

FAST ITERATIVE SOLVER FOR NEURAL NETWORK METHOD: I. 1D DIFFUSION PROBLEMS *

ZHIQIANG CAI[†], ANASTASSIA DOKTOROVA[†], ROBERT D. FALGOUT[‡], AND CÉSAR HERRERA[†]

Abstract. The discretization of the deep Ritz method [18] for the Poisson equation leads to a high-dimensional non-convex minimization problem, that is difficult and expensive to solve numerically. In this paper, we consider the shallow Ritz approximation to one-dimensional diffusion problems and introduce an effective and efficient iterative method, a damped block Newton (dBN) method, for solving the resulting non-convex minimization problem.

The method employs the block Gauss-Seidel method as an outer iteration by dividing the parameters of a shallow neural network into the linear parameters (the weights and bias of the output layer) and the non-linear parameters (the weights and bias of the hidden layer). Per each outer iteration, the linear and the non-linear parameters are updated by exact inversion and one step of a damped Newton method, respectively. Inverses of the coefficient matrix and the Hessian matrix are tridiagonal and diagonal, respectively, and hence the cost of each dBN iteration is $\mathcal{O}(n)$. To move the breakpoints (the non-linear parameters) more efficiently, we propose an adaptive damped block Newton (AdBN) method by combining the dBN with the adaptive neuron enhancement (ANE) method [25]. Numerical examples demonstrate the ability of dBN and AdBN not only to move the breakpoints quickly and efficiently but also to achieve a nearly optimal order of convergence for AdBN. These iterative solvers are capable of outperforming BFGS for select examples.

Key words. Fast iterative solvers, Neural network, Ritz formulation, ReLU activation, Diffusion problems, Elliptic problems, Newton's method

1. Introduction. In the past decade, the use of neural networks (NNs) has surged in popularity, finding applications in artificial intelligence, natural language processing, image recognition, and various other domains within machine learning. Consequently, the use of NNs has extended to diverse fields, including numerically solving partial differential equations (PDEs) [4, 9, 16, 18, 29, 31]. The idea of using NNs to solve PDEs may be traced back to 90s [15, 22, 23] based on some versions of the least-squares principle. Neural networks give rise to a novel class of functions characterized by piecewise linearity, offering the advantage of adaptability in the physical partition [7, 8, 25].

In one dimension, the shallow neural networks produce the same class of approximating functions as the free knots splines introduced in 1960s and studied by many researchers [3, 14, 21, 30]. Spline approximations to non-smooth functions can be improved dramatically with free knots [5]. Nevertheless, determining optimal knot locations (the non-linear parameters) becomes a complicated, computationally intensive non-convex optimization problem and was a subject of many research articles on various optimization methods such as the DFP method [13, 19], the Gauss-Newton method [20], a method moving each knot locally [14, 26], etc. (see, e.g., [21, 28] and references therein).

In the context of the shallow neural network for multi-dimension, the active neuron least squares (ANLS) method was recently introduced in [1, 2] to efficiently update the non-linear parameters. By utilizing both the quadratic structure of the functional and the neural network structure, the structure-guided Gauss-Newton (SgGN) method was newly proposed in [10] for solving the non-linear least-squares problem. For several one and two dimensional least-squares test problems which are difficult for the commonly used training algorithms in machine learning such as BFGS and Adam, the SgGN shows superior convergence. However, the mass matrix for the linear parameters and the layer Gauss-Newton matrix are ill-conditioned even though they are

*This work was supported in part by the National Science Foundation under grant DMS-2110571. This work was performed under the auspices of the U.S. Department of Energy by Lawrence Livermore National Laboratory under Contract DE-AC52-07NA27344 (LLNL-JRNL-862433).

[†]Department of Mathematics, Purdue University, West Lafayette, IN (caiz@purdue.edu, adoktoro@purdue.edu, herre125@purdue.edu).

[‡]Lawrence Livermore National Laboratory, Livermore, CA (rfalgout@llnl.gov)

symmetric and positive definite.

The deep Ritz method introduced in [18] is a discretization method for the Poisson equation using deep neural networks as approximating functions. The method is based on the Ritz minimization formulation of the underlying problem. The resulting discretization is a high-dimensional and computationally intensive non-convex minimization problem.

The purpose of this paper is to develop a fast iterative solver, a damped block Newton (dBN) method, for numerically solving the above mentioned non-convex minimization problem arising from the shallow Ritz discretization of one-dimensional diffusion problems. Like existing approaches, the neural network parameters are assorted as the linear and non-linear parameters. The outer iteration of the dBN method for the resulting algebraic system from the optimality conditions is based on the block Gauss-Seidel method, and the inner iteration requires a linear solver for the linear parameters and a damped Newton iteration for the non-linear parameters.

For the linear parameters, the corresponding coefficient matrix is no longer sparse and depends on the non-linear parameters. Moreover, its condition number is bounded by $\mathcal{O}(n h_{\min}^{-1})$ (see Lemma 3.1), the same as that obtained when using local hat basis functions [17], where n is the number of neurons and $h_{\min}$ is the smallest gap between two neighboring breakpoints. Instead of an iterative solver that requires many matrix and vector multiplications, the coefficient matrix is inverted by an explicit formula per each outer iteration and the inversion is a tridiagonal matrix. For the non-linear parameters, the corresponding Hessian matrix is diagonal. Hence, computational cost of each dBN iteration is $\mathcal{O}(n)$. To move the breakpoints (the non-linear parameters) more efficiently, we propose an adaptive damped block Newton (AdBN) method by combining the dBN with the adaptive neuron enhancement (ANE) method [25]. Numerical examples demonstrate the ability of dBN and AdBN not only to move the breakpoints quickly and efficiently but also to achieve a nearly optimal order of convergence for AdBN. These iterative solvers are capable of outperforming BFGS for select examples.

The remainder of the paper is structured as follows: In section 2 the Poisson equation and the modified Ritz formulation are introduced and an error estimate is presented. Then, in section 3, the linear and non-linear algebraic systems are laid out. Additionally, a bound for the condition number of the coefficient matrix for the linear system is given and the Hessian matrix for the non-linear system is presented. In section 4, the formula for the inverse of the coefficient matrix is revealed. Subsequently, the dBN is outlined in the remainder of the section. This algorithm solves for the linear and non-linear parameters separately while using the derived formulas for the inverses of the matrices. In section 5 an adaptivity scheme is introduced to improve the convergence of the method. Lastly, numerical results are presented in section 6, demonstrating the benefit of adaptivity as well as the ability of dBN to outperform BFGS for those select examples.

2. Poisson's Equation and Neural Network Approximation. Consider the following one-dimensional Poisson equation

$$(2.1) \quad \begin{cases} -(a(x)u'(x))' = f(x), & x \in I = (0, 1), \\ u(0) = \alpha, & u(1) = \beta, \end{cases}$$

where f is a given real-valued function defined on I . Assume that the diffusion coefficient $a(x)$ is bounded below by a positive constant $\mu > 0$ almost everywhere on I .

One of the Dirichlet boundary conditions is enforced strongly and the other can be enforced either algebraically (see Appendix A for the formulation) or weakly by penalizing the energy functional. We opt for the latter; the modified Ritz formulation of problem (2.1) is to find $u \in H^1(I) \cap \{u(0) = \alpha\}$ such that

$$(2.2) \quad J(u) = \min_{v \in H^1(I) \cap \{v(0) = \alpha\}} J(v),$$

where the modified energy functional is given by

$$J(v) = \frac{1}{2} \int_0^1 a(x)(v'(x))^2 dx - \int_0^1 f(x)v(x) dx + \frac{\gamma}{2}(v(1) - \beta)^2.$$

Here, $\gamma > 0$ is a penalization constant.

Denote by $\sigma(x) = \max\{0, x\}$ the ReLU activation function, and let

$$(2.3) \quad \mathcal{M}_n(I) = \left\{ c_{-1} + \sum_{i=0}^n c_i \sigma(x - b_i) : c_i \in \mathbb{R}, 0 \leq b_0 < b_1 < \dots < b_n < b_{n+1} = 1 \right\}$$

be the collection of the shallow neural network functions. Then the Ritz neural network approximation is to find $u_n \in \mathcal{M}_n(I) \cap \{u_n(0) = \alpha\}$ such that

$$(2.4) \quad J(u_n) = \min_{v \in \mathcal{M}_n(I) \cap \{v(0) = \alpha\}} J(v).$$

2.1. Error Estimate. Define the bilinear and linear forms by

$$(2.5) \quad a(u, v) := \int_0^1 a(x)u'(x)v'(x) dx + \gamma u(1)v(1) \quad \text{and} \quad f(v) := \int_0^1 f(x)v(x) dx + \gamma \beta v(1),$$

respectively, and denote the induced norm of the bilinear form by $\|v\|_a^2 = a(v, v)$. Then the modified energy function is given by

$$(2.6) \quad J(v) = \frac{1}{2}a(u, v) - f(v) + \frac{1}{2}\gamma\beta^2.$$

LEMMA 2.1. *Let u and u_n be the solutions of problems (2.2) and (2.4), respectively. Then*

$$(2.7) \quad \|u - u_n\|_a \leq \sqrt{3} \inf_{v \in \mathcal{M}_n(I) \cap \{v(0) = \alpha\}} \|u - v\|_a + \sqrt{2} |a(1)u'(1)| \gamma^{-1/2}.$$

Proof. It is easy to see that the solution u of (2.1) satisfies

$$(2.8) \quad a(u, v) + \alpha a(0)u'(0) - a(1)u'(1)v(1) = f(v)$$

for any $v \in H^1(I) \cap \{v(0) = \alpha\}$. Together with (2.6), we have

$$a(v - u, v - u) = 2(J(v) - J(u)) + a(1)u'(1)(u(1) - v(1)),$$

which, combining with the inequality that $2cd \leq \gamma^{-1}c^2 + \gamma d^2$, implies

$$\|v - u\|_a^2 - |a(1)u'(1)|^2 \gamma^{-1} \leq 4(J(v) - J(u)) \leq 3\|v - u\|_a^2 + |a(1)u'(1)|^2 \gamma^{-1}.$$

Now, together with the fact that $J(u_n) \leq J(v)$ for any $v \in \mathcal{M}_n(I) \cap \{v(0) = \alpha\}$, we have

$$\|u_n - u\|_a^2 \leq 4(J(v) - J(u)) + |a(1)u'(1)|^2 \gamma^{-1} \leq 3\|v - u\|_a^2 + 2|a(1)u'(1)|^2 \gamma^{-1}.$$

This implies (2.7) and proves the lemma. $\square$

The set of the shallow neural network functions, $\mathcal{M}_n(I)$, is equivalent to the set of continuous piecewise linear functions with free knots, i.e., the so-called free knot linear spline functions, (see, e.g., [12, 30]). It is well-known that there exists a constant $C(u)$ depending on u such that

$$(2.9) \quad \inf_{v \in \mathcal{M}_n(I)} \|a^{1/2}(u - v)'\|_{L^2(I)} \leq C(u) n^{-1},$$

provided that u has certain smoothness. Obviously, (2.9) is valid for $u \in H^2(I) = W^{2,2}(I)$ even when knots are fixed and form a quasi-uniform partition of the interval I . It is expected that (2.9) is also valid for $u \in W^{2,1}(I)$ when knots are free.

PROPOSITION 2.1. Let u and u_n be the solutions of the problem (2.2) and (2.4), respectively. Assume that $a \in L^\infty(I)$, then there exists a constant C depending on u such that

$$\|u - u_n\|_a \leq C \left(n^{-1} + \gamma^{-1/2} \right).$$

Proof. The proposition is a direct consequence of Lemma 2.1 and (2.9). $\square$

3. Systems of Algebraic Equations. Let

$$u_n = u_n(x) = u_n(x; \mathbf{c}, \mathbf{b}) = \alpha + \sum_{i=0}^n c_i \sigma(x - b_i)$$

be a solution of (2.4). Then the linear and non-linear parameters

$$\mathbf{c} = (c_0, \dots, c_n)^T \quad \text{and} \quad \mathbf{b} = (b_0, \dots, b_n)^T$$

satisfy the following system of algebraic equations

$$(3.1) \quad \nabla_{\mathbf{c}} J(u_n) = \mathbf{0} \quad \text{and} \quad \nabla_{\mathbf{b}} J(u_n) = \mathbf{0},$$

where $\nabla_{\mathbf{c}}$ and $\nabla_{\mathbf{b}}$ denote the gradients with respect to the respective parameters $\mathbf{c}$ and $\mathbf{b}$.

Let $H(t)$ be the Heaviside step function given by

$$H(t) = \sigma'(t) = \begin{cases} 1, & t > 0, \\ 0, & t < 0. \end{cases}$$

Obviously, we have

$$\nabla_{\mathbf{c}} u'_n(x) = \begin{pmatrix} H(x - b_0) \\ \vdots \\ H(x - b_n) \end{pmatrix}, \quad \nabla_{\mathbf{c}} u_n(x) = \begin{pmatrix} \sigma(x - b_0) \\ \vdots \\ \sigma(x - b_n) \end{pmatrix}, \quad \text{and} \quad \nabla_{\mathbf{c}} u_n(1) = \begin{pmatrix} b - b_0 \\ \vdots \\ b - b_n \end{pmatrix}.$$

Let $\mathbf{d} = \nabla_{\mathbf{c}} u_n(1)$, then the first equation in (3.1) gives

$$\begin{aligned} \mathbf{0} &= \int_0^1 a(x) u'_n(x) \nabla_{\mathbf{c}} u'_n(x) dx - \int_0^1 f(x) \nabla_{\mathbf{c}} u_n(x) dx + \gamma(u_n(1) - \beta) \nabla_{\mathbf{c}} u_n(1) \\ &= \left[\int_0^1 a(x) \nabla_{\mathbf{c}} u'_n(x) (\nabla_{\mathbf{c}} u'_n(x))^T dx \right] \mathbf{c} - \int_0^1 f(x) \nabla_{\mathbf{c}} u_n(x) dx + \gamma(\mathbf{d}^T \mathbf{c} + \alpha - \beta) \mathbf{d} \\ &= \left(A(\mathbf{b}) + \gamma \mathbf{d} \mathbf{d}^T \right) \mathbf{c} - \left(\mathbf{f}(\mathbf{b}) + \gamma(\beta - \alpha) \mathbf{d} \right), \end{aligned}$$

where the coefficient matrix and right hand side vector are defined as

$$A(\mathbf{b}) = \int_0^1 a(x) \nabla_{\mathbf{c}} u'_n(x) (\nabla_{\mathbf{c}} u'_n(x))^T dx \quad \text{and} \quad \mathbf{f}(\mathbf{b}) = \int_0^1 f(x) \nabla_{\mathbf{c}} u_n(x) dx,$$

respectively. Their components are given by

$$a_{ij}(\mathbf{b}) = \int_0^1 a(x) H(x - b_{i-1}) H(x - b_{j-1}) dx \quad \text{and} \quad f_i = \int_0^1 f(x) \sigma(x - b_{i-1}) dx.$$

Hence, the first equation in (3.1) has the form of

$$(3.2) \quad (A(\mathbf{b}) + \gamma \mathbf{d} \mathbf{d}^T) \mathbf{c} = \mathbf{f}(\mathbf{b}) + \gamma(\beta - \alpha) \mathbf{d}.$$

In the rest of the paper, we assume that the breakpoints $\mathbf{b} = (b_0, \dots, b_n)^T$ satisfy

$$0 \leq b_0 < \dots < b_n < b_{n+1} = 1.$$

Let $h_i = b_i - b_{i-1}$ for $i = 1, \dots, n+1$ and $h_{\min} = \min_{1 \leq i \leq n+1} h_i$. The next lemma provides an upper bound for the condition number of the coefficient matrix $A(\mathbf{b})$.

LEMMA 3.1. *Let $a(x) = 1$ in (2.1), then the condition number of the coefficient matrix $A(\mathbf{b})$ is bounded by $\mathcal{O}(n h_{\min}^{-1})$.*

Proof. For any vector $\boldsymbol{\xi} = (\xi_0, \dots, \xi_n)^T \in \mathcal{R}^n$, denote its magnitude by $|\boldsymbol{\xi}| = \left(\sum_{i=0}^n \xi_i^2 \right)^{1/2}$. By the Cauchy-Schwarz inequality, we have

$$(3.3) \quad \begin{aligned} \boldsymbol{\xi}^t A(\mathbf{b}) \boldsymbol{\xi} &= \int_0^1 \left(\sum_{i=0}^n \xi_i H(x - b_i) \right)^2 dx \leq |\boldsymbol{\xi}|^2 \int_0^1 \left(\sum_{i=0}^n H(x - b_i)^2 \right) dx \\ &= |\boldsymbol{\xi}|^2 \sum_{i=0}^n (1 - b_i) < (n+1) |\boldsymbol{\xi}|^2. \end{aligned}$$

To estimate the lower bound of the quadratic form $\boldsymbol{\xi}^t A(\mathbf{b}) \boldsymbol{\xi}$, note that

$$(3.4) \quad \boldsymbol{\xi}^t A(\mathbf{b}) \boldsymbol{\xi} = \sum_{j=0}^n \int_{b_j}^{b_{j+1}} \left(\sum_{i=0}^j \xi_i H(x - b_i) \right)^2 dx = \sum_{j=0}^n h_{j+1} \left(\sum_{i=0}^j \xi_i \right)^2 \geq h_{\min} \sum_{j=0}^n \left(\sum_{i=0}^j \xi_i \right)^2.$$

Now, the lemma is a direct consequence of (3.3), (3.4), and the fact that

$$|\boldsymbol{\xi}|^2 = \sum_{j=0}^n \left(\sum_{i=0}^j \xi_i - \sum_{i=0}^{j-1} \xi_i \right)^2 \leq 2 \sum_{j=0}^n \left(\sum_{i=0}^j \xi_i \right)^2 + 2 \sum_{j=0}^n \left(\sum_{i=0}^{j-1} \xi_i \right)^2 \leq 4 \sum_{j=0}^n \left(\sum_{i=0}^j \xi_i \right)^2.$$

This completes the proof of the lemma. $\square$

In the remainder of this section, we derive the Hessian matrix for the non-linear parameters.

LEMMA 3.1. *For $j = 0, 1, \dots, n$, the j^{th} equation of $\nabla_{\mathbf{b}} J(u_n) = \mathbf{0}$ is given by*

$$\frac{\partial}{\partial b_j} J(u_n) = c_j \left(\int_{b_j}^1 f(x) dx - \gamma u_n(1) - a(b_j) \left(\sum_{i=0}^{j-1} c_i + \frac{c_j}{2} \right) + c_j \gamma \beta \right) = 0.$$

Proof. For each $j = 0, 1, \dots, n$, we have

$$\int_{b_j}^{b_{j+1}} a(x) (u'_n(x))^2 dx = \int_{b_j}^{b_{j+1}} a(x) \left(\sum_{i=0}^j c_i H(x - b_i) \right)^2 dx = \left(\sum_{i=0}^j c_i \right)^2 \int_{b_j}^{b_{j+1}} a(x) dx,$$

which yields

$$\frac{\partial}{\partial b_j} \int_0^1 a(x) (u'(x))^2 dx = a(b_j) \left(\sum_{i=0}^{j-1} c_i \right)^2 - a(b_j) \left(\sum_{i=0}^j c_i \right)^2 = -2c_j a(b_j) \left(\sum_{i=0}^{j-1} c_i + \frac{c_j}{2} \right).$$

It is easy to see that

$$\frac{\partial u_n(x)}{\partial b_j} = -c_j H(x - b_j) \quad \text{and} \quad \frac{\partial u_n(1)}{\partial b_j} = -c_j,$$

which implies

$$\frac{\partial}{\partial b_j} \int_0^1 f(x) u_n(x) dx = -c_j \int_{b_j}^1 f(x) dx \quad \text{and} \quad \frac{\partial}{\partial b_j} (u_n(1) - \beta)^2 = -2c_j (u_n(1) - \beta).$$

Hence,

$$0 = \frac{\partial}{\partial b_j} J(u_n) = -c_j \left[a(b_j) \left(\sum_{i=0}^{j-1} c_i + \frac{c_j}{2} \right) - \int_{b_j}^1 f(x) dx + \gamma (u_n(1) - \beta) \right].$$

This completes the proof of the lemma. $\square$

LEMMA 3.2. For $j = 0, 1, \dots, n$, let

$$g(b_j) = f(b_j) + a'(b_j) \left(\sum_{i=0}^{j-1} c_i + \frac{c_j}{2} \right).$$

Then the Hessian matrix $\nabla_{\mathbf{b}}^2 J(u_n)$ has the form

$$(3.5) \quad \mathcal{H}(\mathbf{c}, \mathbf{b}) = \begin{pmatrix} -c_0 g(b_0) & 0 & 0 & \dots & 0 \\ 0 & -c_1 g(b_1) & 0 & \dots & 0 \\ 0 & 0 & -c_2 g(b_2) & \dots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & 0 & \dots & -c_n g(b_n) \end{pmatrix} + \gamma \mathbf{c} \mathbf{c}^T \\ \equiv -\mathbf{B}(\mathbf{c}, \mathbf{b}) + \gamma \mathbf{c} \mathbf{c}^T.$$

Proof. For $k = 0, 1, \dots, n$, we have

$$\frac{\partial}{\partial b_k} c_j \left[\int_{b_j}^1 f(x) dx - a(b_j) \left(\sum_{i=0}^{j-1} c_i + \frac{c_j}{2} \right) \right] = \begin{cases} -c_j g(b_j), & k = j, \\ 0, & k \neq j. \end{cases}$$

Now (3.5) is a direct consequence of the fact that $\nabla_{\mathbf{b}} u_n(1) = -\mathbf{c}$ and Lemma 3.1. $\square$

4. A Damped Block Newton (dBN) Method. In this section, we introduce a damped block Newton (dBN) method for solving the resulting non-convex minimization problem in (2.4), i.e., the system of algebraic equations in (3.1). The method employs the block Gauss-Seidel method as outer iteration between the linear and non-linear parameters. Per each outer iteration, the linear and the non-linear parameters are updated by exact inversion and one step of a damped Newton method, respectively.

To this end, the coefficient matrix A has the form of

$$A(\mathbf{b}) = \begin{pmatrix} \int_{b_0}^1 a(x) dx & \int_{b_1}^1 a(x) dx & \int_{b_2}^1 a(x) dx & \dots & \int_{b_n}^1 a(x) dx \\ \int_{b_1}^1 a(x) dx & \int_{b_1}^1 a(x) dx & \int_{b_2}^1 a(x) dx & \dots & \int_{b_n}^1 a(x) dx \\ \int_{b_2}^1 a(x) dx & \int_{b_2}^1 a(x) dx & \int_{b_2}^1 a(x) dx & \dots & \int_{b_n}^1 a(x) dx \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \int_{b_n}^1 a(x) dx & \int_{b_n}^1 a(x) dx & \int_{b_n}^1 a(x) dx & \dots & \int_{b_n}^1 a(x) dx \end{pmatrix}.$$

Let $s_i = \int_{b_{i-1}}^{b_i} a(x) dx$ for $i = 1, \dots, n+1$. Then the inverse of $A(\mathbf{b})$ is a tri-diagonal matrix given in the following lemma.

LEMMA 4.1. *The coefficient matrix is invertible and its inverse is given by*

$$(4.1) \quad A(\mathbf{b})^{-1} = \begin{pmatrix} \frac{1}{s_1} & -\frac{1}{s_1} & 0 & 0 & \cdots & 0 & 0 \\ -\frac{1}{s_1} & \frac{1}{s_1} + \frac{1}{s_2} & -\frac{1}{s_2} & 0 & \cdots & 0 & 0 \\ 0 & -\frac{1}{s_2} & \frac{1}{s_2} + \frac{1}{s_3} & -\frac{1}{s_3} & \cdots & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots & \ddots & \vdots & \vdots \\ 0 & 0 & 0 & 0 & \cdots & \frac{1}{s_{n-1}} + \frac{1}{s_n} & -\frac{1}{s_n} \\ 0 & 0 & 0 & 0 & \cdots & -\frac{1}{s_n} & \frac{1}{s_n} + \frac{1}{s_{n+1}} \end{pmatrix}.$$

Proof. It is easy to verify that $A(\mathbf{b})^{-1}A(\mathbf{b}) = I$. $\square$

Let $\mathcal{A}(\mathbf{b}) = A(\mathbf{b}) + \gamma \mathbf{d} \mathbf{d}^T$, by Lemma 4.1 and the Sherman-Morrison formula, we have

$$\mathcal{A}(\mathbf{b})^{-1} = A(\mathbf{b})^{-1} - \frac{\gamma A(\mathbf{b})^{-1} \mathbf{d} \mathbf{d}^T A(\mathbf{b})^{-1}}{1 + \gamma \mathbf{d}^T A(\mathbf{b})^{-1} \mathbf{d}}$$

The following lemma states a sufficient condition for guaranteeing invertibility of the Hessian matrix and an explicit formula for its inversion.

LEMMA 4.1. *If $c_i g(b_i) \neq 0$ for all $i = 0, 1, \dots, n$ and $1 - \gamma \mathbf{c}^T \mathbf{B}^{-1}(\mathbf{c}, \mathbf{b}) \mathbf{c} \neq 0$, then the Hessian matrix $\mathcal{H}(\mathbf{c}, \mathbf{b})$ is invertible. Moreover, its inverse is given by*

$$(4.2) \quad \mathcal{H}^{-1}(\mathbf{c}, \mathbf{b}) = - \left[\mathbf{B}^{-1}(\mathbf{c}, \mathbf{b}) + \frac{\gamma \mathbf{B}^{-1}(\mathbf{c}, \mathbf{b}) \mathbf{c} \mathbf{c}^T \mathbf{B}^{-1}(\mathbf{c}, \mathbf{b})}{1 - \gamma \mathbf{c}^T \mathbf{B}^{-1}(\mathbf{c}, \mathbf{b}) \mathbf{c}} \right].$$

Proof. Clearly, by Lemma 3.2, the assumption implies that $\mathcal{H}(\mathbf{c}, \mathbf{b})$ is invertible. (4.2) is a direct consequence of the Sherman-Morrison formula. $\square$

Now, we are ready to describe the damped block Newton (dBN) method (see Algorithm 4.1 for a pseudocode). Let $(\mathbf{c}^{(k)}, \mathbf{b}^{(k)})$ be the previous iterate, then the current iterate $(\mathbf{c}^{(k+1)}, \mathbf{b}^{(k+1)})$ is computed as follows:

(i) Compute the linear parameters $\mathbf{c}^{(k+1)}$ by

$$\mathbf{c}^{(k+1)} = \mathcal{A}(\mathbf{b}^{(k)})^{-1} \left\{ \mathbf{f}(\mathbf{b}^{(k)}) + \gamma(\beta - \alpha) \mathbf{d}(\mathbf{b}^{(k)}) \right\}.$$

(ii) If $c_i^{(k+1)} g(b_i^{(k)}) \neq 0$ for all $i = 0, \dots, n$, compute the search direction

$$\mathbf{p}^{(k)} = -\mathcal{H}(\mathbf{c}^{(k+1)}, \mathbf{b}^{(k)})^{-1} \nabla_{\mathbf{b}} J(u_n(x; \mathbf{c}^{(k+1)}, \mathbf{b}^{(k)})).$$

(iii) Calculate the stepsize η_k by performing a one-dimensional optimization

$$\eta_k = \underset{\eta \in \mathbb{R}_+}{\operatorname{argmin}} J(u_n(x; \mathbf{c}^{(k+1)}, \mathbf{b}^{(k)} + \eta \mathbf{p}^{(k)})).$$

(iv) Set the non-linear parameters by

$$\mathbf{b}^{(k+1)} = \mathbf{b}^{(k)} + \eta_k \mathbf{p}^{(k)}.$$

Remark 4.1. *In the case that $c_i^{(k+1)} g(b_i^{(k)})$ vanishes for some $i \in \{0, \dots, n\}$, the step (ii) of the method is modified as follows: set $\mathbf{b}_i^{(k+1)} = \mathbf{b}_i^{(k)}$ and the i^{th} neuron is removed when computing the search direction $\mathbf{p}^{(k)}$.*

Algorithm 4.1 A damped block Newton method for (2.4)

Input: Initial network parameters $(\mathbf{c}^{(0)}, \mathbf{b}^{(0)})$, coefficient function $a(x)$, the right-hand side function $f(x)$, boundary data α and β .

Output: network parameters $\mathbf{c}, \mathbf{b}$

```

for  $k = 0, 1, \dots$  do
   $\triangleright$  Linear parameters
   $\mathbf{c}^{(k+1)} \leftarrow \mathcal{A}(\mathbf{b}^{(k)})^{-1} \{ \mathbf{f}(\mathbf{b}^{(k)}) + \gamma(\beta - \alpha) \mathbf{d}(\mathbf{b}^{(k)}) \}$ 
   $\triangleright$  non-linear parameters
   $\mathbf{p}^{(k)} \leftarrow -\mathcal{H}(\mathbf{c}^{(k+1)}, \mathbf{b}^{(k)})^{-1} \nabla_{\mathbf{b}} J(u_n(x; \mathbf{c}^{(k+1)}, \mathbf{b}^{(k)}))$ 
   $\eta_k \leftarrow \operatorname{argmin}_{\eta \in \mathbb{R}_+} J(u_n(x; \mathbf{c}^{(k+1)}, \mathbf{b}^{(k)} + \eta \mathbf{p}^{(k)}))$ 
   $\mathbf{b}^{(k+1)} \leftarrow \mathbf{b}^{(k)} + \eta_k \mathbf{p}^{(k)}$ 
end for

```

The computational cost of inverting the mass and the Hessian matrices in Algorithm 4.1 is $\mathcal{O}(n)$. More specifically, the linear parameters $\mathbf{c}^{(k+1)}$ and the direction vector $\mathbf{p}^{(k)}$ are calculated in $8(n+1)$ and $4(n+1)$ operations, respectively. This is significantly less than Quasi-Newton approaches like BFGS, where the computational cost per iteration is $\mathcal{O}(n^2)$ (see [27], Chapter 6).

Remark 4.2. *The minimization problem in (2.4) is non-convex and has many local and global minimums. The desired one is obtained only if we start from a close enough first approximation. As in [6, 24], in this paper we initialize the non-linear parameters $\mathbf{b}$ to form a uniform partition of the interval $[a, b]$ and the linear parameters $\mathbf{c}$ to be the corresponding solution on this uniform mesh.*

5. An adaptivity scheme. For a fixed number of neurons, the dBN method with the initialization described in Remark 4.2 moves the uniformly distributed breakpoints very efficiently to nearly optimal locations as shown in section 6. Nevertheless, those locations may not be the optimal positions to achieve optimal rate of convergence of the shallow Ritz approximation.

To circumvent this difficulty, we employ the adaptive neuron enhancement (ANE) method [24, 25] that starts with a relatively small neuron network and adaptively adds new neurons based on the previous approximation. Moreover, the newly added neurons are initialized at where the previous approximation is not accurate. At each adaptive step, we use the dBN method to numerically solve the minimization problem in (2.4).

A key component of the ANE is the marking strategy, defining by error indicators, that determines the number of new neurons to be added. Below, we describe the local indicators and the marking strategy used in this paper.

Let $\mathcal{K} = [c, d] \subseteq [0, 1]$ be a subinterval, a modified local indicator of the ZZ type on $\mathcal{K}$ (see, e.g., [11]) is defined by

$$\xi_{\mathcal{K}} = \|a^{-1/2} (G(au'_n) - au'_n)\|_{L^2(\mathcal{K})},$$

where $G(au'_n)$ is the projection of au'_n onto the space of the continuous piecewise linear functions. The corresponding relative error estimator is defined by

$$(5.1) \quad \xi = \frac{\|a^{-1/2} (G(au'_n) - au'_n)\|_{L^2(I)}}{\|u'_n\|_{L^2(I)}}.$$

For a given $u_n \in \mathcal{M}_n(I)$ with the breakpoints

$$0 = b_{-1} < b_0 < \dots < b_n < b_{n+1} = 1,$$

let $\mathcal{K}^i = [b_{i-1}, b_i]$, then $\mathcal{K}_n = \{\mathcal{K}^i\}_{i=0}^{n+1}$ is a partition of the interval $[0, 1]$. Define a subset $\hat{\mathcal{K}}_n \subset \mathcal{K}_n$

by using the following average marking strategy:

$$(5.2) \quad \widehat{\mathcal{K}}_n = \left\{ K \in \mathcal{K}_n : \xi_K \geq \frac{1}{\#\mathcal{K}_n} \sum_{K \in \mathcal{K}_n} \xi_K \right\},$$

where $\#\mathcal{K}_n$ is the number of elements in $\mathcal{K}_n$. Adaptive damped block Newton (AdBN) method is described in [Algorithm 5.1](#).

Algorithm 5.1 Adaptive damped block Newton (AdBN) method

Input: Initial number of neurons n_0 , parameters $a(x)$, $f(x)$, α , and β , tolerance ϵ ,

- (1) Compute an approximation to the solution u_n of the optimization problem in (2.4) by the dBN method;
 - (2) Compute the estimator $\xi_n = \left(\sum_{K \in \mathcal{K}} \xi_K^2 \right)^{1/2} / |u_n|_{H^1(I)}$;
 - (3) If $\xi_n \leq \epsilon$, then stop; otherwise go to step (4);
 - (4) Mark elements in $\widehat{\mathcal{K}}_n$ and denote by $\#\widehat{\mathcal{K}}_n$ the number of elements in $\widehat{\mathcal{K}}_n$;
 - (5) Add $\#\widehat{\mathcal{K}}_n$ new neurons to the network and initialize them at the midpoints of elements in $\widehat{\mathcal{K}}_n$, then go to step (1)
-

For numerical experiments presented in the subsequent section, the dBN method in (1) of [Algorithm 5.1](#) is stopped if the difference of the estimators for two consecutive iterates is less than a prescribed tolerance δ .

6. Numerical experiments. This section presents numerical results of the dBN and AdBN methods for solving (2.1). In all the experiments, the penalization parameter γ was set to 10^4 . For the AdBN method, a refinement occurred when the difference of the estimators for two consecutive iterates was less than 10^{-7} . In order to evaluate the performance of the iterative solvers described in [Algorithm 4.1](#) and [Algorithm 5.1](#), we compare the resulting approximations to the true solution. For each test problem, let u and u_n be the exact solution and its approximation in $\mathcal{M}_n(I)$, respectively. Denote the relative error by

$$e_n = \frac{|u - u_n|_{H^1(I)}}{|u|_{H^1(I)}}.$$

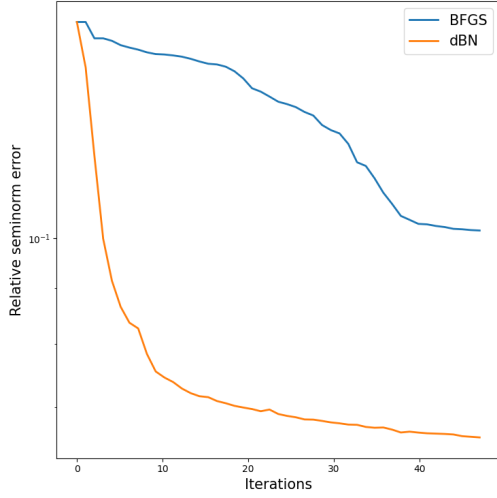
6.1. Exponential solution. The first test problem involves the function

$$(6.1) \quad u(x) = x \left(\exp \left(-\frac{(x - \frac{1}{3})^2}{0.01} \right) - \exp \left(-\frac{4}{9 \times 0.01} \right) \right),$$

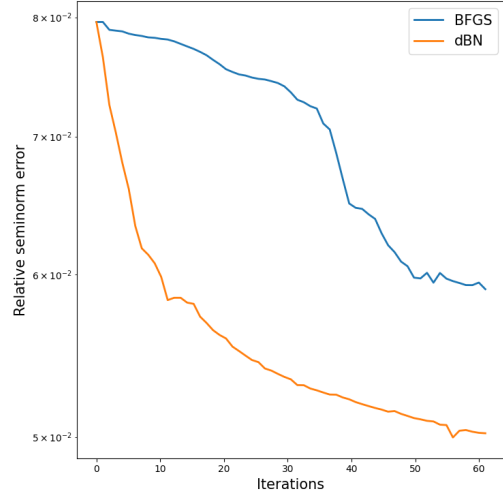
serving as the exact solution of (2.1).

We start by comparing the dBN method with a commonly used method: BFGS. In this comparison, we utilized a Python BFGS implementation from ‘scipy.optimize’. The initial network parameters for both algorithms were set to be the uniform mesh for $\mathbf{b}^{(0)}$ with $\mathbf{c}^{(0)}$ given by solving (3.1). We see in [Figure 1](#) that dBN requires about 15 iterations to get an accuracy that BFGS cannot achieve before stopping. Recall that the computational cost per iteration of dBN is $\mathcal{O}(n)$ while each iteration of BFGS has cost $\mathcal{O}(n^2)$. Not only does dBN decrease the relative error much more quickly than BFGS, but dBN also achieves a lower final error, which is further emphasized as the number of neurons increases.

In [Figure 2](#) (a), we present the initial neural network approximation of the exact solution in (6.1), obtained by using uniform breakpoints and determining the linear parameters through the



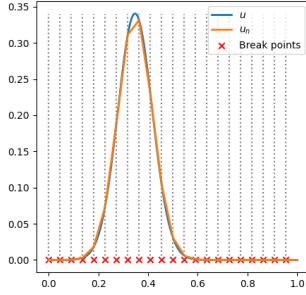
(a) Relative error e_n vs number of iterations using 32 neurons, the ratio between the final errors is 0.673



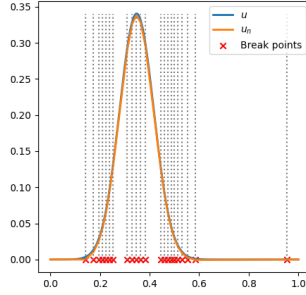
(b) Relative error e_n vs number of iterations using 64 neurons, the ratio between the final errors is 0.783

Fig. 1: Comparison between BFGS and dBN for approximating function (6.1)

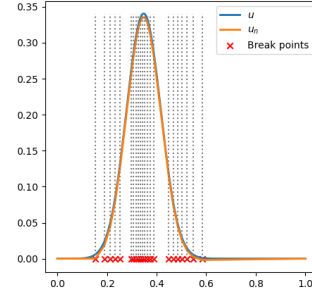
solution of (3.2). The approximations generated by the dBN and AdBN methods are shown in Figure 2 (b) and (c), respectively. We observe that our methods influence the movement of breakpoints, enhancing the overall approximation. Worth noting is that the placement of breakpoints by AdBN appears to be more optimal than that achieved by dBN in terms of relative error.



(a)



(b)



(c)

Fig. 2: Results of using ReLU networks to approximate the exact solution in (6.1): (a) initial NN model with 22 uniform breakpoints, $e_n = 0.227$, (b) optimized NN model with 22 breakpoints, 500 iterations, $e_n = 0.100$, (c) adaptive NN model with 22 breakpoints: 10 initial breakpoints, 2 refinements (13, 22 neurons), $e_n = 0.083$

To verify this observation, we estimate the order of convergence for the approximation to (6.1) in Algorithm 4.1. Recall that according to Proposition 2.1, it is theoretically possible to achieve an order of convergence of $\mathcal{O}(n^{-1})$. However, since (2.4) is a non-convex minimization problem,

the existence of local minimums makes it challenging to achieve this order. Therefore, given the neural network approximation u_n to u provided by the dBN method, assume that

$$e_n = \left(\frac{1}{n}\right)^r,$$

for some $r > 0$.

The larger this r is, the better the approximation. In [Table 1](#), dBN is applied for 1000 iterations for different values n , and the resulting r is calculated. Since

$$0.74 < r < 0.8,$$

it can be inferred that the algorithm gets stuck in a local minimum. Henceforth, to improve this order of convergence, we can use the AdBN method.

n	e_n	r
60	4.65×10^{-2}	0.749
90	3.41×10^{-2}	0.751
120	2.49×10^{-2}	0.771
150	1.97×10^{-2}	0.783
180	1.78×10^{-2}	0.775
210	1.59×10^{-2}	0.775
240	1.27×10^{-2}	0.796
270	1.22×10^{-2}	0.787
300	1.09×10^{-2}	0.792
330	1.01×10^{-2}	0.792
360	9.94×10^{-3}	0.783
390	9.54×10^{-3}	0.780
420	8.07×10^{-3}	0.798

Table 1: Relative errors e_n and powers r for n neurons after 1000 iterations.

In fact, adding adaptivity improves the r value. [Table 2](#) illustrates AdBN starting with 19 neurons, refining 8 times, and reaching a final count of 280 neurons. The stopping tolerance was set to $\epsilon = 0.01$. The recorded data in [Table 2](#) includes the relative seminorm error and the error estimator for each iteration of the adaptive process.

Additionally, [Table 2](#) provides the results for dBN with a fixed 187 and 280 neurons. Comparing these results to the adaptive run with the same number of neurons, we observe a significant improvement in rate, error estimator, and seminorm error within the adaptive run. The experiments confirm that AdBN improves the overall error. Furthermore, the order of convergence notably improves, especially with a larger number of neurons.

6.2. Non-smooth solution. The second test problem has the exact solution

$$(6.2) \quad u(x) = x^{2/3},$$

which belongs to $H^{1+\frac{1}{6}-\epsilon}(I)$ for any $\epsilon > 0$. We highlight that the order of convergence for approximating this solution with n uniform breakpoints is at most $\mathcal{O}(n^{-1/6})$. However, it is more optimal to concentrate more mesh points to the left side, where the curve is steeper. This adjustment is evident in [Figure 3](#) (b) and (c); the breakpoints shift to the left where the function exhibits the most curvature. This adjustment significantly improves the approximation compared to using uniform breakpoints, as illustrated in [Figure 3](#) (a). The relative errors confirm that the

NN (n neurons)	e_n	ξ_n	r
Adaptive (19)	1.04×10^{-1}	0.149	0.768
Adaptive (27)	6.55×10^{-2}	0.089	0.827
Adaptive (32)	5.71×10^{-2}	0.075	0.826
Adaptive (45)	4.21×10^{-2}	0.054	0.832
Adaptive (68)	2.76×10^{-2}	0.032	0.851
Adaptive (94)	2.01×10^{-2}	0.023	0.860
Adaptive (137)	1.39×10^{-2}	0.015	0.869
Adaptive (187)	1.04×10^{-2}	0.011	0.873
Adaptive (280)	6.93×10^{-3}	0.007	0.882
Fixed (187)	1.87×10^{-2}	0.020	0.761
Fixed (280)	1.12×10^{-2}	0.013	0.781

Table 2: Comparison of an adaptive network with fixed networks for relative errors e_n , relative error estimators ξ_n , and powers r

order of convergence improves substantially when the breakpoints are moved according to the steepness of the function.

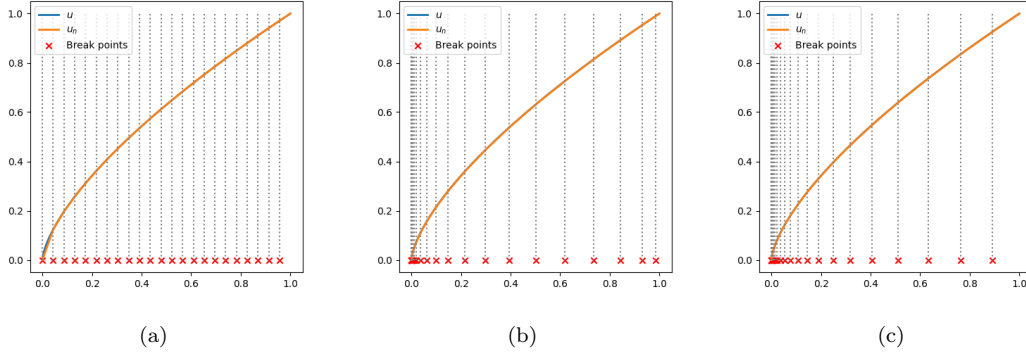


Fig. 3: Results of using ReLU networks for approximating function (6.2): (a) initial NN model with 23 uniform breakpoints, $e_n = 0.284$, (b) optimized NN model with 23 breakpoints, 500 iterations, $e_n = 0.056$, (c) adaptive NN model with 23 breakpoints: 9 initial breakpoints, 2 refinements (14, 23 neurons), $e_n = 0.042$

7. Discussion and Conclusion. The resulting discretization from the shallow Ritz approximation to one-dimensional diffusion problem is a high-dimensional minimization problem in (2.4). This is a computationally challenging problem to solve numerically due its non-convexity.

This paper proposed the dBN method to numerically solve this non-convex minimization problem. The dBN uses the block Gauss-Seidel method for the linear and non-linear parameters as outer iteration and inverts the corresponding coefficient and Hessian matrices exactly for inner iteration. Even though the coefficient matrix is dense and ill-conditioned, its inverse is tri-diagonal. Moreover, the Hessian matrix is diagonal. Hence, computational cost per each iteration of the dBN method is $\mathcal{O}(n)$ which is significantly faster than $\mathcal{O}(n^2)$ for the commonly used “second-order” methods. Numerical results demonstrated that the dBN method moves the uniformly distributed

breakpoints to nearly optimal locations very efficiently.

In order to achieve optimal rate of convergence of the shallow Ritz approximation, the paper proposed the AdBN method that empolys the adaptive neuron enhancement (ANE) method for adaptively adding neurons, initialized at where the previous approximation is inaccurate, and uses the dBN method for solving the minimization problem. Numerical results show that AdBN improves the rate of convergence.

When using the shallow ReLU neural network, the condition number of the coefficient matrix for the diffusion problem is bounded by $\mathcal{O}(n h_{\min}^{-1})$ (see [Lemma 3.1](#)), the same as that obtained when using local hat basis functions. For applications to general elliptic partial differential equations and least-squares data fitting problems, the corresponding dBN method requires inversion of the mass matrix. However, the condition number of the mass matrix is extremely large. This difficulty will be addressed in a forthcoming paper.

Appendix A. Enforcing Dirichlet Boundary Condition Algebraically. Another way to make a function $u_n \in \mathcal{M}_n(I) \cap \{u_n(0) = \alpha\}$ satisfy the Dirichlet boundary condition $u_n(1) = \beta$ is by enforcing this algebraically. Consider the energy functional given by

$$\mathcal{J}(v) = \frac{1}{2} \int_0^1 a(x)(v'(x))^2 dx - \int_0^1 f(x)v(x)dx.$$

Let $\mathbf{h} = \mathbf{h}(\mathbf{b}) = (\sigma(1 - b_0), \sigma(1 - b_1), \dots, \sigma(1 - b_n))^T$ and consider the Lagrangian function

$$\begin{aligned} \mathcal{L}(\mathbf{c}, \mathbf{b}, \lambda) &= \mathcal{J}(u_n(x; \mathbf{c}, \mathbf{b})) + \lambda(u_n(1) - \beta) \\ &= \mathcal{J}(u_n(x; \mathbf{c}, \mathbf{b})) + \lambda(\mathbf{h}^T \mathbf{c} + \alpha - \beta), \end{aligned}$$

hence, if u_n minimizes $\mathcal{J}(v)$ for $v \in \mathcal{M}_n(I) \cap \{u_n(0) = \alpha\}$, subject to the constraint $u_n(1) = \beta$, then by the KKT conditions:

$$(A.1) \quad \nabla_{\mathbf{c}} \mathcal{L}(u_n) = \mathbf{0}, \quad \frac{\partial}{\partial \lambda} \mathcal{L}(u_n) = 0 \quad \text{and} \quad \nabla_{\mathbf{b}} \mathcal{L}(u_n) = \mathbf{0}.$$

The first two equations in [\(A.1\)](#) can be written as

$$\begin{pmatrix} A(\mathbf{b}) & \mathbf{h} \\ \mathbf{h}^T & 0 \end{pmatrix} \begin{pmatrix} \mathbf{c} \\ \lambda \end{pmatrix} = \begin{pmatrix} \mathbf{f}(\mathbf{b}) \\ \beta - \alpha \end{pmatrix},$$

which can be solved efficiently, by writting the matrix in the left hand side as

$$(A.2) \quad \begin{pmatrix} A(\mathbf{b}) & \mathbf{h} \\ \mathbf{h}^T & 0 \end{pmatrix} = \begin{pmatrix} I & 0 \\ \mathbf{h}^T A(\mathbf{b})^{-1} & 1 \end{pmatrix} \begin{pmatrix} A(\mathbf{b}) & \mathbf{h} \\ 0 & -\mathbf{h}^T A(\mathbf{b})^{-1} \mathbf{h} \end{pmatrix},$$

and since

$$\begin{pmatrix} I & 0 \\ \mathbf{h}^T A(\mathbf{b})^{-1} & 1 \end{pmatrix}^{-1} = \begin{pmatrix} I & 0 \\ -\mathbf{h}^T A(\mathbf{b})^{-1} & 1 \end{pmatrix},$$

it follows that [\(A.2\)](#) is equivalent to

$$\begin{pmatrix} A(\mathbf{b}) & \mathbf{h} \\ 0 & -\mathbf{h}^T A(\mathbf{b})^{-1} \mathbf{h} \end{pmatrix} \begin{pmatrix} \mathbf{c} \\ \lambda \end{pmatrix} = \begin{pmatrix} I & 0 \\ -\mathbf{h}^T A(\mathbf{b})^{-1} & 1 \end{pmatrix} \begin{pmatrix} \mathbf{f}(\mathbf{b}) \\ \beta - \alpha \end{pmatrix}$$

which can be solved by finding λ first and then solving for $\mathbf{c}$. According to [\(4.1\)](#), the computational cost is $\mathcal{O}(n)$.

On the other hand, for the non-linear parameters $\mathbf{b}$, the Hessian matrix $\nabla_{\mathbf{b}}^2 \mathcal{L}(u_n) = -\mathbf{B}(\mathbf{c}, \mathbf{b})$, where $\mathbf{B}(\mathbf{c}, \mathbf{b})$ is defined in [\(3.5\)](#). Therefore, as described in [section 4](#), we solve [\(2.1\)](#) iteratively, alternating between solving exactly for $(\mathbf{c}^T, \lambda)^T$ and performing a Newton step for $\mathbf{b}$.

REFERENCES

- [1] M. Ainsworth and Y. Shin. Plateau phenomenon in gradient descent training of ReLU networks: Explanation, quantification and avoidance. *SIAM Journal on Scientific Computing*, 43:A3438–A3468, 2020.
- [2] M. Ainsworth and Y. Shin. Active neuron least squares: A training method for multivariate rectified neural networks. *SIAM Journal on Scientific Computing*, 44(4):A2253–A2275, 2022.
- [3] D. L. Barrow, C. K. Chui, P. W. Smith, and J. D. Ward. Unicity of best mean approximation by second order splines with variable knots. *Mathematics of Computation*, 32(144):1131–1143, 1978.
- [4] J. Berg and K. Nyström. A unified deep artificial neural network approach to partial differential equations in complex geometries. *Neurocomputing*, 317:28–41, 2018.
- [5] H. G. Burchard. Splines (with optimal knots) are better. *Applicable Analysis*, 3:309–319, 1974.
- [6] Z. Cai, J. Chen, and M. Liu. Least-squares ReLU neural network (LSNN) method for linear advection-reaction equation. *Journal of Computational Physics*, 443 (2021) 110514.
- [7] Z. Cai, J. Chen, and M. Liu. Least-squares ReLU neural network (LSNN) method for scalar nonlinear hyperbolic conservation law. *Applied Numerical Mathematics*, 174:163–176, 2022.
- [8] Z. Cai, J. Chen, and M. Liu. Least-squares neural network (LSNN) method for scalar nonlinear hyperbolic conservation laws: Discrete divergence operator. *Journal of Computational and Applied Mathematics*, 433:115298, 2023.
- [9] Z. Cai, J. Chen, M. Liu, and Xinyu Liu. Deep least-squares methods: An unsupervised learning-based numerical method for solving elliptic PDEs. *Journal of Computational Physics*, 420:109707, 2020.
- [10] Z. Cai, T. Ding, M. Liu, X. Liu, and J. Xia. A structure-guided gauss-newton method for shallow ReLU neural network. *arXiv:2404.05064v1 [cs.LG]*, 2024.
- [11] Z. Cai and S. Zhang. Recovery-based error estimators for interface problems: conforming linear elements. *SIAM Journal on Numerical Analysis*, 47(3):2132–2156, 2009.
- [12] I. Daubechies, R. DeVore, S. Foucart, B. Hanin, and G. Petrova. Nonlinear approximation and (deep) ReLU networks. *Constructive Approximation*, 55(1):127–172, February 2022.
- [13] W. C. Davidon. Variable metric method for minimization. *SIAM Journal on Optimization*, 1(1):1–17, 1991.
- [14] C. deBoor and J. R. Rice. Least squares cubic spline approximation ii: variable knots. *Department of Computer Sciences Purdue University, CSD TR 21*, 1968.
- [15] 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.
- [16] T. Dockhorn. A discussion on solving partial differential equations using neural networks. *arXiv:1904.07200 [cs.LG]*, abs/1904.07200, 2019.
- [17] Q. Du, D. Wang, and L. Zhu. On mesh geometry and stiffness matrix conditioning for general finite element spaces. *SIAM Journal on Numerical Analysis*, 47(2):1421–1444, 2009.
- [18] W. E and B. Yu. The deep Ritz method: A deep learning-based numerical algorithm for solving variational problems. *Communications in Mathematics and Statistics*, 6(1):1–12, March 2018.
- [19] R. Fletcher and M. J. D. Powell. A rapidly convergent descent method for minimization. *The Computer Journal*, 6(2):163–168, 1963.
- [20] G. H. Golub and V. Pereyra. The differentiation of pseudo-inverses and nonlinear least squares problems whose variables separate. *SIAM Journal on Numerical Analysis*, 10(2):413–432, 1973.
- [21] D. L. B. Jupp. Approximation to data by splines with free knots. *SIAM Journal on Numerical Analysis*, 15(2):328–343, 1978.
- [22] I. E. Lagaris, A. C. Likas, and D. G. Papageorgiou. Neural-network methods for boundary value problems with irregular boundaries. *IEEE Trans. on Neural Networks*, 11(5):1041–1049, 2000.
- [23] 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, 1998.
- [24] M. Liu and Z. Cai. Adaptive two-layer ReLU neural network: II. Ritz approximation to elliptic pdes. *Computers & Mathematics with Applications*, 113:103–116, May 2022.
- [25] M. Liu, Z. Cai, and J. Chen. Adaptive two-layer ReLU neural network: I. best least-squares approximation. *Computers & Mathematics with Applications*, 113:34–44, May 2022.
- [26] P. D. Loach and A. J. Wathen. On the best least squares approximation of continuous functions using linear splines with free knots. *IMA Journal of Numerical Analysis*, 11(3):393–409, July 1991.
- [27] J. Nocedal and S. J. Wright. *Numerical Optimization*. Springer, New York, NY, USA, 2e edition, 2006.
- [28] M. R. Osborne. Some special nonlinear least squares problems. *SIAM Journal on Numerical Analysis*, 12:571–592, 1975.
- [29] 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.
- [30] L. Schumaker. Spline functions: Basic theory. *Wiley, New York*, 1981.
- [31] J. Sirignano and K. Spiliopoulos. DGM: A deep learning algorithm for solving partial differential equations. *Journal of Computational Physics*, 375:1339–1364, 2018.