

STRUCTURE PROBING NEURAL NETWORK DEFLATION

YIQI GU

DEPARTMENT OF MATHEMATICS, NATIONAL UNIVERSITY OF SINGAPORE, 10 LOWER KENT
RIDGE ROAD, SINGAPORE, 119076 (MATGUY@NUS.EDU.SG)

CHUNMEI WANG

DEPARTMENT OF MATHEMATICS & STATISTICS, TEXAS TECH UNIVERSITY, 1108 MEMORIAL
CIRCLE, LUBBOCK, TX 79409, USA (CHUNMEI.WANG@TTU.EDU)

HAIZHAO YANG

DEPARTMENT OF MATHEMATICS, PURDUE UNIVERSITY, 150 N UNIVERSITY ST, WEST
LAFAYETTE, IN 47907, USA (HAIZHAO@PURDUE.EDU)

Abstract. Deep learning is a powerful tool for solving nonlinear differential equations, but usually, only the solution corresponding to the flattest local minimizer can be found due to the implicit regularization of stochastic gradient descent. This paper proposes a network-based structure probing deflation method to make deep learning capable of identifying multiple solutions that are ubiquitous and important in nonlinear physical models. First, we introduce deflation operators built with known solutions to make known solutions no longer local minimizers of the optimization energy landscape. Second, to facilitate the convergence to the desired local minimizer, a structure probing technique is proposed to obtain an initial guess close to the desired local minimizer. Together with neural network structures carefully designed in this paper, the new regularized optimization can converge to new solutions efficiently. Due to the mesh-free nature of deep learning, the proposed method is capable of solving high-dimensional problems on complicated domains with multiple solutions, while existing methods focus on merely one or two-dimensional regular domains and are more expensive in operation counts. Numerical experiments also demonstrate that the proposed method could find more solutions than exiting methods.

Key words. Neural Networks Deflation; Structure Probing; Nonlinear Differential Equations; High Dimension; Deep Least-Square Method; Convergence.

AMS subject classifications. 65M75; 65N75; 62M45;

1. Introduction.

1.1. Problem statement. Nonlinear differential equations are ubiquitous in various important physical models such as fluid dynamics, plasma physics, solid mechanics, and quantum field theory [30, 17, 22, 49], as well as chemical and biological models [76, 18]. Solving nonlinear differential equations has been a very challenging problem especially when it is important to find multiple distinct solutions. The nonlinearity of the differential equation may cause traditional iterative solvers to stop at a spurious solution if the initial guess is not close to a physically meaningful solution. When multiple distinct solutions are of interest, a naive strategy is to try different initial guesses as many as possible so that iterative solvers can return distinct solutions as many as possible. However, most of the initial guesses would lead to either spurious solutions or repeated solutions, making this approach usually time-consuming and inefficient unless a priori estimate of the solutions is available.

Neural network-based optimization has become a powerful tool for solving nonlinear differential equations, dating back to 1980s [67] and 1990s [52, 33, 23, 51], and recently revisited in high-dimensional spaces [38, 7, 46, 12, 83, 37, 71, 48, 47, 81, 68]. As a form of nonlinear parametrization through compositions of simple functions [34], deep neural networks (DNNs) can efficiently approximate various useful classes of functions without or lessening the curse of dimensionality [6, 60, 62, 70, 80, 57, 61, 44] and achieve exponential approximation rates [79, 62, 57, 54, 27, 66]. Therefore, applying DNNs to parametrize the solution space of differential equations (including

boundary value problems, initial value problems, and eigenvalue problems) and seeking a solution via energy minimization from variational formulation have become a popular choice, e.g., the least-square method [7, 46, 43] as a special case of variational formulation, the Ritz method [28], the Nitsche method [56].

However, neural network-based optimization usually can only find the smoothest solution with the fastest decay in the frequency domain due to the implicit regularization of network structures and the stochastic gradient descent (SGD) for solving the minimization problem, no matter how the initial guess is randomly selected. It was shown through the frequency principle of neural networks [84, 85, 58] and the neural tangent kernel [13] that neural networks have an implicit bias towards functions that decay fast in the Fourier domain and the gradient descent method tends to fit a low-frequency function better than a high-frequency function. Through the analysis of the optimization energy landscape of SGD, it was shown that SGD with small batches tends to converge to the flattest minimum [63, 53, 20]. Though the above optimization and generalization analysis work only for regression problems, they can be generalized to PDE problems. Recently in [59], the optimization convergence and generalization analysis of two-layer neural networks for general second-order linear PDEs with variable coefficients on a bounded domain in an arbitrary dimension has been investigated. Global convergence of the gradient descent optimization in the over-parametrization regime is proved using neural tangent kernels and the generalization error with a regularized loss using a Barron-norm is analyzed. Later in [73], the neural tangent kernel of network-based PDE solvers using two-layer neural networks for one-dimensional Poisson equation is also discussed including the analog of the spectral bias for regression problems proposed in [13]. Therefore, designing an efficient algorithm for neural network-based optimization to find distinct solutions as many as possible is a challenging problem.

To tackle the challenging problem just above and find distinct solutions as many as possible, we propose a network-based structure probing deflation method in this paper. The key idea of the deflation method is to introduce deflation operators built with known solutions to regularize deep learning optimization, making known solutions no longer local minimizers of the optimization energy landscape while preserving unknown solutions as local minimizers. In particular, we introduce a deflation functional mapping known solutions to infinity. We multiply this deflation functional to the original optimization loss function, then the known solutions will be removed from consideration and unknown solutions can be found by optimizing the regularized loss function via SGD. Furthermore, to facilitate the convergence of SGD, we propose special network structures incorporating boundary conditions of differential equations to simplify the optimization loss function. Finally, a novel structure-probing algorithm is proposed to initialize the deflation optimization making it more powerful to identify distinct solutions with desired structures.

As a general framework, the deflation method can be applied to all neural network-based optimization methods for differential equations. In this paper, we will take the example of boundary value problem (BVP) and the least-square method without loss of generality. The generalization to other problems and methods is similar. Consider the boundary value problem (BVP)

$$(1.1) \quad \begin{aligned} \mathcal{D}u(\mathbf{x}) &= f(u(\mathbf{x}), \mathbf{x}), \text{ in } \Omega, \\ \mathcal{B}u(\mathbf{x}) &= g(\mathbf{x}), \text{ on } \partial\Omega, \end{aligned}$$

where $\mathcal{D} : \Omega \rightarrow \Omega$ is a differential operator that is either linear or nonlinear, $f(u(\mathbf{x}), \mathbf{x})$ can be a nonlinear function in u , Ω is a bounded domain in $\mathbb{R}^d$, and $\mathcal{B}u = g$ characterizes the boundary condition. Other types of problems like initial value problems can also be formulated as a BVP as discussed in [37]. Then least-square method seeks a solution $u(\mathbf{x}; \boldsymbol{\theta})$ as a neural network with a

parameter set θ via the following optimization problem

$$(1.2) \quad \min_{\theta} L_{\text{LS}} := \|\mathcal{D}u(\mathbf{x}; \theta) - f(u, \mathbf{x})\|_{L^2(\Omega)}^2 + \lambda \|\mathcal{B}u(\mathbf{x}; \theta) - g(\mathbf{x})\|_{L^2(\partial\Omega)}^2,$$

where L_{LS} is the loss function measuring the L^2 norms of the differential equation residual $\mathcal{D}u(\mathbf{x}; \theta) - f(u, \mathbf{x})$ and the boundary residual $\mathcal{B}u(\mathbf{x}; \theta) - g(\mathbf{x})$, and $\lambda > 0$ is a regularization parameter.

As we shall see, the neural network deflation method enjoys four main advantages compared to traditional deflation methods not based on deep learning:

- Numerical examples show that the network-based method can identify more solutions than other existing methods, e.g., see Test Case 5 in Section 5.
- The network-based method can be applied to solve high-dimensional nonlinear differential equations with multiple solutions while existing methods are only applicable to low-dimensional problems. For example, there is a 6-dimensional Yamabe's equation in Test Case 6 in Section 5.
- The network-based method can be applied to problems with complex domains due to the flexibility of neural network parameterization, e.g., see Test Cases 5 & 6 in Section 5.
- As we shall discuss in Section 3.4, the network-based method admits lower computational complexity in each iteration compared to existing methods like the original deflation method in [29].

1.2. Related work. The deflation technique is traced back to the last century for identifying distinct roots of scalar polynomials [75]. This technique was extended to find roots of systems of nonlinear algebraic equations by Brown and Gearhart in [9], where deflation matrices were constructed with old roots to transform the residual of a system of nonlinear algebraic equations so that iterative methods applied to the new residual will only converge to a new root. In [29], Ferrell et al. extended the theoretical framework of Brown and Gearhart [9] to the case of infinite-dimensional Banach spaces with new classes of deflation operators, enabling the Newton-Kantorovitch iteration to converge to several distinct solutions of nonlinear differential equations even with the same initial guess.

Another well-established method for distinct solutions of differential equations is based on the numerical continuation [5, 4, 14, 16], where the basic idea of which is to transform the known solutions of a simple start system gradually to the desired solutions of a difficult target system. For example, [1] proposed coefficient-parameter polynomial continuation for computing all geometrically isolated solutions to polynomial systems. [39] put forward a bootstrapping approach for computing multiple solutions of differential equations using a homotopy continuation method with domain decomposition to speed up computation. For more examples of homotopy-based methods and theory in the literature, the reader is referred to [55].

The third kind of methods to identify distinct solutions of nonlinear systems is the numerical integration of the Davidenko differential equation associated with the original nonlinear problem [8, 21]. The basic idea is to introduce an artificial time parameter s such that solving the original nonlinear equation $F(u(\mathbf{x})) = 0$ to identify a solution $u_0(\mathbf{x})$ is equivalent to finding a steady state solution of a time-dependent nonlinear equation $\frac{dF(u(s, \mathbf{x}))}{ds} + F(u(s, \mathbf{x})) = 0$, which provides a gradient flow of $u(s, \mathbf{x})$. The gradient flow forms an ordinary differential equation with a solution converging to a solution to the original problem, i.e., $\lim_{s \rightarrow \infty} u(s, \mathbf{x}) = u_0(\mathbf{x})$. This method is indeed a broad framework containing the Newton's method as a special example.

1.3. Organization. This paper is organized as follows. In Section 2, we will review the fully connected feed-forward neural network, introduce the formulation of the least-square method for BVP, and design special network structures for four types of boundary conditions. In Section 3,

the detailed formulation and implementation of the proposed method will be presented. In Section 4, the structure probing initialization is introduced. Various numerical experiments are provided in Section 5 to verify the efficiency of the proposed method. Finally, we conclude this paper in Section 6.

2. Network-based Methods for Differential Equations. In this section, we introduce the network-based least-square method based on fully connected feed-forward neural networks and (1.2) for solving the BVP (1.1). Moreover, special network structures for common boundary conditions are introduced to simplify the loss function in (1.2) to facilitate the convergence to the desired PDE solution. Vectors are written in bold font to distinguish from scalars in our presentation.

2.1. Fully connected feed-forward neural network (FNN). FNNs are one of the most popular DNNs and are widely applied to network-based methods for differential equations. Mathematically speaking, for a fixed nonlinear activation function σ , FNN is the composition of L simple nonlinear functions, called hidden layer functions, in the following formulation:

$$\phi(\mathbf{x}; \boldsymbol{\theta}) := \mathbf{a}^T \mathbf{h}_L \circ \mathbf{h}_{L-1} \circ \cdots \circ \mathbf{h}_1(\mathbf{x}) \quad \text{for } \mathbf{x} \in \mathbb{R}^d,$$

where $\mathbf{a} \in \mathbb{R}^{N_L}$; $\mathbf{h}_\ell(\mathbf{x}_\ell) := \sigma(\mathbf{W}_\ell \mathbf{x}_\ell + \mathbf{b}_\ell)$ with $\mathbf{W}_\ell \in \mathbb{R}^{N_\ell \times N_{\ell-1}}$ and $\mathbf{b}_\ell \in \mathbb{R}^{N_\ell}$ for $\ell = 1, \dots, L$. With the abuse of notations, $\sigma(\mathbf{x})$ means that σ is applied entry-wise to a vector $\mathbf{x}$ to obtain another vector of the same size. Usual choices of σ include the rectified linear unit (ReLU) function $\sigma(x) = \max\{x, 0\}$, its cubic polynomial $\sigma(x) = \max\{x^3, 0\}$, a hyperbolic tangent function $\sigma(x) = \tanh(x)$, etc. N_ℓ is the width of the ℓ -th layer and L is the depth of the FNN. $\boldsymbol{\theta} := \{\mathbf{a}, \mathbf{W}_\ell, \mathbf{b}_\ell : 1 \leq \ell \leq L\}$ is the set of all parameters in ϕ to determine the underlying neural network. Other kinds of neural networks are also suitable in our proposed methods, but we will adopt FNNs for simplicity.

2.2. Least-square method. The least-square method is an optimization approach to solve general differential equations. Specifically, let $u(\mathbf{x}; \boldsymbol{\theta})$ be a neural network to approximate the solution $u(\mathbf{x})$ of BVP (1.1), then the least-square method is formulated as

$$(2.1) \quad \min_{\boldsymbol{\theta}} L_{\text{LS}}(\boldsymbol{\theta}) := \|\mathcal{D}u(\mathbf{x}; \boldsymbol{\theta}) - f(\mathbf{x})\|_{L^2(\Omega)}^2 + \lambda \|\mathcal{B}u(\mathbf{x}; \boldsymbol{\theta}) - g(\mathbf{x})\|_{L^2(\partial\Omega)}^2,$$

where L_{LS} is the loss function measuring the weighted magnitude of the differential equation residual $\mathcal{D}u(\mathbf{x}; \boldsymbol{\theta}) - f(\mathbf{x})$ and the boundary residual $\mathcal{B}u(\mathbf{x}; \boldsymbol{\theta}) - g(\mathbf{x})$ in the sense of L^2 -norm with a weight parameter $\lambda > 0$.

The goal of (2.1) is to find an appropriate set of parameters $\boldsymbol{\theta}$ such that the network $u(\mathbf{x}; \boldsymbol{\theta})$ minimizes the loss L_{LS} . If the loss L_{LS} is minimized to zero with some $\boldsymbol{\theta}$, then $u(\mathbf{x}; \boldsymbol{\theta})$ satisfies $\mathcal{D}u(\mathbf{x}; \boldsymbol{\theta}) - f(\mathbf{x}) = 0$ in Ω and $\mathcal{B}u(\mathbf{x}; \boldsymbol{\theta}) - g(\mathbf{x}) = 0$ on $\partial\Omega$, implying that $u(\mathbf{x}; \boldsymbol{\theta})$ is exactly a solution of (1.1). If L_{LS} is minimized to a nonzero but small positive number, $u(\mathbf{x}; \boldsymbol{\theta})$ is close to the true solution as long as (1.1) is well-posed (e.g. the elliptic PDE with Neumann boundary condition, see Theorem 4.1 in [37]).

In general, the optimization problem (2.1) is solved by stochastic gradient descent (SGD) method or its variants (e.g. Adagrad [25], Adam [50] and AMSGrad [72]) in the deep-learning framework. The optimization and mesh-free setting of the least-square method with neural networks admit several advantageous features that led to its great success and popularity including but not limited to 1) the capacity to solve high-dimensional problems; 2) the flexibility to solve equations of various forms on complicated problem domains; 3) the simple and high-performance implementation with automatic differential programming in existing open-source software.

2.3. Special network structures for boundary conditions. In numerical implementations, the least-square loss function in (2.1) relies on the selection of a suitable weight parameter λ and a suitable initial guess. If λ is not appropriate, it may be difficult to identify a reasonably good minimizer of (2.1). For instance, in the BVP (1.1) with $g \equiv 0$, if we solve (2.1) by SGD with an initial guess θ^0 such that $u(\mathbf{x}; \theta^0) \approx 0$, SGD might converge to a local minimizer corresponding to a solution neural network close to a constant zero, which is far away from the desired solution, especially when the differential operator $\mathcal{D}$ is highly nonlinear or λ is too large. The undesired local minimizer is due to the fact that the boundary residual $\|\mathcal{B}u(\mathbf{x}; \theta) - g(\mathbf{x})\|$ overwhelms the equation residual $\|\mathcal{D}u(\mathbf{x}; \theta) - f(\mathbf{x})\|$ in the loss function.

The idea just above motivates us to design special networks $u(\mathbf{x}; \theta)$ that satisfy the boundary condition $\mathcal{B}u(\mathbf{x}; \theta) = g(\mathbf{x})$ automatically and hence we can simplify the least-square loss function from (2.1) to

$$(2.2) \quad \min_{\theta} L_{\text{LS}}(\theta) := \|\mathcal{D}u(\mathbf{x}; \theta) - f(\mathbf{x})\|_{L^2(\Omega)}^2.$$

As we shall see in the numerical section, our numerical tests show that such simplification can help SGD to converge to desired solutions rather than spurious solutions. The design of these special neural networks depends on the type of boundary conditions. We will discuss four common types of boundary conditions by taking one-dimensional problems defined in the domain $\Omega = [a, b]$ as an example. Network structures for more complicated boundary conditions in high-dimensional domains can be constructed similarly. In what follows, denote by $\hat{u}(x; \theta)$ a generic deep neural network with trainable parameters θ . We will augment $\hat{u}(x; \theta)$ with several specially designed functions to obtain a final network $u(x; \theta)$ that satisfies $\mathcal{B}u(x; \theta) = g(x)$ automatically.

Case 1. Dirichlet boundary condition $u(a) = a_0$, $u(b) = b_0$.

In this case, we can introduce two special functions $h_1(x)$ and $l_1(x)$ to augment $\hat{u}(x; \theta)$ to obtain the final network $u(x; \theta)$:

$$(2.3) \quad u(x; \theta) = h_1(x)\hat{u}(x; \theta) + l_1(x).$$

Note $h_1(x)$ and $l_1(x)$ are chosen such that $u(x; \theta)$ automatically satisfies the Dirichlet $u(a; \theta) = a_0$, $u(b; \theta) = b_0$ no matter what θ is. Then $u(x; \theta)$ is used to approximate the true solution of the differential equation and is trained through (2.2).

For the purpose, $l_1(x)$ is set as a lifting function which satisfies the given Dirichlet boundary condition, i.e. $l_1(a) = a_0$, $l_1(b) = b_0$; $h_1(x)$ is set as a special function which satisfies the homogeneous Dirichlet boundary condition, i.e. $h_1(a) = 0$, $h_1(b) = 0$. A straightforward choice of $l_1(x)$ is the linear function given by

$$l_1(x) = (b_0 - a_0)(x - a)/(b - a) + a_0.$$

For $h_1(x)$, we can set it as a (possibly fractional) polynomial with roots a and b , namely,

$$h_1(x) = (x - a)^{p_a}(x - b)^{p_b},$$

with $0 < p_a, p_b \leq 1$. To obtain an accurate approximation, p_a and p_b should be chosen to be consistent with the orders of a and b of the true solution, hence no singularity will be brought into the network structure. For regular solutions, we take $p_a = p_b = 1$; for singular solutions, p_a and p_b would take fractional values. For instance, in the case of a fractional Laplace equation $(-\Delta)^s u = f$ for $0 < s < 1$ on the domain $\Omega = [-1, 1]$ with boundary conditions $u(\pm 1) = 0$, the true solution $u(x)$ has the property that $u(x) = (x - 1)^s(x + 1)^s v(x)$ with $v(x)$ as a smooth function [2, 26]. Then in the construction of $u(x; \theta)$, it is reasonable to choose $h_1(x) = (x - 1)^s(x + 1)^s$ and $l_1(x) = 0$.

Case 2. one-sided condition $u(a) = a_0, u'(a) = a_1$.

Similarly to Case 1, the special network is constructed by $u(x; \boldsymbol{\theta}) = h_2(x)\hat{u}(x; \boldsymbol{\theta}) + l_2(x)$, where the lifting function $l_2(x)$ is given by

$$l_2(x) = a_1(x - a) + a_0,$$

and $h_2(x)$ is set as

$$(2.4) \quad h_2(x) = (x - a)^{p_a},$$

with $1 < p_a \leq 2$. Such p_a guarantees $h_2(x)\hat{u}(x; \boldsymbol{\theta})$ and its first derivative both vanish at $x = a$.

Case 3. mixed boundary condition $u'(a) = a_0, u(b) = b_0$.

In this case, the special network is constructed by $u(x; \boldsymbol{\theta}) = \tilde{u}(x; \boldsymbol{\theta}) + l_3(x)$ with a lifting function $l_3(x)$ chosen as a linear function satisfying the mixed boundary conditions, e.g.,

$$l_3(x) = a_0x + b_0 - a_0b,$$

and $\tilde{u}(x; \boldsymbol{\theta})$ satisfying the homogeneous mixed boundary conditions. In the construction of $\tilde{u}(x; \boldsymbol{\theta})$, it is inappropriate to naively take $\tilde{u}(x; \boldsymbol{\theta}) = (x - a)^{p_a}(x - b)^{p_b}$ with $1 < p_a \leq 2$ and $0 < p_b \leq 1$, following the approaches in the preceding two cases, because such $\tilde{u}(x; \boldsymbol{\theta})$ satisfies a redundant condition $\tilde{u}(a; \boldsymbol{\theta}) = 0$. Instead, we assume

$$(2.5) \quad \tilde{u}(x; \boldsymbol{\theta}) = (x - a)^{p_a}\hat{u}(x; \boldsymbol{\theta}) + c,$$

where $1 < p_a \leq 2$ and c is a network-related constant to be determined. Clearly, (2.5) implies $\tilde{u}'(a; \boldsymbol{\theta}) = 0$, whereas $\tilde{u}(a; \boldsymbol{\theta})$ has not been specified. Next, the constraint $\tilde{u}(b; \boldsymbol{\theta}) = 0$ gives $c = -(b - a)^{p_a}\hat{u}(b; \boldsymbol{\theta})$. Therefore, the special network for mixed boundary conditions is constructed via

$$(2.6) \quad u(x; \boldsymbol{\theta}) = (x - a)^{p_a}\hat{u}(x; \boldsymbol{\theta}) - (b - a)^{p_a}\hat{u}(b; \boldsymbol{\theta}) + l_3(x).$$

Case 4. Neumann boundary condition $u'(a) = a_0, u'(b) = b_0$.

Similarly to Case 3, we construct the network by $u(x; \boldsymbol{\theta}) = \tilde{u}(x; \boldsymbol{\theta}) + l_4(x)$ with a lifting function $l_4(x)$ satisfying the Neumann boundary condition given by

$$l_4(x) = \frac{(b_0 - a_0)}{2(b - a)}(x - a)^2 + a_0x.$$

And $\tilde{u}(x; \boldsymbol{\theta})$ satisfying the homogeneous Neumann boundary condition is assumed to be

$$(2.7) \quad \tilde{u}(x; \boldsymbol{\theta}) = (x - a)^{p_a}\check{u}(x; \check{\boldsymbol{\theta}}) + c_1,$$

where $1 < p_a \leq 2$, $\check{u}(x; \check{\boldsymbol{\theta}})$ is an intermediate network to be determined later, and c_1 is a network parameter to be trained together with $\check{\boldsymbol{\theta}}$. It is easy to check that $\tilde{u}'(a; \boldsymbol{\theta}) = 0$. Next, by the constraint $\tilde{u}'(b; \boldsymbol{\theta}) = p_a(b - a)^{p_a-1}\check{u}(b; \check{\boldsymbol{\theta}}) + (b - a)^{p_a}\check{u}'(b; \check{\boldsymbol{\theta}}) = 0$, we have

$$p_a\check{u}(b; \check{\boldsymbol{\theta}}) + (b - a)\check{u}'(b; \check{\boldsymbol{\theta}}) = 0,$$

which can be reformulated as

$$\left(\exp\left(\frac{p_ax}{b-a}\right)\check{u}(x; \check{\boldsymbol{\theta}}) \right)'_{x=b} = 0.$$

Therefore, we have

$$(2.8) \quad \exp\left(\frac{p_a x}{b-a}\right) \check{u}(x; \check{\theta}) = (x-b)^{p_b} \hat{u}(x; \hat{\theta}) + c_2,$$

where $1 < p_b \leq 2$ and c_2 is another network parameter to be trained together with $\hat{\theta}$. Finally, by combining (2.7) and (2.8), we obtain the following special network satisfying the given Neumann condition, i.e.

$$(2.9) \quad u(x; \theta) = \exp\left(\frac{p_a x}{a-b}\right) (x-a)^{p_a} ((x-b)^{p_b} \hat{u}(x; \hat{\theta}) + c_2) + c_1 + l_4(x),$$

where $\theta = \{\hat{\theta}, c_1, c_2\}$.

Finally, we would like to remark that it is difficult to construct special neural networks to automatically satisfy boundary conditions when the PDE domain is irregular. In this case, the conventional penalty method in (2.1) is more preferable. Though we will show in our numerical experiments that special neural networks satisfying boundary conditions are better than penalty methods to identify distinct solutions. This does not exclude the possibility that penalty methods, or other advanced optimization algorithms for constrained optimization, can also work well with well-tuned parameters.

3. Neural Network Deflation. In this section, we propose the general formulation, the detailed implementation, and the computational complexity of the proposed method. As we shall see, our method is easy to implement on high-dimensional and complex domains with a lower computational cost per iteration than other traditional deflation methods.

3.1. Formulation. A nonlinear BVP (1.1) might have multiple distinct solutions and each solution is a minimizer of the corresponding network-based optimization, say

$$(3.1) \quad \min_{\theta} L(u(\mathbf{x}; \theta)),$$

where L is a generic loss function for solving differential equations. One example of L is the residual loss in (2.2), and L can also be other loss functions. Due to the implicit regularization of SGD and neural networks, only local minimizers in flat energy basins are likely to be found. Hence, no matter how to initialize the SGD and how to choose hyper-parameters, usually, only a few solutions can be found by minimizing (3.1) directly.

The neural network deflation is therefore introduced, the main idea of which is to construct a modified loss function L_{ND} with two properties: First, a candidate minimizer of L_{ND} is also a minimizer of L . Second, the minimizers that are already found by the network-based optimization (3.1) will not be minimizers of L_{ND} again. Following this idea, L_{ND} is constructed by multiplying L with a deflation term introduced in [29] such that the energy landscape of L is modified. Specifically, suppose the minimum value of L is zero. Let $u_k(\mathbf{x})$ ($k = 1, \dots, K$) be the solutions already found by (3.1), then the neural network deflation is formulated as the following optimization problem,

$$(3.2) \quad \min_{\theta} L_{\text{ND}} := \left(\sum_{k=1}^K \|u(\mathbf{x}; \theta) - u_k(\mathbf{x})\|_{L^2(\Omega)}^{-p_k} + \alpha \right) L(u(\mathbf{x}; \theta)),$$

where p_k are positive powers for $k = 1, \dots, K$, and $\alpha > 0$ is a shift constant. Here, we name $u_k(\mathbf{x})$ ($k = 1, \dots, K$) as deflation sources. Indeed, the modified loss function L_{ND} satisfies the two properties discussed above. First, any minimizer of L_{ND} such that $L_{\text{ND}} = 0$ also ensure $L = 0$ and, hence, is also a minimizer of L . Second, for all $k = 1, \dots, K$, the term $\|u(\mathbf{x}; \theta) - u_k(\mathbf{x})\|_{L^2(\Omega)}^{-p_k}$

acts a penalty term that excludes u_k as a minimizer, since it approaches infinity as u goes to u_k . The introduction of a positive α is help to eliminate spurious solutions in practice. If $\alpha = 0$, the modified loss function L_{ND} would approach zero when u is far from all u_k 's, which leads to many spurious solutions. For a more detailed discussion of the deflation term, the reader can refer to [29].

3.2. Deflation with a varying shift. The original deflation operator introduced in [29] fixes the shift α in (3.2) as a constant. In this paper, we propose a new variant of deflation operators with a varying shift α along with the SGD iteration. Note that when α is equal or close to 0, the deflation term $\sum_{k=1}^K \|u(\mathbf{x}; \boldsymbol{\theta}) - u_k(\mathbf{x})\|_{L^2(\Omega)}^{-p_k}$ dominates the loss and hence gradient descent tends to converge to what is far away from the known solutions. When α is moderately large, the original loss function $L(u(\mathbf{x}; \boldsymbol{\theta}))$ dominates the loss and the gradient descent process tends to converge to a solution with a smaller residual. Therefore, α in this paper is set to be a monotonically increasing function of the SGD iteration. In the early stage, α is chosen to be close to 0 such that the current solution will be pushed away from known solutions. During this stage, a large learning rate is preferable. In the latter stage when the current solution is roughly stable, α is set to be large and a small learning rate is used to obtain a small residual loss.

In practice, one heuristic choice is to increase α exponentially with a linearly growing power when the iteration increases. For example, in the n -th iteration, α is set as α_n defined below

$$(3.3) \quad \alpha_n = 10^{p_0 + n(p_1 - p_0)/N_I},$$

where p_0 and p_1 are two prescribed powers with $p_0 \leq p_1$, and N_I is the total number of iterations. Note that the exponentially varying formula is also widely used in setting the learning rates of SGD.

3.3. Discretization. In the implementation, the continuous loss functions in (2.2) and (3.2) are approximately evaluated by stochastic sampling. The L^2 -norm can be interpreted as an expectation of a random function with a random variable $\mathbf{x}$ in a certain domain. Hence, the expectation is approximated by sampling $\mathbf{x}$ several times and computing the average function value as an approximant. Let us take $\|u(\mathbf{x})\|_{L^2(\Omega)}$ as an example. We generate N_p random samples $\mathbf{x}_i$, $i = 1, \dots, N_p$, which are uniformly distributed in Ω . Denote $\mathbf{X} := \{\mathbf{x}_i\}_{i=1}^{N_p}$, then $\|u(\mathbf{x})\|_{L^2(\Omega)}$ is evaluated as the discrete L^2 -norm denoted as $\|u(\mathbf{x})\|_{L^2(\mathbf{X})}$ via

$$(3.4) \quad \|u(\mathbf{x})\|_{L^2(\mathbf{X})} := \left(\frac{1}{N_p} \sum_{\mathbf{x}_i \in \mathbf{X}} |u(\mathbf{x}_i)|^2 \right)^{\frac{1}{2}}.$$

The discretization technique above is applied to discretize the L^2 -norms in all loss functions in this paper. In the n -th iteration of gradient descent, assuming that the shift α is set to be α_n and the set of random samples is denoted as $\mathbf{X}_n$, the discrete deflation loss function is calculated by

$$\hat{L}_{\text{ND}}^{(n)}(\boldsymbol{\theta}) := \left(\sum_{k=1}^K \|u(\mathbf{x}; \boldsymbol{\theta}) - u_k(\mathbf{x})\|_{L^2(\mathbf{X}_n)}^{-p_k} + \alpha_n \right) \hat{L}(u(\mathbf{x}; \boldsymbol{\theta})),$$

where $\hat{L}(u(\mathbf{x}; \boldsymbol{\theta}))$ is a discrete approximation to $L(u(\mathbf{x}; \boldsymbol{\theta}))$ using the same set of samples, e.g.,

$$\hat{L}(u(\mathbf{x}; \boldsymbol{\theta})) = \|\mathcal{D}u(\mathbf{x}; \boldsymbol{\theta}) - f(\mathbf{x})\|_{L^2(\mathbf{X}_n)}^2$$

when the least-square loss in (2.2) is applied. Then the network parameter $\boldsymbol{\theta}$ is updated by

$$\boldsymbol{\theta} \leftarrow \boldsymbol{\theta} - \tau_n \nabla_{\boldsymbol{\theta}} \hat{L}_{\text{ND}}^{(n)}(\boldsymbol{\theta}),$$

where $\tau_n > 0$ is the learning rate in the n -th iteration. In our implementation, $\mathbf{X}_n$ is renewed in every iteration. Note that the gradient of the loss function can be evaluated using PyTorch built-in function autograd that essentially compute the gradient using a sequence of chain rules, since the network is the composition of several simple functions with explicit formulas.

3.4. Computational Complexity. Let us estimate the computational complexity of the SGD algorithm for deflation optimization (3.2) with least-square loss function (2.2). Recall that N_p denotes the number of random samples in each iteration. Assume that the FNN has L layers and N neurons in each hidden layer. Note that evaluating the FNN or computing its derivative with respect to its parameters or input $\mathbf{x}$ via the forward or backward propagation takes $O(dN + LN^2)$ FLOPS (floating point operations per second) for each sample $\mathbf{x}$. Moreover, as in most existing approaches, we assume $f(\mathbf{x})$ in the BVP can be evaluated with $O(d)$ FLOPS for a single $\mathbf{x}$. Therefore, $L(u(\mathbf{x}; \boldsymbol{\theta}))$ in (2.2) and its derivative $\nabla_{\boldsymbol{\theta}} L(u(\mathbf{x}; \boldsymbol{\theta}))$ can be calculated with $O(N_p(dN + LN^2))$ FLOPS using the discrete L^2 -norm in (3.4), if the differential operator $\mathcal{D}$ is evaluated through finite difference approximation. Similarly, assuming the number of known solutions K is $O(1)$ and the known solutions $\{u_k(\mathbf{x})\}_{k=1}^K$ are stored as neural networks of width N and depth L , then the deflation factor and its gradient with respect to $\boldsymbol{\theta}$ can also be calculated with $O(N_p(dN + LN^2))$ FLOPS. Finally, the total complexity in each gradient descent iteration of the deflation optimization is $O(N_p(dN + LN^2))$.

In existing methods [29, 19, 3], a given nonlinear differential equation is discretized via traditional discretization techniques, e.g. FDM and FEM, resulting in a nonlinear system of algebraic equations. The solutions of the system of algebraic equations provide numerical solutions to the original nonlinear differential equation. By multiplying different deflation terms to the nonlinear system of algebraic equations, existing methods can identify distinct solutions via solving the deflated system by Newton’s iteration. The number of algebraic equations N_e derived by FDM is exactly the number of grid points; and the number of equations derived by FEM is exactly the number of trial functions in the Galerkin formulation.

Now we compare neural network deflation with existing deflation methods in [29, 19, 3] in terms of the computational complexity under the assumption that the degrees of freedom of these methods are equal, i.e., the number of grid points or trial functions in existing methods is equal to the number of parameters in the neural network deflation, which guarantees that these methods have almost the same accuracy to find a solution. Denote the degree of freedom of these methods by W . Then by the above discussion, we have $W = N_e = O(dN + LN^2)$. Therefore, the total computational complexity in each iteration is $O(N_p W)$, where N_p is usually chosen as a hyper-parameter much smaller than W . In existing methods, the Jacobian matrix in each Newton’s iteration is a low-rank matrix plus a sparse matrix of size W by W . Typically, each iteration of Newton’s method requires solving a linear system of the Jacobian matrix, which usually requires $O(W^2)$ FLOPS. If a good preconditioner exists or a sparse direct solver for inverting the Jacobian matrix exists, the operation count may be reduced. Consequently, the total complexity in each iteration of existing methods would be more expensive than the neural network deflation depending on the performance of preconditioners.

4. Structure Probing Initialization. The initialization of parameters plays a critical role in training neural networks and has a significant impact on the ultimate performance. In the training of a general FNN, network parameters are usually randomly initialized using normal distributions with zero-mean. One popular technique is the Xavier initialization [32]: for each layer ℓ , the weight matrix $\mathbf{W}_\ell \in \mathbb{R}^{N_\ell \times N_{\ell-1}}$ is chosen randomly from a normal distribution with mean 0 and variance $1/N_{l-1}$; the bias vector $\mathbf{b}_\ell$ is initialized to be zero. As a variant of Xavier initialization, the He initialization [41] takes a slightly different variance $2/(N_{l-1} + N_l)$ for W_ℓ and $2/N_{l-1}$ for

b_ℓ . In general, FNNs initialized randomly have a smooth function configuration, and hence their Fourier transform coefficients decay quickly.

The least-squares optimization problem, either for regression problems or solving linear partial differential equations, with over-parameterized FNNs and random initialization admits global convergence by gradient descent with a linear convergence rate [45, 24, 82, 15, 59]. However, the speed of convergence depends on the spectrum of the target function. The training of a randomly initialized DNN tends to first capture the low-frequency components of a target solution quickly. The high-frequency fitting error cannot be improved significantly until the low-frequency error has been eliminated, which is referred to as F-principle [78]. Related works on the learning behavior of DNNs in the frequency domain is further investigated in [58, 85, 84, 13]. In the case of nonlinear differential equations where multiple solutions exist, these theoretical works imply that deep learning-based solvers converge to solutions in the low-frequency domain unless the DNN is initialized near a solution with high-frequency components.

The discussion just above motivates us to propose the structure probing initialization that helps the training converge to multiple structured solutions. The structure probing initialization incorporates desired structures in the initialization and training of DNNs. For example, to obtain oscillatory solutions of a differential equation, we initialize the DNN with high-frequency components to make the initialization closer to the desired oscillatory solution. During the optimization process, the magnitudes of these high-frequency components will be optimized to fit the desired solution. One choice to probe an oscillatory solution is to take a linear combination of structure probing functions with various frequencies, e.g., $\{\xi_j(\mathbf{x}) = e^{i\mathbf{k}_j \cdot \mathbf{x}}, |\mathbf{k}_j| = j, j = 1, \dots, J\}$ with $\mathbf{k}_j$ randomly selected. Then the following network u_J with a set of random parameters $\boldsymbol{\theta}$ can serve as an oscillatory initial guess:

$$(4.1) \quad u_J(\mathbf{x}; \boldsymbol{\theta}_J) = u(\mathbf{x}; \boldsymbol{\theta}) + \sum_{j=1}^J c_j \xi_j(\mathbf{x}),$$

where $\boldsymbol{\theta}_J := \{\boldsymbol{\theta}, \{c_j\}_{j=1}^J\}$ is trainable after initialization. In the initialization, $\{c_j\}$ are set as random numbers or manually determined hyper-parameters with large magnitudes. The idea of adding planewaves has been applied in [10, 11] to obtain high-frequency solutions. But the goal and detailed formulations are different. Instead of planewaves, radial basis functions are also a popular structure in the solution of differential equations. In this case, we can choose $\{\xi_j(\mathbf{x}) = \sin(j\pi|\mathbf{x}|), j = 1, \dots, J\}$ for example. The idea of structure probing initialization is not limited to the above two types of structures and is application dependent.

The above paragraph has sketched out the main idea of the structure probing initialization. Now we are ready to discuss its special cases when we need to make the structure probing network u_J in (4.1) satisfy the boundary condition $\mathcal{B}u_J = g$ in the BVP (1.1), which is important for the convergence of deep learning-based solvers as discussed in Section 2.3. For this purpose, we first construct a special network $u(\mathbf{x}; \boldsymbol{\theta})$ such that $\mathcal{B}u(\mathbf{x}; \boldsymbol{\theta}) = g$ by the approaches described in Section 2.3. Next, the structured probing functions $\{\xi_j(\mathbf{x})\}$ are specifically chosen to satisfy $\mathcal{B}\xi_j(\mathbf{x}) = 0$ for each j . As an example, let us take the one-dimensional mixed boundary condition on $[a, b]$:

$$(4.2) \quad u'(a) = a_0, \quad u(b) = b_0$$

for any constants a_0 and b_0 . Then a feasible choice of $\xi_j(\mathbf{x})$ is $\xi_j(x) = \cos(\frac{(2j-1)\pi(x-a)}{2(b-a)})$. Finally, it is easy to check that $\mathcal{B}u_J(\mathbf{x}; \boldsymbol{\theta}) = g$.

5. Numerical Examples. In this section, several numerical examples are provided to show the performance of network-based structure probing deflation in solving BVP (1.1). We

choose the least-square loss function as the general loss function $L(u(\mathbf{x}; \boldsymbol{\theta}))$ in (3.2), then the neural network deflation is formulated as

$$(5.1) \quad \min_{\boldsymbol{\theta}} L_{\text{ND}}(\boldsymbol{\theta}) := \left(\sum_{k=1}^K \|u(\mathbf{x}; \boldsymbol{\theta}) - u_k(\mathbf{x})\|_{L^2(\Omega)}^{-p_k} + \alpha \right) \|\mathcal{D}u(\mathbf{x}; \boldsymbol{\theta}) - f(\mathbf{x})\|_{L^2(\Omega)}^2,$$

where $u(\mathbf{x}; \boldsymbol{\theta})$ is the neural network of the approximate solution to be determined. Remark that the optimization problem can also be formulated by other optimization-based methods instead of least squares.

To verify the effectiveness of special networks that satisfy boundary conditions automatically, we use the deflation without the special network for boundary conditions as a comparison, where the loss function of the deflation becomes

$$(5.2) \quad \min_{\boldsymbol{\theta}} L_{\text{ND}}(\boldsymbol{\theta}) := \left(\sum_{k=1}^K \|u(\mathbf{x}; \boldsymbol{\theta}) - u_k(\mathbf{x})\|_{L^2(\Omega)}^{-p_k} + \alpha \right) \cdot \left(\|\mathcal{D}u(\mathbf{x}; \boldsymbol{\theta}) - f(u, \mathbf{x})\|_{L^2(\Omega)}^2 + \lambda \|\mathcal{B}u(\mathbf{x}; \boldsymbol{\theta}) - g(\mathbf{x})\|_{L^2(\partial\Omega)}^2 \right).$$

The overall setting for all examples is summarized as follows.

- **Environment.** The experiments are performed in Python 3.7 environment. We utilize the PyTorch library for neural network implementation and CUDA 10.0 toolkit for GPU-based parallel computing. One-dimensional examples (Test Case 1-4) are implemented on a laptop and high-dimensional examples (Test Case 5-7) are implemented on a scientific workstation.
- **Optimizer.** In all examples, the optimization problems are solved by *adam* subroutine from PyTorch library with default hyper parameters. This subroutine implements the Adam algorithm in [50].
- **Learning rate.** In each example, the learning rate is set to decay exponentially with linearly decreasing powers. Specifically, the learning rate in the n -th iteration is set as

$$\tau_n = 10^{q_0 + n(q_1 - q_0)/N_I},$$

where $q_0 > q_1$ are the initial and final powers, respectively, and N_I denotes the total number of iterations.

- **Network setting.** In each example, we construct a special network that satisfies the given boundary condition as discussed in Section 2.3. The special network involves a generic FNN, denoted by $\hat{u}$. In all examples, we set the depth and width of $\hat{u}$ as fixed numbers $L = 3$ and $N = 100$. Unless specified particularly, all weights and biases of $\hat{u}$ are initialized by $\mathbf{W}_l, \mathbf{b}_l \sim U(-\sqrt{N_{l-1}}, \sqrt{N_{l-1}})$. The activation function of $\hat{u}$ is chosen as $\sigma(x) := \max(0, x^3)$.
- **Varying shifts in deflation operators.** In one-dimensional examples (Test Case 1-4), using constant shifts is sufficient to find all solutions. In high-dimensional examples (Test Case 5-7), varying shifts will help to find more distinct solutions. In these examples, we set varying shifts according to (3.3).

We also summarize the numerical examples in this section in Table 5.1 below, which could help the reader to better understand how the extensive numerical examples demonstrate the advantages of our new ideas in this paper: 1) neural network deflation (ND); 2) structure probing initialization (SP); 3) special network for boundary conditions (BC); 4) varying shifts in deflation operators (VS).

Test Case	ND	SP	BC	VS	Justified Ideas
Case 1	1/0	0	1/0	0	ND and BC
Case 2	1/0	0	1/0	0	ND and BC
Case 3	1/0	0	1/0	0	ND and BC
Case 4	1/0	1/0	1/0	0	ND, SP and BC
Case 5	1/0	1/0	1	1/0	ND, SP, and VS
Case 6	1/0	1/0	1	1/0	ND, SP, and VS
Case 7	1/0	0	1	1	ND

Table 5.1: Summary of numerical examples and goals. In this table, “1” represents an idea is used and “0” means the idea is not used. “1/0” indicates that a comparison with/without the idea is tested.

In each example, necessary parameters to obtain each solution are listed in a table right next to the example. In these tables, we use N_p , N_I , and I_r to denote the width for $\hat{u}$, the batch size, the number of iterations, and the range of learning rates (i.e. $[10^{q_1}, 10^{q_0}]$), respectively. In each iteration of the optimization, N_p random samples will be renewed. The value of the shift α for each solution found by the deflation is listed in the table as a constant for a fixed α or an interval $[10^{p_0}, 10^{p_1}]$ for a varying α .

5.1. Numerical tests in one-dimension. First of all, we will provide four numerical tests for problems in one-dimension. These numerical tests show that the proposed neural network deflation works as well as existing methods [29, 39].

Test Case 1. We consider second-order the Painlevé equation [42, 31, 65] that seeks $u(x)$ satisfying

$$(5.3) \quad \frac{d^2 u}{dx^2} = 100u^2 - 1000x, \quad \text{in } \Omega = (0, 1),$$

$$(5.4) \quad u(0) = 0, \quad u(1) = \sqrt{10}.$$

It has been shown in [40] that the Painlevé equation (5.3)-(5.4) has exactly two solutions, denoted by u_1 and u_2 , which satisfy $u'_1(0) > 0$ and $u'_2(0) < 0$, respectively.

In our experiments, we take the following special network

$$(5.5) \quad u(x; \boldsymbol{\theta}) = x(x-1)\hat{u}(x; \boldsymbol{\theta}) + \sqrt{10}x$$

that automatically satisfies the boundary conditions and use parameters in Table 5.2. The initial guess of $\boldsymbol{\theta}$ is randomly initialized as mentioned previously. The first solution u_1 is easily found by the least-square method using (2.2), and the second solution u_2 is found by deflation with u_1 as the deflation source and $p_1 = 2$. Other parameters associated with these solutions are listed in Table 5.2. Figure 5.1 visualizes the identified solutions u_1 and u_2 with the same function configurations as in [29].

To verify the effectiveness of special networks that satisfy the boundary conditions (5.4), we use the deflation without special networks for boundary conditions as a comparison. Hence, the loss function is given by (5.2) with a solution network $u(\mathbf{x}; \boldsymbol{\theta})$ as a generic FNN of the same structure as $\hat{u}$ in (5.5). To show that the results of (5.2) are quite independent of the weight λ , $\lambda = 1$ and $\lambda = 100$ are used and the corresponding solutions are denoted as $\bar{u}_1$ and $\bar{u}_2$, respectively. As listed in Table 5.2, other parameters to identify these two

	u_1	u_2	$\bar{u}_1$	$\bar{u}_2$
N_I	10000	10000	10000	10000
N_p	1000	1000	1000	1000
I_{lr}	$[10^{-3}, 10^{-2}]$	$[10^{-3}, 10^{-2}]$	$[10^{-3}, 10^{-2}]$	$[10^{-3}, 10^{-2}]$
deflation source	/	u_1 ($p_1 = 2$)	u_1 ($p_1 = 2$)	u_1 ($p_1 = 2$)
α	/	1	1	1

Table 5.2: *Parameters for 1-D Painlevé equations (5.3)-(5.4). “/” means the corresponding item is not used (the same as below)*

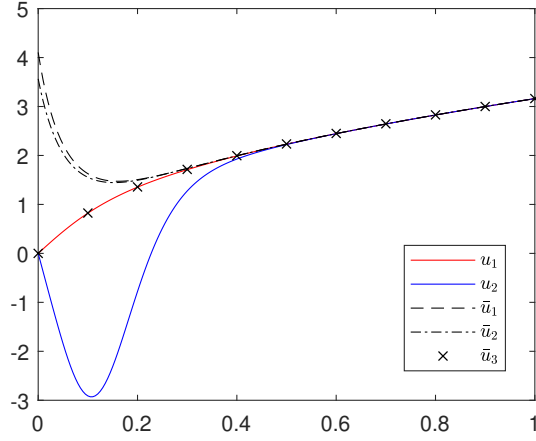


Fig. 5.1: *Identified solutions of the 1-D Painlevé equations (5.3)-(5.4) by the least squares method and neural network deflation. All correct solutions, u_1 and u_2 , are identified with special networks for boundary conditions. Spurious solutions, $\bar{u}_1$ and $\bar{u}_2$, are found if the special networks are not used. Another solution, $\bar{u}_3$, is found when the deflation fails with an inappropriate power p_1 .*

solutions are the same as those for identifying u_2 for a fair comparison. It is clear that these two solutions do not satisfy the boundary condition at the endpoint $x = 0$ (see Figure 5.1). This verifies the importance of using special networks that satisfy the boundary conditions automatically.

Moreover, we test the effectiveness of the deflation when smaller powers of deflation sources are used. We basically repeat the same experiment as the previous one for computing u_2 in Figure 5.1. The only difference is that we use a larger power $p_1 = 2$ in the previous experiment, but now we use a smaller power $p_1 = 1$. The solution by neural network deflation with $p_1 = 1$ is denoted as $\bar{u}_3$ and visualized in Figure 5.1. Note that $\bar{u}_3$ is almost the same as the deflation source u_1 by visual inspection. This result is not surprising even if we have used u_1 as the deflation source. The loss function L_{ND} in (3.2) can still be very small at $u = \bar{u}_3$ even if the deflation term is large, since the loss function L in (3.2) can be much smaller than one over the deflation term at $u = \bar{u}_3$ close to u_1 . This example indicates that an appropriate power p_1 is necessary to exclude spurious solutions close to u_1 .

	u_1	u_2	$\bar{u}_1$	$\bar{u}_2$
N_I	5000	5000	5000	5000
N_p	1000	1000	1000	1000
I_{lr}	$[10^{-3}, 10^{-2}]$	$[10^{-2}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$
deflation source	/	u_1 ($p_1 = 1$)	u_1 ($p_1 = 1$)	u_1 ($p_1 = 1$)
α	/	1	1	1

Table 5.3: *Parameters for the equation (5.6)-(5.7).*

Test Case 2. We consider a fourth-order nonlinear BVP that seeks u such that

$$(5.6) \quad \frac{d^4 u}{dx^4} = \beta x(1 + u^2) \quad \text{in } \Omega = (0, 1),$$

$$(5.7) \quad u(0) = u'(1) = u''(1) = 0, \quad u''(0) - u''(\gamma) = 0,$$

where β and γ are two given constants. Graef et al. [36, 35] have proven that the problem (5.6)-(5.7) has at least two positive solutions when $\beta = 10$ and $\gamma = 1/5$.

The three-point boundary condition (5.7) is more complicated than usual. Accordingly, we construct the following special network for it,

$$(5.8) \quad u(x; \boldsymbol{\theta}) = (x - 1)^3 \hat{u}(x; \boldsymbol{\theta}) + \hat{u}(0; \boldsymbol{\theta}) + c_\gamma x(x - 1)^3,$$

where

$$(5.9) \quad c_\gamma = \frac{1}{-12\gamma^2 + 18\gamma} \left(\frac{d^2}{dx^2} ((x - 1)^3 \hat{u}(x; \boldsymbol{\theta}))|_{x=\gamma} - \frac{d^2}{dx^2} ((x - 1)^3 \hat{u}(x; \boldsymbol{\theta}))|_{x=0} \right).$$

It can be verified that (5.8) indeed satisfies the boundary condition (5.7) independent of $\boldsymbol{\theta}$.

In our experiment, we find the first solution, denoted by u_1 , by applying the least-square method (2.2). With deflation source u_1 ($p_1 = 1$), we find the second solution, denoted by u_2 , by using the deflation (5.1). The parameters and solutions are demonstrated in Table 5.3 and Figure 5.2.

Similarly, we test the deflation without special networks for boundary conditions as a comparison under the same setting as Test Case 1. We find two solutions, denoted by $\bar{u}_1$ and $\bar{u}_2$, from $\lambda = 1$ and $\lambda = 100$, respectively (see Figure 5.2). It is clear that both solutions are spurious since their configurations do not take the prescribed boundary value 0 at $x = 0$ (see Figure 5.2), which implies the effectiveness of using special networks for boundary conditions.

Test Case 3. We consider the fourth-order nonlinear equation describing the steady laminar flow of a viscous incompressible fluid in a porous channel [77]. For simplicity, we consider the one-dimensional problem that seeks u such that

$$(5.10) \quad \frac{d^4 u}{dx^4} + \gamma \left(x \frac{d^3 u}{dx^3} + 3 \frac{d^2 u}{dx^2} \right) + R \left(u \frac{d^3 u}{dx^3} - \frac{du}{dx} \frac{d^2 u}{dx^2} \right) = 0, \quad 0 < x < 1,$$

$$(5.11) \quad u(0) = 0, \quad u''(0) = 0, \quad u(1) = 1, \quad u'(1) = 0,$$

where R is the cross-flow Reynolds number and γ is a physical constant related to the wall expansion ratio. Xu et al. [77] have proven that the problem (5.10)-(5.11) admits multiple

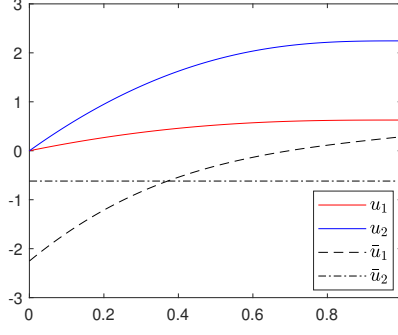


Fig. 5.2: Identified solutions of the equation (5.6)-(5.7) by least square or neural network deflation. All correct solutions, u_1 and u_2 , are identified with special networks for boundary conditions. Spurious solutions, $\bar{u}_1$ and $\bar{u}_2$, are found if the special networks are not used.

	u_1	u_2	u_3	$\hat{u}_1$	$\hat{u}_2$
N_I	20000	10000	20000	10000	10000
N_p	1000	1000	1000	1000	1000
I_{lr}	$[10^{-3}, 10^{-2}]$	$[10^{-3}, 10^{-2}]$	$[10^{-3}, 10^{-2}]$	$[10^{-3}, 10^{-2}]$	$[10^{-3}, 10^{-2}]$
deflation source	/	u_1 ($p_1 = 2$)	u_1, u_2 ($p_1 = p_2 = 2$)	u_1 ($p_1 = 2$)	u_1 ($p_1 = 2$)
α	/	1	1	1	1

Table 5.4: Parameters for the channel flows equation (5.10)-(5.11).

solutions by analytic approaches. Three solutions were found by homotopy analysis method (HAM) in [55] for the setting $R = -11$ and $\gamma = 1.5$.

In our experiments, we take the same R and γ as in [55]. The special network for the boundary condition (5.11) is chosen as

$$(5.12) \quad u(x; \boldsymbol{\theta}, \hat{c}) = x(x-1)^2(x^2\hat{u}(x; \boldsymbol{\theta}) + c)e^{2x} + \sin(\pi x/2),$$

where c is a network parameter to be trained together with $\boldsymbol{\theta}$. In this case, we initialize $\mathbf{b}_3 = \mathbf{0}$ and $c \sim U[-5, 0]$. Other network parameters are initialized as mentioned above. Firstly, one solution u_1 is found by the least-square method (2.2). Next, the second solution u_2 is obtained by the deflation (5.1) with deflation source u_1 ($p_1 = 2$). Moreover, the third solution u_3 is obtained by the deflation (5.1) with deflation sources u_1 and u_2 ($p_1 = p_2 = 2$). Corresponding parameters are shown in Table 5.4. The three found solutions and their first derivatives are plotted in Figure 5.3, which are the same solutions found in [55].

Also, a comparison test is performed to seek u_2 by the deflation (5.2) with the same setting as above, except for using a generic solution network without special structures for boundary conditions. We find two solutions, denoted by $\hat{u}_1$ and $\hat{u}_2$, using $\lambda = 1$ and $\lambda = 100$. Neither of them takes the prescribed boundary value 0 at $x = 0$ or 1 at $x = 1$ and, hence, they are spurious solutions (see Figure 5.3).

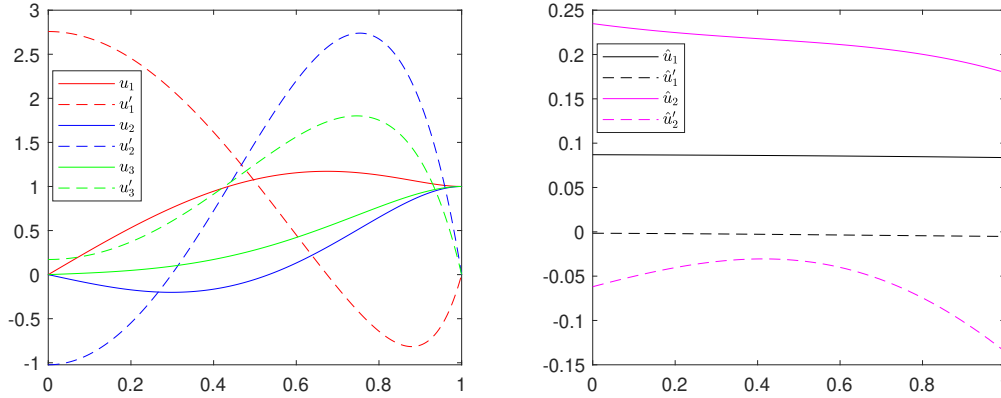


Fig. 5.3: Identified solutions and their derivatives of the channel flows equation (5.10)-(5.11) by least square or neural network deflation. All correct solutions, u_1 , u_2 and u_3 , are identified with special networks for boundary conditions. Spurious solutions, $\bar{u}_1$ and $\bar{u}_2$, are found if the special networks are not used.

Test Case 4. We consider the following second-order problem that seeks u such that

$$(5.13) \quad \frac{d^2 u}{dx^2} = f(u), \quad 0 < x < 1,$$

$$(5.14) \quad u'(0) = 0, \quad u(1) = 0,$$

where $f(u)$ is a polynomial function of u . The existence of multiple solutions for the problem (5.13) has been studied by the bootstrapping method [39].

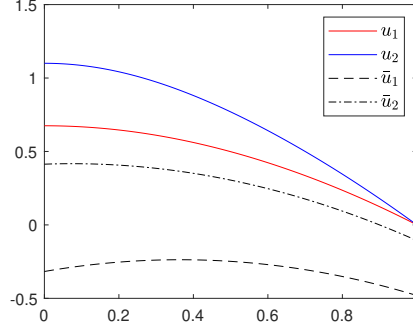
First, we set the right-hand side of the problem (5.13) as $f(u) = \lambda(1 + u^4)$. It is shown in [39] that there are two solutions for $0 < \lambda < \lambda^* = 1.30107$. In our experiments, we take $\lambda = 1.2$. The special network for the boundary condition (5.14) is given by

$$(5.15) \quad u(x; \boldsymbol{\theta}) = x^2 \hat{u}(x; \boldsymbol{\theta}) - \hat{u}(1; \boldsymbol{\theta}).$$

The first solution u_1 is found by the least-square method (2.2) and the second solution u_2 is found by the deflation (5.1) with deflation source u_1 ($p_1 = 2$). Similarly to preceding cases, we perform a comparison test without the special network structure for boundary conditions and two spurious solutions $\hat{u}_1$ (for $\lambda = 1$) and $\hat{u}_2$ (for $\lambda = 100$) are found by the deflation (5.2). The parameters for all these solutions are shown in Table 5.5 and all solutions are plotted in Figure 5.4.

Second, we repeat the test by choosing $f(u) = -\frac{\pi^2}{4}u^2(u^2 - 10)$. [39] has proved that there exist eight solutions in total. Note that $u_0 = 0$ is a trivial solution. In this case, we start from the deflation (5.1) with the special network (5.15) and the deflation source u_0 ($p_0 = 2$) to find the first solution u_1 , which is quite close to u_0 . We would like to emphasize that it is sufficient to use the deflation without structure probing initializations to identify u_0 and u_1 . However, we were not able to identify any other solutions without the structure probing initialization even if we tried our best to tune parameters and use different random initialization. To perform a wider search for other solutions, we employ the following

	u_1	u_2	$\hat{u}_1$	$\hat{u}_2$
N_I	10000	10000	10000	10000
N_p	1000	1000	1000	1000
I_{lr}	$[10^{-3}, 10^{-2}]$	$[10^{-3}, 10^{-2}]$	$[10^{-3}, 10^{-2}]$	$[10^{-3}, 10^{-2}]$
deflation source	/	u_1 ($p_1 = 2$)	u_1 ($p_1 = 2$)	u_1 ($p_1 = 2$)
α	/	1	1	1

Table 5.5: Parameters for the nonlinear problem (5.13)-(5.14) with $f(u) = \lambda(1 + u^4)$.Fig. 5.4: Identified solutions of the nonlinear problem (5.13)-(5.14) with $f(u) = \lambda(1 + u^4)$ by least square or neural network deflation. All correct solutions, u_1 and u_2 , are identified with special networks for boundary conditions. Spurious solutions, $\bar{u}_1$ and $\bar{u}_2$, are found if the special networks are not used.

structure probing initialization

$$(5.16) \quad u_J(x; \boldsymbol{\theta}, \hat{c}_j) = x^2 \hat{u}(x; \boldsymbol{\theta}) - \hat{u}(1; \boldsymbol{\theta}) + \sum_{j=1}^J \hat{c}_j \cos((2j-1)\pi x/2),$$

with initial setting $c_j = 0$ for $j = 1, \dots, J-1$ and $c_J \sim U(-5, 5)$. Two solutions, denoted by u_2 and u_3 , are found by the deflation (5.1) with source u_0 ($p_0 = 2$) and the structure probing network (5.16) with $J = 1$. Another two solutions, denoted by u_4 and u_6 , are found by the deflation (5.1) with source u_0 ($p_0 = 2$) and the network (5.16) with $J = 2$. Two more solutions, denoted by u_5 and u_7 , are found by the deflation (5.1) with deflation sources u_4 ($p_4 = 2$) and u_6 ($p_6 = 2$), respectively, and the network (5.16) with $J = 2$. Corresponding parameters, including the initial value of c_J actually randomized for each solution, are listed in Table 5.6. All the 7 nontrivial solutions are plotted in Figure 5.5.

5.2. Numerical tests in high-dimension. In this subsection, we will provide numerical tests in high-dimensional domains ($d \geq 2$).

Test Case 5. We consider 2-D Yamabe's equation that seeks u such that

$$(5.17) \quad \begin{aligned} -8\Delta u - 0.1u + \frac{u^5}{|\mathbf{x}|^3} &= 0, \quad \text{in } \Omega = \{\mathbf{x} \in \mathbb{R}^2 : r < |\mathbf{x}| < R\}, \\ u &= 1, \quad \text{on } \partial\Omega, \end{aligned}$$

	u_1	u_2	u_3	u_4
N_I	5000	5000	10000	20000
N_p	1000	1000	1000	1000
I_{lr}	$[10^{-3}, 10^{-2}]$	$[10^{-3}, 10^{-2}]$	$[10^{-4}, 10^{-3}]$	$[10^{-4}, 10^{-3}]$
J	/	1	1	2
initial c_J	/	-3.48	4.61	-3.67
deflation source	u_0 ($p_0 = 2$)	u_0 ($p_0 = 2$)	u_0 ($p_0 = 2$)	u_0 ($p_0 = 2$)
	u_5	u_6	u_7	
N_I	20000	20000	20000	
N_p	1000	1000	1000	
I_{lr}	$[10^{-4}, 10^{-3}]$	$[10^{-4}, 10^{-3}]$	$[10^{-4}, 10^{-3}]$	
J	2	2	2	
initial c_J	-4.12	3.64	3.44	
deflation source	u_4 ($p_4 = 2$)	u_0 ($p_0 = 2$)	u_6 ($p_6 = 2$)	

Table 5.6: Parameters for the nonlinear problem (5.13)-(5.14) with $f(u) = -\frac{\pi^2}{4}u^2(u^2 - 10)$.

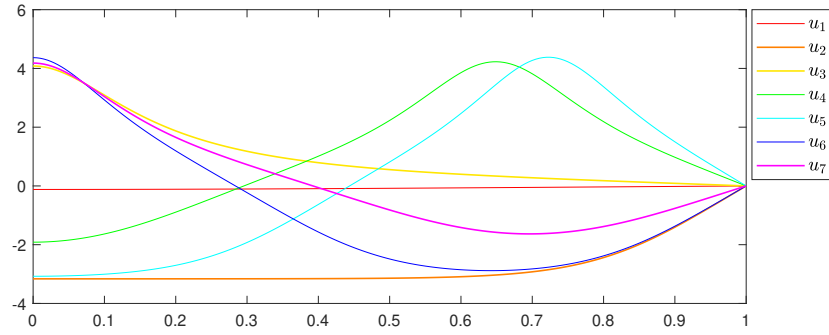


Fig. 5.5: Identified solutions of the nonlinear equation (5.13)-(5.14) with $f(u) = -\frac{\pi^2}{4}u^2(u^2 - 10)$ by the deflation.

where r and R are set as 1 and 100. Nine solutions were found by using non-network deflation techniques and various initial guesses in [29].

In our experiments, the solutions are approximated by the following special network

$$(5.18) \quad u_J(\mathbf{x}; \boldsymbol{\theta}) = \hat{u}(\mathbf{x}; \boldsymbol{\theta}) \sin\left(\pi \frac{|\mathbf{x}| - r}{R - r}\right) + 1$$

if the random initialization without the structure probing technique is used, or the following network

$$(5.19) \quad u_J(\mathbf{x}; \boldsymbol{\theta}, c_j) = \hat{u}(\mathbf{x}; \boldsymbol{\theta}) \sin\left(\pi \frac{|\mathbf{x}| - r}{R - r}\right) + \sum_{j=1}^J c_j \sin(j\pi \frac{|\mathbf{x}| - r}{R - r}) + 1$$

with the structure probing initialization, where the initial values are $c_j = 0$ for $j = 1, \dots, J-1$ and $c_J \sim U(-1, 1)$. Note that both (5.18) and (5.19) satisfy the given boundary condition automatically.

	u_1	u_2	u_3	u_4	u_5
N_I	2000	2000	2000	5000	2000
N_p	10000	10000	10000	10000	10000
I_{lr}	$[10^{-3}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$
network	(5.18)	(5.18)	(5.18)	(5.18)	(5.18)
α	/	1	1	$[0.01, 100]$	$[0.01, 100]$
deflation source	/	u_1	u_2	u_3	u_1

	u_6	u_7	u_8	u_9	u_{10}
N_I	5000	10000	20000	20000	10000
N_p	10000	10000	10000	20000	10000
I_{lr}	$[10^{-2}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$
network	(5.18)	(5.18)	(5.18)	(5.18)	(5.18)
α	$[0.01, 10]$	$[0.01, 10]$	$[0.01, 10]$	$[0.01, 10]$	$[0.01, 10]$
deflation source	u_1, u_4	u_1, u_2	u_1, u_2	u_1, u_2	u_9

	u_{11}	u_{12}	u_{13}	u_{14}	
N_I	2000	10000	10000	10000	
N_p	10000	10000	10000	10000	
I_{lr}	$[10^{-3}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$	10^{-2}	$[10^{-2}, 10^{-1}]$	
network	(5.19) ($J = 4$)	(5.19) ($J = 4$)	(5.19) ($J = 4$)	(5.18)	
α	/	0.01	$[0.01, 10]$	1	
deflation source	/	u_{11}	u_8, u_{11}	u_{11}	

Table 5.7: *Parameters for the 2-D Yamabe’s equation (5.17) ($p_k = 2$ for all deflation sources for the solutions obtained by the deflation).*

In our proposed framework of the network-based structure probing deflation with a varying shift, we always follow the four steps: 1) use the least-square method (2.2) to find the first few solutions; 2) use neural network deflation without structure probing and varying shifts to find other solutions; 3) use structure probing deflation without varying shifts to find more distinct solutions; 4) finally, use structure probing deflation with varying shifts to find extra distinct solutions. Following these procedures, we find 14 solutions in total for the 2-D Yamabe’s equation as plotted in Figure 5.6 with parameters specified in Table 5.7.

More precisely, u_1 and u_{11} are found by the least-square method (2.2) and the others are found by the deflation (5.1) with previously found solutions as deflation sources ($p_k = 2$ for all k). In deflation, we employ the technique of varying shifts in deflation operators, which helps to find more distinct solutions. All solutions are found by using networks (5.18) or (5.19) (specified in Table 5.7) with their corresponding initialization as mentioned previously, except that we take the network (5.18) with $2 - u_9$ as the initial guess to find u_{10} . We would like to remark that both the structure probing initialization and the varying shifts are key techniques to find more distinct solutions for high-dimensional problems. Without any of them, we cannot find 14 distinct solutions even if we have tried our best to tune parameters with commonly used random initialization in the literature.

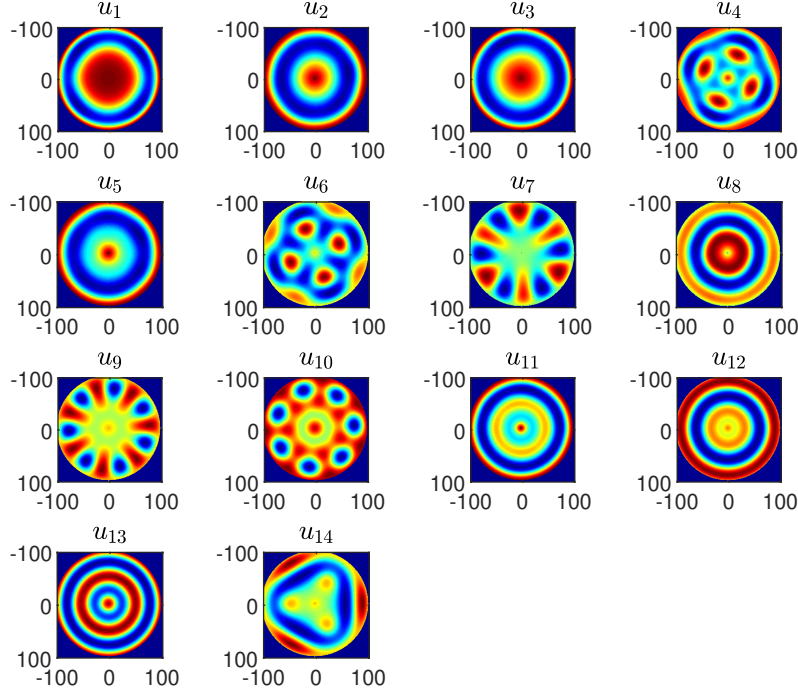


Fig. 5.6: Identified solutions of the 2-D Yamabe's equation (5.17).

Test Case 6. The high-dimensional Yamabe's equation seeks u such that

$$(5.20) \quad \begin{aligned} -\frac{4(d-1)}{(d-2)}\Delta u - 0.125u + \frac{u^{\frac{d+2}{d-2}}}{|\mathbf{x}|^3} &= 0, \quad \text{in } \Omega = \{1 < |\mathbf{x}| < 100\}, \\ u &= 1, \quad \text{on } \partial\Omega, \end{aligned}$$

where $d \geq 3$ is the dimension of the problem.

We continue applying the network (5.18) without structure probing initialization and the network (5.19) with the structure probing initialization as solution networks to solve Yamabe's equation when $d = 3$ and $d = 6$. The initialization parameters are the same as in the 2-D case.

Again, in our proposed framework of the network-based structure probing deflation with a varying shift, we follow the four steps: 1) use the least-square method (2.2) to find the first few solutions; 2) use the deflation without structure probing and varying shifts to find other solutions; 3) use structure probing deflation without varying shifts to find more distinct solutions; 4) finally, use structure probing deflation with varying shifts to find extra distinct solutions. Following these procedures, we obtain 11 solutions when $d = 3$ and 9 solutions when $d = 6$. The corresponding parameters are shown in Tables 5.8 and 5.9 for $d = 3$ and $d = 6$, respectively. The solutions are visualized in Figures 5.7 and 5.8 for $d = 3$ and $d = 6$, respectively. We would like to remark that both the structure probing initialization and the varying shifts are key techniques to find more distinct solutions for high-dimensional

	u_1	u_2	u_3	u_4	u_5
N_I	20000	20000	20000	20000	20000
N_p	10000	10000	10000	10000	10000
I_{lr}	$[10^{-2}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$
network	(5.18)	(5.18)	(5.18)	(5.18)	(5.18)
α	/	$[0.01, 10]$	1	0.1	0.01
deflation source	/	u_1	u_1, u_2	u_1, u_2	u_1, u_2
	u_6	u_7	u_8	u_9	u_{10}
N_I	20000	20000	20000	20000	20000
N_p	10000	10000	10000	10000	10000
I_{lr}	$[10^{-2}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$
network	(5.19) ($J = 4$)	(5.19) ($J = 6$)	(5.19) ($J = 4$)	(5.18)	(5.18)
α	0.01	0.1	$[0.01, 10]$	$[0.01, 10]$	$[0.01, 10]$
deflation source	u_1, u_2	u_1, u_2	u_1, u_4	u_1, u_2, u_3	u_1, u_2, u_5
	u_{11}				
N	100				
N_I	20000				
N_p	10000				
I_{lr}	$[10^{-2}, 10^{-1}]$				
network	(5.19) ($J = 4$)				
α	$[0.01, 10]$				
deflation source	u_1, u_2, u_6				

Table 5.8: *Parameters for the 3-D Yamabe's equation (5.20) ($p_k = 2$ for all deflation sources for the solutions obtained by the deflation).*

problems. Without any of them, we cannot find several distinct solutions even if we have tried our best to tune parameters with commonly used random initialization in the literature.

In these tests, the deflation powers p_k are set as 2 for all k whenever deflation is used. In the case of $d = 3$, most networks are initialized using (5.18) or (5.19), except for u_8 , u_9 and u_{10} , which are found by using initial guesses $2 - u_4$, $2 - u_3$ and $2 - u_5$, respectively. In the case of $d = 6$, we also try the initialization with a constant minus a known solution. However, this initialization method does not lead to new solutions.

Test Case 7. In the last example, we consider the following reaction-diffusion system applied in the modeling of the chemical reaction with two components [64] and irregular patterns [69],

$$(5.21) \quad \begin{cases} \mathcal{D}_1(u, v) := \varepsilon_u \Delta u - uv^2 + F(1 - u) = 0 \\ \mathcal{D}_2(u, v) := \varepsilon_v \Delta v + uv^2 - (F + k)v = 0 \end{cases} \quad \text{in } \Omega,$$

with Dirichlet boundary conditions

$$(5.22) \quad u = 1 \text{ and } v = 0 \text{ on } \partial\Omega.$$

In this case, Ω is set as a more complicated domain in $\mathbb{R}^3$ formulated by

$$(5.23) \quad \Omega = \{\mathbf{x} \in \mathbb{R}^3 : |\mathbf{x}| < \rho(\mathbf{x}) := 1 + 0.1 \sin(5\theta(x_1 + ix_2))\},$$

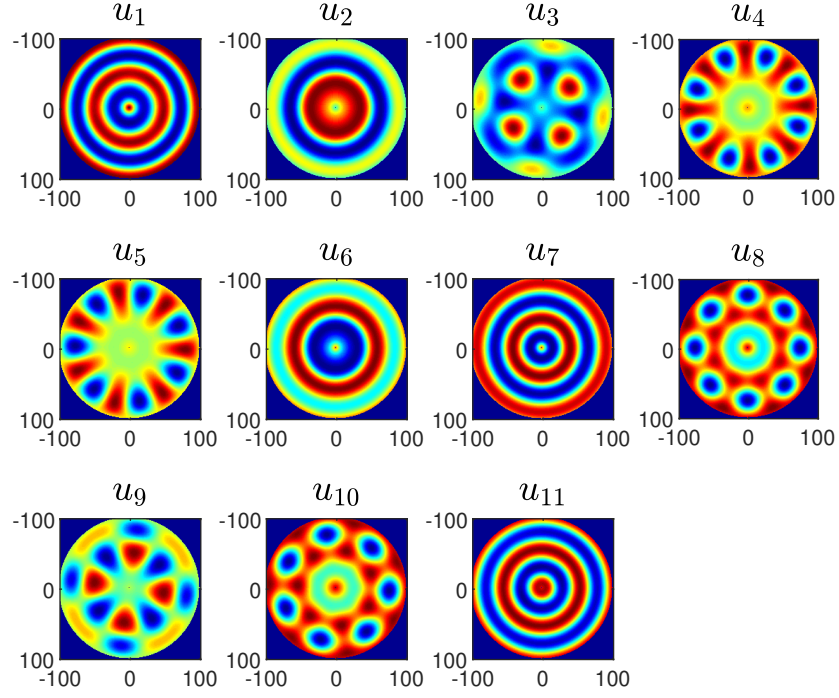


Fig. 5.7: Identified solutions of the 3-D Yamabe's equation (5.20). We visualize these solutions by projecting them in the first two coordinates.

	u_1	u_2	u_3	u_4	u_5
N_I	20000	20000	20000	20000	20000
N_P	10000	10000	10000	10000	10000
I_{lr}	$[10^{-2}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$	$[10^{-3}, 10^{-1}]$	$[10^{-3}, 10^{-1}]$
network	(5.18)	(5.18)	(5.18)	(5.18)	(5.18)
α	/	$[0.01, 10]$	0.1	10	$[0.01, 10]$
deflation source	/	u_1	u_1	u_1	u_1, u_2
	u_6	u_7	u_8	u_9	
N_I	20000	20000	20000	20000	
N_P	10000	10000	10000	10000	
I_{lr}	$[10^{-3}, 10^{-2}]$	$[10^{-3}, 10^{-2}]$	$[10^{-2}, 10^{-1}]$	$[10^{-2}, 10^{-1}]$	
network	(5.18)	(5.19) ($J = 6$)	(5.19) ($J = 6$)	(5.19) ($J = 6$)	
α	$[0.1, 1]$	/	$[0.01, 10]$	1	
deflation source	u_1, u_2	/	u_7	u_7	

Table 5.9: Parameters for the 6-D Yamabe's equation (5.20) ($p_k = 2$ for all deflation sources for the solutions obtained by the deflation).

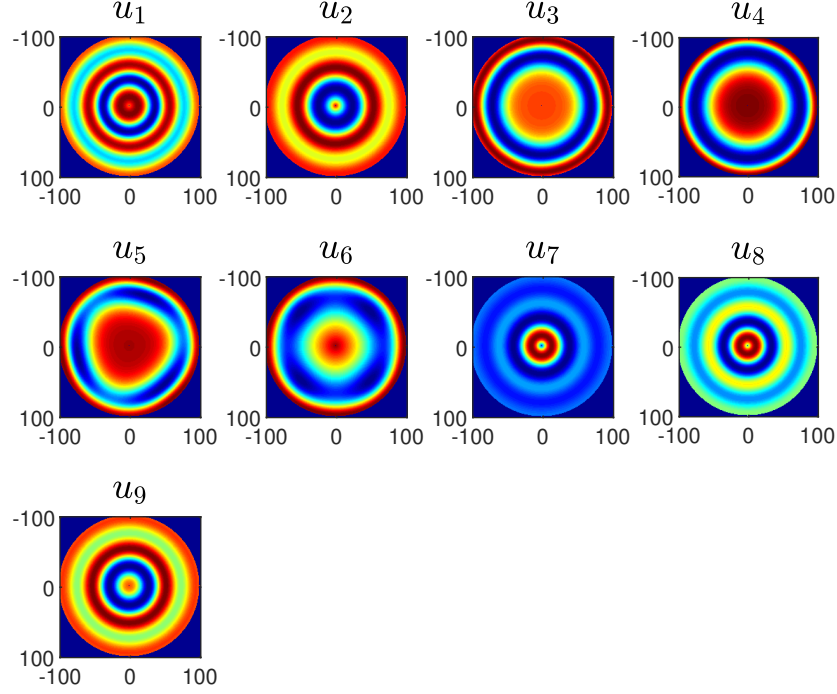


Fig. 5.8: Identified solutions of the 6-D Yamabe's equation (5.20). We visualize these solutions by projecting them in the first two coordinates.

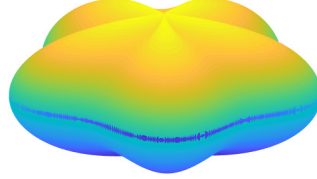


Fig. 5.9: The problem domain of the 3-D reaction-diffusion system.

where $\theta(z)$ means the argument of a complex number z . See Fig. 5.9 for the visualization of Ω . Note the system (5.21) has a pair of trivial solutions $u_0 \equiv 1$ and $v_0 \equiv 0$.

In [74], the authors solve the problem (5.21) in a 2-D square by a spectral collocation method, obtaining a vast number of solutions with residuals less than 10^{-9} . However, it is quite challenging to solve the problem (5.21) in a 3-D complicated domain by most conventional approaches (e.g., FDM and spectral methods).

Our network-based strategy is to construct two special networks $u(\mathbf{x}; \boldsymbol{\theta}_u)$ and $v(\mathbf{x}; \boldsymbol{\theta}_v)$ to approximate u and v , respectively. Specifically, we let

$$(5.24) \quad u(\mathbf{x}; \boldsymbol{\theta}_u) = \hat{u}(\mathbf{x}; \boldsymbol{\theta}_u)(|\mathbf{x}|^2 - \rho^2(\mathbf{x})) + 1,$$

$$(5.25) \quad v(\mathbf{x}; \boldsymbol{\theta}_v) = \hat{v}(\mathbf{x}; \boldsymbol{\theta}_v)(|\mathbf{x}|^2 - \rho^2(\mathbf{x})),$$

which automatically satisfy $u(\mathbf{x}; \boldsymbol{\theta}_u) = 1$ and $v(\mathbf{x}; \boldsymbol{\theta}_v) = 0$ on $\partial\Omega$. If we use the original least

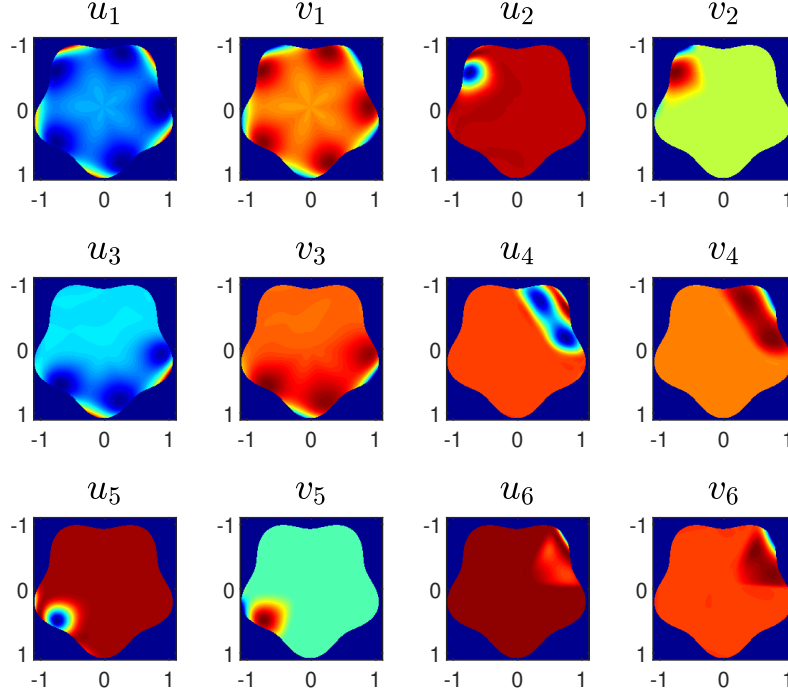


Fig. 5.10: *Selected solution pairs (u, v) of the 3-D reaction-diffusion system (5.21). We visualize these solutions by projecting them in the first two coordinates.*

squares method in (2.1), only the trivial solutions can be found. Therefore, we train the networks by the following deflation

$$(5.26) \quad \min_{\theta_u, \theta_v} L_{\text{ND}}(\theta_u, \theta_v) := \left(\sum_{k=1}^K \left(\|u(\mathbf{x}; \theta_u) - u_k(\mathbf{x})\|_{L^2(\Omega)}^{-p_k} + \|v(\mathbf{x}; \theta_v) - v_k(\mathbf{x})\|_{L^2(\Omega)}^{-p_k} \right) + \alpha \right) \cdot \left(\|\mathcal{D}_1(u(\mathbf{x}; \theta_u), v(\mathbf{x}; \theta_v))\|_{L^2(\Omega)}^2 + \|\mathcal{D}_2(u(\mathbf{x}; \theta_u), v(\mathbf{x}; \theta_v))\|_{L^2(\Omega)}^2 \right),$$

where $\{u_k(\mathbf{x}), v_k(\mathbf{x})\}_{k=1}^K$ are K pairs of solutions that have already been obtained. We start the search by taking the trivial solutions u_0 and v_0 as deflation sources, and then take identified solutions as new deflation sources for the next search. Hyper-parameters are set as $N_I = 10000$, $N_p = 10000$, and $I_r = [10^{-5}, 10^{-2}]$. Besides, we use a varying α with a range $[10^{-2}, 1]$. Deflation powers are set as $p_k = 2$ for all sources. Finally, we find more than 100 distinct solutions, some of which are shown in Fig. 5.10. The residual errors of all identified solutions for Equation (5.21) are below 1.0×10^{-3} and the corresponding values of the loss function in (5.26) are below 0.5×10^{-3} .

6. Conclusion. In this paper, we proposed the structure probing neural network deflation to find distinct solutions to nonlinear differential equations. The original optimization energy landscape of network-based methods is regularized by neural network deflation

so that known solutions are no longer local minimizers while preserving unknown solutions as local minimizers. To obtain a new solution with the desired features, a structure probing algorithm is applied to obtain an initial guess that is close to the desired solution. Finally, special network structures that satisfy various boundary conditions automatically are introduced to simplify the objective function of network-based methods. These techniques form a new framework for identifying distinct solutions of nonlinear differential equations. Compared to existing methods, the proposed neural network deflation is capable of solving high-dimensional problems on complex domains with a lower computational cost and can identify more distinct solutions. As a neural network-based PDE solver, structure probing neural network deflation may not provide highly accurate solutions. But these solutions are usually accurate enough for industrial applications and serve as a good initial guess for conventional methods as in [43] to obtain highly accurate solutions efficiently.

Structure probing neural network deflation relies on the deflation operator proposed in [29] based on conventional discretization methods. Although the application of neural networks has conquered some disadvantages of the conventional deflation method, e.g., we can solve high-dimensional problems on complex domains and identify more solutions, the proposed method in this paper still inherits some disadvantages of the conventional deflation method. For example, when two solutions are very close to each other, the optimization landscape of the deflated loss using one solution as the deflation source becomes very steep at the other solution, making it very challenging to identify another solution. As in the conventional deflation method, it is crucial to choose appropriate powers p_k for deflation sources as shown in our numerical tests. However, the parameter selection is still heuristic and problem-dependent. Learning how to choose parameters automatically is an important future direction. Network-based methods in general might need extra effort to deal with boundary conditions, which is not an issue of conventional methods. Designing more advanced optimization algorithms for constrained optimization in network-based methods would also be interesting in the future.

Acknowledgments. Y. G. was partially supported by the Ministry of Education in Singapore under the grant MOE2018-T2-2-147. C. W. is partially supported by the US National Science Foundation Award DMS-1849483. H. Y. was partially supported by the US National Science Foundation under award DMS-1945029.

REFERENCES

- [1] A. J. Sommese A. P. Morgan. Coefficient-parameter polynomial continuation. *Appl. Math. Comput.*, 29(2):123–160, 1989.
- [2] G. Acosta, J.P. Borthagaray, O. Bruno, and M. Maas. Regularity theory and high order numerical methods for the (1d)-Fractional Laplacian. *arXiv e-prints*, arXiv:1608.08443, 2016.
- [3] J. H. Adler, D. B. Emerson, P. E. Farrell, and S. P. MacLachlan. A deflation technique for detecting multiple liquid crystal equilibrium states. *arXiv e-prints*, arXiv:1601.07383, 2016.
- [4] E. L. Allgower and K. Georg. Continuation and path following. *Acta Numer.*, 2:1–64, 1993.
- [5] E. L. Allgower and K. Georg. *Introduction to Numerical Continuation Methods*. Classics Appl. Math., SIAM, Philadelphia, 2003.
- [6] A. R. Barron. Universal approximation bounds for superpositions of a sigmoidal function. *IEEE Trans. Inf. Theory*, 39(3):930–945, 1993.
- [7] J. Berg and K. Nyström. A unified deep artificial neural network approach to partial differential equations in complex geometries. *Neurocomputing*, 317:28 – 41, 2018.
- [8] F. H. Branin. Widely convergent method for finding multiple solutions of simultaneous nonlinear equations. *IBM J. Res. Develop.*, 16:504–522, 1972.
- [9] K. M. Brown and W. B. Gearhart. Deflation techniques for the calculation of further solutions of a nonlinear

- system. *Numer. Math.*, 16:334–342, 1971.
- [10] W. Cai, X. Li, and L. Liu. A phase shift deep neural network for high frequency approximation and wave problems. *arXiv e-prints*, arXiv:1909.11759, 2019.
 - [11] W. Cai, X. Li, and L. Liu. PhaseDNN - a parallel phase shift deep neural network for adaptive wideband learning. *arXiv e-prints*, arXiv:1905.01389, 2019.
 - [12] W. Cai and Z. J. Xu. Multi-scale deep neural networks for solving high dimensional PDEs. *arXiv e-prints*, arXiv:1910.11710, 2019.
 - [13] Yuan Cao, Zhiying Fang, Yue Wu, Ding-Xuan Zhou, and Quanquan Gu. Towards understanding the spectral bias of deep learning, 2020.
 - [14] K.-S. Chao, D.-K. Liu, and C.-T. Pan. A systematic search method for obtaining multiple solutions of simultaneous nonlinear equations. *IEEE Trans. Circuits Systems*, 22:748–753, 1975.
 - [15] Z. Chen, Y. Cao, D. Zou, and Q. Gu. How much over-parameterization is sufficient to learn deep relu networks? *CoRR*, arXiv:1911.12360, 2019.
 - [16] M.-J. Chien. Searching for multiple solutions of nonlinear systems. *IEEE Trans. Circuits Systems*, 26:817–827, 1979.
 - [17] C. J. Chyan and J. Henderson. Multiple solutions for 2mth order sturmliouville boundary value problems. *Comput. Math. Appl.*, 40:231–237, 2000.
 - [18] L. H. Clark, P. M. Schlosser, and J. F. Selgrade. Multiple stable periodic solutions in a model for hormonal control of the menstrual cycle. *Bull. Math. Biol.*, 65:157173, 2003.
 - [19] M. Croci and P. E. Farrell. Distinct solutions of finite-dimensional complementarity problems. *arXiv e-prints*, arXiv:1510.02433, 2015.
 - [20] X. Dai and Y. Zhu. Towards theoretical understanding of large batch training in stochastic gradient descent. *CoRR*, abs/1812.00542, 2018.
 - [21] D. F. Davidenko. On a new method of numerical solution of systems of nonlinear equations. *Dokl. Akad. Nauk SSSR*, 88:601–602, 1953.
 - [22] J. M. Davis, L. H. Erbe, and J. Henderson. Multiplicity of positive solutions for higher order sturm-liouville problems. *Rocky Mountain J. Math.*, 31:169–184, 2001.
 - [23] M. W. M. G. Dissanayake and N. Phan-Thien. Neural-network-based approximations for solving partial differential equations. *Communications in Numerical Methods in Engineering*, 10(3):195–201, 1994.
 - [24] S. S. Du, X. Zhai, B. Póczos, and A. Singh. Gradient descent provably optimizes over-parameterized neural networks. *arXiv e-prints*, arXiv:1810.02054, 2018.
 - [25] J. Duchi, E. Hazan, and Y. Singer. Adaptive subgradient methods for online learning and stochastic optimization. *J. Mach. Learn. Res.*, 12:2121–2159, 2011.
 - [26] B. Dyda, A. Kuznetsov, and M. Kwaśnicki. Eigenvalues of the fractional laplace operator in the unit ball. *J. Lond. Math. Soc.*, 95(2):500–518, 2017.
 - [27] W. E and Q. Wang. Exponential convergence of the deep neural network approximation for analytic functions. *CoRR*, abs/1807.00297, 2018.
 - [28] W. E and B. Yu. The deep ritz method: a deep learning-based numerical algorithm for solving variational problems. *Commun. Math. Stat.*, 6:1–12, 2018.
 - [29] P. E. Farrell, Á. Birkisson, and S. W. Funke. Deflation techniques for finding distinct solutions of nonlinear partial differential equations. *SIAM J. Sci. Comput.*, 37(4):A2026–A2045, 2015.
 - [30] M. C. Ferris and J. S. Pang. Engineering and economic applications of complementarity problems. *SIAM Rev.*, 39:669–713, 1997.
 - [31] B. Fornberg and J.A.C. Weideman. A numerical methodology for the Painlevé equations. *J. Comput. Phys.*, 230(15):5957 – 5973, 2011.
 - [32] X. Glorot and Y. Bengio. Understanding the difficulty of training deep feedforward neural networks. In Yee Whye Teh and Mike Titterton, editors, *Proceedings of the Thirteenth International Conference on Artificial Intelligence and Statistics*, volume 9 of *Proceedings of Machine Learning Research*, pages 249–256, Chia Laguna Resort, Sardinia, Italy, 2010. PMLR.
 - [33] D. Gobovic and M. E. Zaghloul. Analog cellular neural network with application to partial differential equations with variable mesh-size. In *Proceedings of IEEE International Symposium on Circuits and Systems - ISCAS '94*, volume 6, pages 359–362 vol.6, May 1994.
 - [34] I. Goodfellow, Y. Bengio, and A. Courville. *Deep Learning*. MIT Press, Cambridge, 2016.
 - [35] J. R. Graef, C. Qian, and B. Yang. Multiple positive solutions of a boundary value problem for ordinary differential equations. *Electron. J. Qual. Theo. [electronic only]*, 2003:Paper No. 11, 13 p., electronic only–Paper No. 11, 13 p., electronic only, 2003.
 - [36] J. R. Graef, C. Qian, and B. Yang. A three point boundary value problem for nonlinear fourth order differential equations. *J. Math. Anal. Appl.*, 287(1):217 – 233, 2003.
 - [37] Y. Gu, H. Yang, and C. Zhou. SelectNet: Self-paced learning for high-dimensional partial differential equations.

- arXiv e-prints*, arXiv:2001.04860, 2020.
- [38] J. Han, A. Jentzen, and W. E. Solving high-dimensional partial differential equations using deep learning. *Proc. Natl. Acad. Sci.*, 115(34):8505–8510, 2018.
 - [39] W. Hao, J. D. Hauenstein, B. Hu, and A. J. Somme. A bootstrapping approach for computing multiple solutions of differential equations. *J. Comput. Appl. Math.*, 258:181–190, 2014.
 - [40] S. P. Hastings and W. C. Troy. On some conjectures of turcotte, spence, bau, and holmes. *SIAM J. Math. Anal.*, 20(3):634–642, 1989.
 - [41] K. He, X. Zhang, S. Ren, and J. Sun. Delving deep into rectifiers: Surpassing human-level performance on imagenet classification. In *The IEEE International Conference on Computer Vision (ICCV)*, December 2015.
 - [42] P. Holmes and D. Spence. On a Painlevé-type boundary-value problem. *Q. J. Mech. Appl. Math.*, 37(4):525–538, 1984.
 - [43] Jianguo Huang, Haoqin Wang, and Haizhao Yang. Int-deep: A deep learning initialized iterative method for nonlinear problems. *Journal of Computational Physics*, 419:109675, 2020.
 - [44] M. Hutzenthaler, A. Jentzen, Th. Kruse, and T. A. Nguyen. A proof that rectified deep neural networks overcome the curse of dimensionality in the numerical approximation of semilinear heat equations. Technical Report 2019-10, Seminar for Applied Mathematics, ETH Zürich, Switzerland, 2019.
 - [45] A. Jacot, F. Gabriel, and C. Hongler. Neural tangent kernel: Convergence and generalization in neural networks. *CoRR*, abs/1806.07572, 2018.
 - [46] S. Justin and S. Konstantinos. Dgm: A deep learning algorithm for solving partial differential equations. *J. Comput. Phys.*, 375:1339–1364, 2018.
 - [47] A. Karpatne, G. Atluri, J. H. Faghmous, M. Steinbach, A. Banerjee, A. Ganguly, S. Shekhar, N. Samatova, and V. Kumar. Theory-guided data science: A new paradigm for scientific discovery from data. *IEEE Transactions on Knowledge and Data Engineering*, 29:2318–2331, 2017.
 - [48] Y. Khoo, J. Lu, and L. Ying. Solving parametric pde problems with artificial neural networks. *arXiv: Numerical Analysis*, 2017.
 - [49] A. A. Kilbas, H. M. Srivastava, and J. J. Trujillo. *Theory and Applications of Fractional Differential Equations*. Elsevier Science, Amsterdam, The Netherlands, 2006.
 - [50] D. P. Kingma and J. Ba. Adam: a method for stochastic optimization. *arXiv e-prints*, arXiv:1412.6980, 2014.
 - [51] I. E. Lagaris, A. Likas, and D. I. Fotiadis. Artificial neural networks for solving ordinary and partial differential equations. *IEEE Transactions on Neural Networks*, 9(5):987–1000, Sep. 1998.
 - [52] Hyuk Lee and In Seok Kang. Neural algorithm for solving differential equations. *Journal of Computational Physics*, 91(1):110 – 131, 1990.
 - [53] D. Lei, Z. Sun, Y. Xiao, and W. Y. Wang. Implicit regularization of stochastic gradient descent in natural language processing: observations and implications. *arXiv e-prints*, arXiv:1811.00659, 2018.
 - [54] S. Liang and R. Srikant. Why deep neural networks? *CoRR*, abs/1610.04161, 2016.
 - [55] S. Liao. *Homotopy Analysis Method in Nonlinear Differential Equations*. Springer, Berlin, Heidelberg, 2012.
 - [56] Y. Liao and P. Ming. Deep Nitsche method: Deep Ritz method with essential boundary conditions. *arXiv e-prints*, arXiv:1912.01309, 2019.
 - [57] J. Lu, Z. Shen, H. Yang, and S. Zhang. Deep network approximation for smooth functions. *arXiv e-prints*, arXiv:2001.03040, 2020.
 - [58] T. Luo, Z. Ma, Z. J. Xu, and Y. Zhang. Theory of the frequency principle for general deep neural networks. *CoRR*, abs/1906.09235, 2019.
 - [59] T. Luo and H. Yang. Two-layer neural networks for partial differential equations: Optimization and generalization theory. *arXiv e-prints*, arXiv:2006.15733, 2020.
 - [60] H. Montanelli and Q. Du. New error bounds for deep networks using sparse grids. *arXiv e-prints*, arXiv:1712.08688, 2017.
 - [61] H. Montanelli and H. Yang. Error bounds for deep relu networks using the kolmogorovarnold superposition theorem. *Neural Networks*, 129:1–6, 2020.
 - [62] H. Montanelli, H. Yang, and Q. Du. Deep ReLU networks overcome the curse of dimensionality for bandlimited functions. *arXiv e-prints*, arXiv:1903.00735, 2019.
 - [63] B. Neyshabur, R. Tomioka, R. Salakhutdinov, and N. Srebro. Geometry of optimization and implicit regularization in deep learning. *arXiv e-prints*, arXiv:1705.03071, 2017.
 - [64] G. Nicolis and I. Prigogine. *SelfOrganization in Nonequilibrium Systems*. Wiley, New York, 1977.
 - [65] V. A. Noonburg. A separating surface for the Painlevé differential equation $x'' = x^2 - t$. *J. Math. Anal. Appl.*, 193(3):817 – 831, 1995.
 - [66] J. A. A. Opschoor, C. Schwab, and J. Zech. Exponential relu dnn expression of holomorphic maps in high dimension. Technical report, Zurich, 2019.
 - [67] Owens and Filkin. Efficient training of the backpropagation network by solving a system of stiff ordinary

- differential equations. In *International 1989 Joint Conference on Neural Networks*, pages 381–386 vol.2, 1989.
- [68] S. Pan and K. Duraisamy. Physics-informed probabilistic learning of linear embeddings of nonlinear dynamics with guaranteed stability. *SIAM Journal on Applied Dynamical Systems*, 19:480–509, 2020.
 - [69] J. E. Pearson. Complex patterns in a simple system. *Science*, 261:189–192, 1993.
 - [70] T. Poggio, H. N. Mhaskar, L. Rosasco, B. Miranda, and Q. Liao. Why and when can deep—but not shallow—networks avoid the curse of dimensionality: A review. *International Journal of Automation and Computing*, 14:503–519, 2017.
 - [71] M. Raissi, P. Perdikaris, and G. E. Karniadakis. Physics-informed neural networks: A deep learning framework for solving forward and inverse problems involving nonlinear partial differential equations. *Journal of Computational Physics*, 378:686 – 707, 2019.
 - [72] S. J. Reddi, S. Kale, and S. Kumar. On the convergence of Adam and beyond. *arXiv e-prints*, arXiv:1904.09237, 2019.
 - [73] Sifan Wang, Xinling Yu, and Paris Perdikaris. When and why pinns fail to train: A neural tangent kernel perspective. *arXiv:2007.14527*, 2020.
 - [74] Y. Wang, W. Hao, and G. Lin. Two-level spectral methods for nonlinear elliptic equations with multiple solutions. *SIAM Journal on Scientific Computing*, 40:B1180–B1205, 2018.
 - [75] J. H. Wilkinson. Rounding errors in algebraic processes. *Natl. Phys. Lab. Notes Appl. Sci.*, 32:334–342, 1963.
 - [76] L. R. Williams and R. W. Leggett. Unique and multiple solutions of a family of differential equations modeling chemical reactions. *SIAM J. Math. Anal.*, 13:122–133, 1982.
 - [77] H. Xu, Z. Lin, S. Liao, J. Wu, and J. Majdalani. Homotopy based solutions of the navier-stokes equations for a porous channel with orthogonally moving walls. *Phys. Fluids*, 22:053601, 2010.
 - [78] Z. J. Xu, Y. Zhang, and Y. Xiao. Training behavior of deep neural network in frequency domain. In *Neural Information Processing*, pages 264–274. Springer International Publishing, 2019.
 - [79] D. Yarotsky. Error bounds for approximations with deep relu networks. *Neural Networks*, 94:103–114, 2017.
 - [80] D. Yarotsky and A. Zhevnerchuk. The phase diagram of approximation rates for deep neural networks. *arXiv e-prints*, arXiv:1906.09477, 2019.
 - [81] Y. A. Yucesan and F. A. C. Viana. A physics-informed neural network for wind turbine main bearing fatigue. *International Journal of Prognostics and Health Management*, 11, 2020.
 - [82] Z. Song Z. A.-Zhu, Y. Li. A convergence theory for deep learning via over-parameterization. In Kamalika Chaudhuri and Ruslan Salakhutdinov, editors, *Proceedings of the 36th International Conference on Machine Learning*, volume 97 of *Proceedings of Machine Learning Research*, pages 242–252, Long Beach, California, USA, 2019. PMLR.
 - [83] Y. Zang, G. Bao, X. Ye, and H. Zhou. Weak adversarial networks for high-dimensional partial differential equations. *J. Comput. Phys.*, 411:109409, 2020.
 - [84] Y. Zhang, Z. J. Xu, T. Luo, and Z. Ma. Explicitizing an implicit bias of the frequency principle in two-layer neural networks. *CoRR*, abs/1905.10264, 2019.
 - [85] Y. Zhang, Z. J. Xu, T. Luo, and Z. Ma. A type of generalization error induced by initialization in deep neural networks. *CoRR*, abs/1905.07777, 2019.