

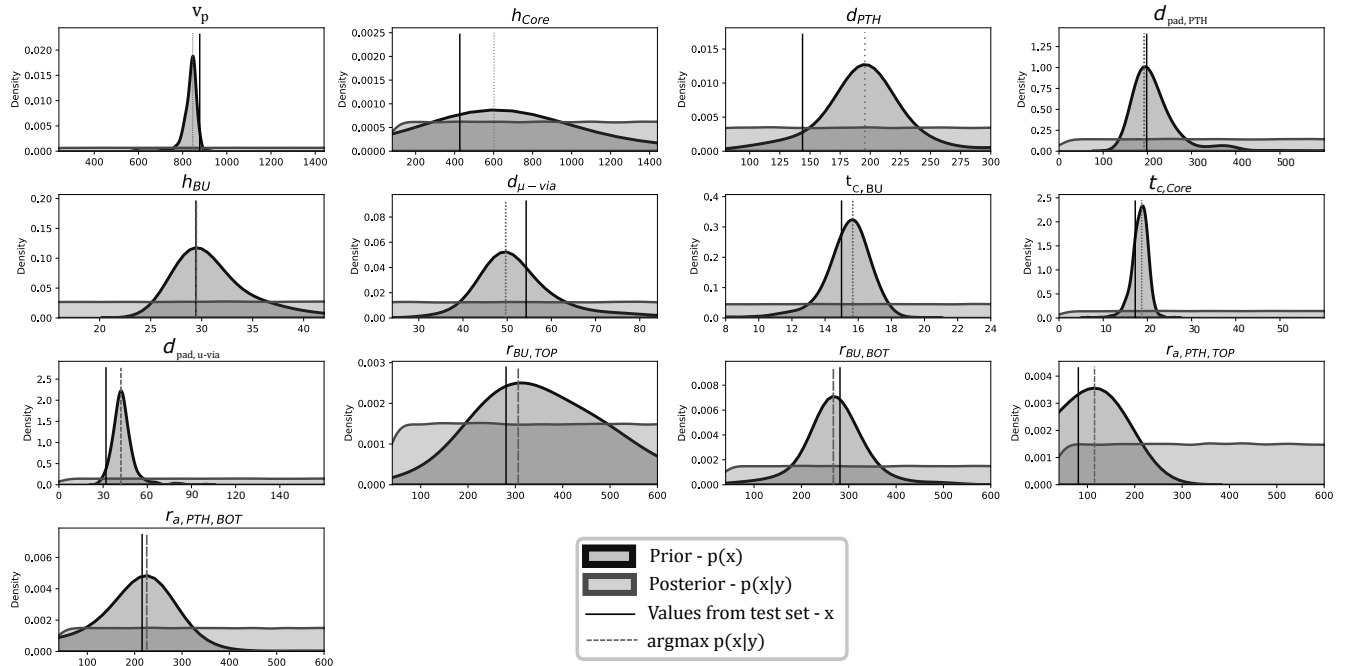


Fig. 5. Inverse posterior distributions $p(\mathbf{x}|y_{target})$, black vertical line shows values from the test set

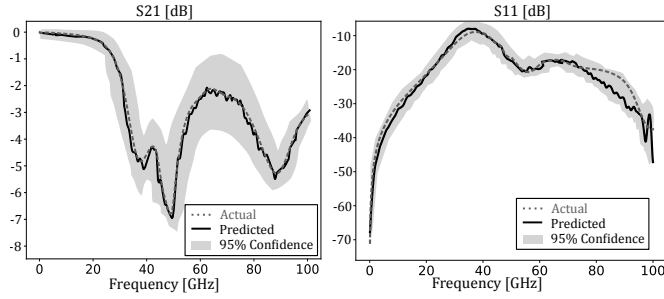


Fig. 6. Forward simulation results comparison for INN predictions with 3D EM solvers

to Hakki Mert Torun for providing help. The authors wish to acknowledge the assistance and support of the NEMO Steering Committee.

REFERENCES

- [1] M. Swaminathan, H. M. Torun, H. Yu, J. A. Hejase, and W. D. Becker, "Demystifying machine learning for signal and power integrity problems in packaging," *IEEE Transactions on Components, Packaging and Manufacturing Technology*, vol. 10, no. 8, pp. 1276–1295, 2020.
- [2] H. Kabir, Y. Wang, M. Yu, and Q.-J. Zhang, "Neural network inverse modeling and applications to microwave filter design," *IEEE Transactions on Microwave Theory and Techniques*, vol. 56, no. 4, pp. 867–879, 2008.
- [3] M. Mirza and S. Osindero, "Conditional generative adversarial nets," *CoRR*, vol. abs/1411.1784, 2014.
- [4] D. P. Kingma and M. Welling, "Auto-encoding variational bayes," *arXiv preprint arXiv:1312.6114*, 2013.
- [5] O. W. Bhatti, O. Akinwande, N. Ambasana, X. Yang, P. R. Paladhi, W. D. Becker, and M. Swaminathan, "Comparison of invertible architectures for high speed channel design," in *2021 IEEE Electrical Design of Advanced Packaging and Systems (EDAPS)*, pp. 1–3, 2021.
- [6] L. Ardizzone, J. Kruse, S. Wirkert, D. Rahner, E. W. Pellegrini, R. S. Klessen, L. Maier-Hein, C. Rother, and U. Köthe, "Analyzing inverse problems with invertible neural networks," 2019.
- [7] L. Dinh, J. Sohl-Dickstein, and S. Bengio, "Density estimation using real nvp," 2017.
- [8] H. Yu, H. M. Torun, M. U. Rehman, and M. Swaminathan, "Design of siw filters in d-band using invertible neural nets," in *2020 IEEE/MTT-S International Microwave Symposium (IMS)*, pp. 72–75, 2020.
- [9] O. W. Bhatti, N. Ambasana, and M. Swaminathan, "Inverse design of power delivery networks using invertible neural networks," in *2021 IEEE 30th Conference on Electrical Performance of Electronic Packaging and Systems (EPEPS)*, pp. 1–3, 2021.
- [10] H. M. Torun, A. C. Durgun, K. Aygün, and M. Swaminathan, "Enforcing causality and passivity of neural network models of broadband s-parameters," in *2019 IEEE 28th Conference on Electrical Performance of Electronic Packaging and Systems (EPEPS)*, pp. 1–3, 2019.