

WAN DISCRETIZATION OF PDEs: BEST APPROXIMATION, STABILIZATION, AND ESSENTIAL BOUNDARY CONDITIONS*

SILVIA BERTOLUZZA[†], ERIK BURMAN[‡], AND CUIYU HE[§]

Abstract. In this paper, we provide a theoretical analysis of the recently introduced weakly adversarial networks (WAN) method, used to approximate partial differential equations in high dimensions. We address the existence and stability of the solution, as well as approximation bounds. We also propose two new stabilized WAN-based formulas that avoid the need for direct normalization. Furthermore, we analyze the method’s effectiveness for the Dirichlet boundary problem that employs the implicit representation of the geometry. We also devise a pseudotime XNODE neural network for static PDE problems, yielding significantly faster convergence results than the classical deep neural networks.

Key words. weak adversarial network, pseudotime XNODE, Cea’s lemma

MSC codes. 65M12, 65N12

DOI. 10.1137/23M1588196

1. Introduction. Recently there has been a vast interest in approximating partial differential equations (PDE) using neural networks and machine learning techniques. In this note, we will consider the weak adversarial networks (WAN) method introduced by Zang et al. [22]. The idea is to rewrite the weak form of the PDE as a saddle point problem whose solution is obtained by approximating both the trial (primal) and the test (adversarial) space through neural networks. In [22], the method was tested on various PDEs, tackling different challenging issues such as high dimension, nonlinearity, and nonconvexity of the domain. It was subsequently applied for the inverse problems in high dimension [1] and for the parabolic problems [16], with quite promising results.

However, as often happens for neural network methods for numerical PDEs, rigorous theoretical results on the capability of WANs to approximate the solution of a given PDE still need to be improved. The most critical issues that must be addressed are the discrete solution’s existence, stability, and approximation properties. Due to the inherent nature of neural network function classes, even the issue of the existence of a discrete solution is far from a trivial one. Indeed, fixed architecture neural network classes are generally neither convex nor closed [18, 14]. Therefore, a global minimum for a cost functional in one of such classes might not exist. Unsurprisingly, as we are ultimately dealing with a saddle point problem, a suitable choice of the test (adversarial) network class will play a vital role in the analysis. The lack of linearity of the trial (primal) network class will imply the need for a strengthened

*Submitted to the journal’s Machine Learning Methods for Scientific Computing section July 20, 2023; accepted for publication (in revised form) August 23, 2024; published electronically December 17, 2024.

<https://doi.org/10.1137/23M1588196>

Funding: The second author was partially funded by EPSRC grants EP/W007460/1, EP/V050400/1, and EP/T033126/1. The first author was partially funded by PRIN grant 20204LN5N5 AdPoly.

[†]Istituto di Matematica Applicata e Tecnologie Informatiche, CNR, 27100 Pavia, Italy (silvia.bertoluzza@imati.cnr.it).

[‡]Department of Mathematics, University College London, London, UK (e.burman@ucl.ac.uk).

[§]Department of Mathematics, Oklahoma State University, Stillwater, OK 74078 USA, and Department of Mathematics, University of Georgia, Athens, GA 30602 USA (cuiyu.he@okstate.edu, cuiyu.he@uga.edu).

inf-sup condition (see (2.14) in the following), which, however, will not, in general, be enough to guarantee the existence and uniqueness of a global minimizer. Indeed, due to the nonclosedness of neural network classes, it might not be possible to attain the minimum with an element belonging to the class.

What we can prove, under suitable assumptions (see (2.6) and (2.14)), in a general abstract framework, is (1) the existence of at least one weakly converging minimizing sequence for the WAN cost functional, and (2) that all weak limits of weakly converging minimizing sequences satisfy a quasi-optimality bound similar to Céa's lemma. More importantly, we further prove that a similar approximation bound will hold for the elements of the minimizing sequences sufficiently close to convergence. Combined with approximation bounds by the deep neural networks (DNNs) [10], this will guarantee that the WAN can, in principle, provide an arbitrarily good approximation to the continuous PDE solution. Another crucial issue relates to the convergence of the optimization scheme used to solve the minimization problem. Also this task is made difficult by the inherent topological properties of neural network classes. It is worth mentioning (see [18]) that the function class of DNNs lacks inverse stability in the L^p and $W^{s,p}$ norms. In simple terms, the norm of the elements of the DNN function class does not control the norm of the associated parameter vector. As the optimization schemes indirectly act on the function class through the parameter space, this will negatively affect the minimization process. In particular, when, the weak limit of the minimizing sequence does not belong to the function class, it can be proved that the sequence of the Euclidean norms of the corresponding parameter vectors explodes [18].

In the WAN framework, the aforementioned problems are integrated with the problems related to the inexact evaluation of the cost functional, which is defined as a supremum over the elements of the adversarial network and requires solving an optimization problem that, for the classical WAN method, becomes ill-posed due to the presence of direct normalization, and is therefore subject to a possibly relevant error. If this error becomes comparable, or even dominant, compared to the value of the cost functional itself, the overall optimization procedure will lose effectiveness and likely display oscillations. To mitigate this phenomenon, developing more stable and accurate methods for evaluating the operator norm is crucial. In the framework of WAN, we propose two alternative ways of evaluating the operator that avoid direct normalization and improve the overall convergence of the minimization procedure.

We then exploit the results for the second-order elliptic PDEs with essential boundary conditions. These are notoriously challenging as the construction of neural networks exactly vanishing on the boundary of a domain is extremely difficult, if not impossible. On the other hand, standard techniques, such as Nitsche's method, that impose Dirichlet boundary conditions weakly, rely on inverse inequalities that do not generally hold in the neural network framework. Adapting a strategy introduced, for finite elements, in [7], we propose to approximate the test space $H_0^1(\Omega)$ with a class of functions obtained by multiplying the elements of a given neural network class with a level set type weight, thus strongly enforcing the homogeneous boundary conditions on the test function class. Nonhomogeneous boundary conditions are then imposed by penalization with a suitable boundary norm. We can show that the resulting discrete schemes fall in our abstract setting, thus obtaining Céa's lemma type quasi-optimal H^1 error bounds.

As the architecture of neural networks plays a crucial role in their performance, we test the newly proposed methods on different function classes of various structures. In particular, besides DNN, we focus on residual-related networks, whose usage [11] was initially proposed to enhance image processing capabilities. These networks have also

found application in various domains, including numerical PDEs [13, 23]. In a recent work by Oliva et al. [16], the XNODE network is proposed to solve parabolic equations. Numerical experiments have demonstrated that, compared to classical DNN networks, XNODE can substantially reduce the number of iterations required for optimization. This rapid convergence can be attributed to the structure of the XNODE model, which emulates a residual network, and to the direct incorporation of the initial condition into the model. Besides testing our framework with XNODE architecture on a parabolic problem, we also introduce a new variant of the XNODE network, which we refer to as the pseudotime XNODE method for stationary problems. Remarkably fast convergence is observed in the numerical results, even for nonlinear and high-dimensional static elliptic PDEs.

The paper is organized as follows. In section 2, we prove quasi-best approximation results, and in section 3 we propose two more stable equivalent formulations. In section 4, we leverage our approach to allow for Dirichlet boundary conditions. Finally, the numerical results are provided in section 6. We devote the remaining part of this section to discussing the standard WAN in an abstract setting. Our framework covers a large class of problems without symmetry or coercivity assumptions, allowing for standard well-posed problems and certain nonstandard data assimilation problems. We also cover a very general class of discretization spaces: while we have in mind neural networks, the only a priori assumptions that we make on our trial and test function classes is that they are function sets containing the identically vanishing function so that our results potentially apply to a much wider range of methods, provided that the inf-sup conditions (2.6) and (2.14) hold.

Throughout the paper, we assume that all forms, linear and bilinear, are evaluated exactly and that the resulting nonlinear optimization problems can be solved with sufficient accuracy. Needless to say, these problems are crucial for the actual performance of the method. Nevertheless, the quasi-best approximation results proved herein are a cornerstone for its reliability.

1.1. The abstract setting. We consider a PDE set in some open, connected set $\Omega \subset \mathbb{R}^d$ ($d \geq 1$). We assume that the problem can be cast in the following general abstract weak form. Let W and V be two reflexive separable Banach spaces. Define a bounded bilinear form $\mathcal{A}: W \times V \mapsto \mathbb{R}$, satisfying

$$(1.1) \quad \mathcal{A}(w, v) \leq M \|w\|_W \|v\|_V \quad \forall w \in W, v \in V,$$

and let $\mathcal{F}: V \mapsto \mathbb{R}$ be a bounded linear form. We consider the abstract problem: find $u \in W$ such that

$$(1.2) \quad \mathcal{A}(u, v) = \mathcal{F}(v) \quad \forall v \in V.$$

As in [1], we rewrite (1.2) as the following minimization problem:

$$(1.3) \quad u = \operatorname{argmin}_{w \in W} \sup_{v \in V, v \neq 0} \frac{\mathcal{F}(v) - \mathcal{A}(w, v)}{\|v\|_V} = \operatorname{argmin}_{w \in W} \|u - w\|_{op},$$

where we define

$$(1.4) \quad \|w\|_{op} := \sup_{v \in V, v \neq 0} \frac{\mathcal{A}(w, v)}{\|v\|_V}.$$

We assume (1.2) admits a unique solution, satisfying the following stability estimate:

$$(1.5) \quad \|u\|_W \leq C \|\mathcal{F}\|_{V'}.$$

This is, for instance, the case if the form satisfies the assumptions of the Banach–Necas–Babuska theorem, or if it satisfies the more general condition of the Lions theorem, complemented by suitable compatibility conditions on $\mathcal{F}$ (see [8, Theorem 2.6 and Lemma A.40]). It is straightforward to show that, under such an assumption, the solution of problem (1.2) coincides with the unique minimizer of (1.3).

In principle, for any function class W_{θ} , parametrized by a parameter set $\mathcal{P}_{\theta}$, we can approximate the solution u by solving the semidiscrete problem:

$$(1.6) \quad \tilde{u}_{\theta}^* = \operatorname{argmin}_{w_{\theta} \in W_{\theta}} \|u - w_{\theta}\|_{op}.$$

Note that allowing the test space V to be different from the space W , where the solution is sought, makes the above formulation extremely flexible, allowing it to cover a wide range of situations, such as the ones where a partial differential equation, written in the form

$$(1.7) \quad a(u, w) = f(w) \quad \forall w \in W_0 \subset W$$

(W_0 denoting some closed subspace of W), is complemented by a constraint,

$$(1.8) \quad b(u, \chi) = g(\chi) \quad \forall \chi \in X,$$

where X is a third reflexive separable Banach space. Such a situation falls in our abstract framework, with $V = W_0 \times X$, if we set, for $v = (w_0, \chi) \in V$,

$$\mathcal{A}(w, v) = a(w, w_0) + \beta b(w, \chi), \quad \mathcal{F}(v) = f(w_0) + \beta g(\chi)$$

(β being a parameter weighting the constraint with respect to the equation). In such a case, the $\|\cdot\|_{op}$ norm satisfies

$$\|w\|_{op} = \sup_{(v, \chi) \in W_0 \times X} \frac{\mathcal{A}(w, (v, \chi))}{(\|v\|_{W_0}^2 + \|\chi\|_X^2)^{1/2}} \simeq \sup_{v \in W_0} \frac{a(w, v)}{\|v\|_{W_0}} + \beta \sup_{\chi \in X} \frac{b(w, \chi)}{\|\chi\|_X}.$$

Typically, as we shall see below, (1.8) could represent the imposition of essential boundary conditions. It could also represent some other form of constraint, such as the ones encountered in data assimilation problems subject to the heat or wave equation (see [2] or [3], where b is the L^2 -scalar product over some subset $\omega \subset \Omega$ [4, 15]).

1.2. The WAN method. Let W denote, throughout this section, the $H^1(\Omega)$ space with norm $\|v\|_W$ defined as $\|v\|_W^2 = (\nabla v, \nabla v)_{\Omega} + (v, v)_{\Omega}$, where $(\cdot, \cdot)_{\Omega}$ denotes the $L^2(\Omega)$ scalar product. We consider an elliptic partial differential equation, endowed with a Dirichlet boundary condition, that we write in the form

$$(1.9) \quad A(u) = f, \quad B(u) = g,$$

where A is a second-order partial differential operator and B is the trace operator. Bao et al. propose in [1] to rewrite (1.9) as a minimization problem in a suitable dual space. To this aim, the so-called operator norm is introduced, defined as

$$(1.10) \quad \|A(v)\|_{H^{-1}(\Omega), op} := \sup_{\substack{\varphi \in H_0^1(\Omega) \\ \varphi \neq 0}} \frac{a(v, \varphi)}{\|\varphi\|_W},$$

where $a : H^1(\Omega) \times H_0^1(\Omega) \mapsto \mathbb{R}$, $a(w, \varphi) = (A(w), \varphi)_{\Omega}$, is the bilinear form corresponding to the operator A . We immediately see that, provided the form a is continuous on

$H^1(\Omega) \times H_0^1(\Omega)$, such a norm is well defined (indeed, it coincides with the standard $H^{-1}(\Omega)$ norm). The idea of [1] was then to combine the residual in such a norm with a boundary penalization term aimed at weakly imposing the boundary conditions (rather than enforcing them exactly), and consider the following minimization problem:

$$(1.11) \quad u^* = \operatorname{argmin}_{w \in W} (\|A(u-w)\|_{H^{-1}(\Omega),op} + \beta \|g-w\|_{L^2(\partial\Omega)}).$$

Setting $V = H_0^1(\Omega) \times L^2(\partial\Omega)$, and

$$\mathcal{A}(w, [v, \chi]) = a(w, v) + \beta(w, \chi)_{\partial\Omega}, \quad \mathcal{F}([v, \chi]) = (f, v)_{\Omega} + \beta(g, \chi)_{\partial\Omega},$$

this problem can be rewritten in the form (1.3). At the continuous level, problem (1.11) is, in some sense, equivalent to (1.9). Indeed, we observe that the unique solution of (1.9) annihilates both $\|A(u-w)\|_{H^{-1}(\Omega),op}$ and $\|g-w\|_{L^2(\partial\Omega)}$, implying existence. Then, the value of the minimum is zero, and any other $H^1(\Omega)$ function minimizing the boundary penalized residual can be easily seen to be the solution to (1.9), thus obtaining uniqueness. Trivially, as it coincides with the solution of (1.9), the solution of (1.11) satisfies $\|u^*\|_W \lesssim \|f\|_{H^{-1}(\Omega)} + \|g\|_{H^{1/2}(\partial\Omega)} \lesssim \|f\|_{H^{-1}(\Omega)} + C(g)\|g\|_{L^2(\partial\Omega)}$, with $C(g) = \|g\|_{H^{1/2}(\partial\Omega)}/\|g\|_{L^2(\partial\Omega)}$, which is a stability bound of the form (1.5), though with a constant depending on g (whether such a constant is large or not depends on the frequency content of g : if g is not oscillating, such a constant is of order one, but it can be large if g presents high frequency oscillations). However, such a formulation does not entirely fall in the abstract setting of section 1.1, since, for $V = H_0^1(\Omega) \times L^2(\partial\Omega)$, the bilinear form $\mathcal{A}$ does not satisfy the boundedness assumption (1.1). It is therefore natural to consider the following minimization problem, where the boundary penalization term is measured in the $H^{1/2}(\partial\Omega) = (H^{-1/2}(\partial\Omega))'$ norm:

$$(1.12) \quad u^* = \operatorname{argmin}_{w \in W} (\|A(u-w)\|_{H^{-1}(\Omega),op} + \beta \|g-w\|_{H^{1/2}(\partial\Omega)}).$$

It is not difficult to see that also this problem can be written in the form (1.3), this time with $V = H_0^1(\Omega) \times H^{-1/2}(\partial\Omega)$. Thanks to the choice of the correct norm for the boundary penalization term, problem (1.12) falls within our abstract framework of subsection 1.1, it is well-posed, and it is equivalent to (1.9). It will serve as a starting point for the boundary condition treatment we will propose in section 4.

Remark 1.1. In the very first version of the WAN method (see [22]), the authors actually proposed a different definition of the operator norm, namely they defined the dual norm involved in the minimization problem as

$$\|A(v)\|_{L^2(\Omega),op} := \sup_{\varphi \in H_0^1(\Omega), \varphi \neq 0} \frac{a(v, \varphi)}{\|\varphi\|_{L^2(\Omega)}}.$$

It should be noted that this norm is not generally well defined at the continuous level, and to remedy this, the different normalization in (1.10) was proposed in [1]. We remark that the notation used for such a norm in [22] was $\|A(v)\|_{op}$, while we use the notation $\|\cdot\|_{op}$ with a different meaning; see (1.4).

Remark 1.2. We remark that replacing the natural norms $H^1(\Omega)$ and $H^{1/2}(\partial\Omega)$ in, respectively, (1.10) and (1.11) with the corresponding L^2 -norm results in two “variational crimes” with fairly different features. In both cases, the natural norm is replaced by a weaker norm, but in the first case the replacement happens in the

denominator. The resulting term $\|A(w)\|_{L^2(\Omega),op}$, $w \in W$, is not necessarily well defined (as this would require $w \in H^2(\Omega)$). This is essentially the same residual quantity as that minimized in so-called PINN (physically informed neural networks) methods [5, 19]. In the second case, the “variational crime” is somewhat less severe: all the quantities involved in the minimization problem (1.11) are well defined, though, as we already pointed out, the boundedness assumption (1.1) does not hold.

In the WAN method, the discretization for either (1.11) or (1.12) is performed by replacing the spaces $H^1(\Omega)$ and $H_0^1(\Omega)$ by, respectively, their discrete counterparts $W_\theta \subset H^1(\Omega)$ and $V_\eta \subset H_0^1(\Omega)$, where W_θ and V_η are two fixed architecture neural network function classes, parameterized by parameter sets $\mathcal{P}_\theta$ and $\mathcal{P}_\eta$. The discretization is carried out via a discrete operator norm, defined, for any $w \in H^1(\Omega)$, as

$$(1.13) \quad \|A(w)\|_{H^{-1}(\Omega),op,\eta} := \sup_{\substack{v_\eta \in V_\eta \\ \|v_\eta\|_V \neq 0}} \frac{a(w, v_\eta)}{\|v_\eta\|_V}.$$

The discrete method can then be written, for X being either $L^2(\partial\Omega)$ or $H^{1/2}(\partial\Omega)$,

$$(1.14) \quad u_\theta^* = \operatorname{argmin}_{w_\theta \in W_\theta} (\|A(u - w_\theta)\|_{H^{-1}(\Omega),op,\eta} + \beta \|w_\theta - g\|_X).$$

Exactly evaluating the functional on the right-hand side is very difficult since functions in W_θ and V_η may have very different geometric structures. In practice, the integrals are approximated using fixed sample points or a Monte Carlo integration method [12]. The optimization is then performed using a stochastic gradient descent method, e.g., Adam, over the parameter sets $\mathcal{P}_\theta$ and $\mathcal{P}_\eta$. We also note that, due to the normalization in (1.13), when w is close to u , the maximization problem in v_η becomes ill-posed, resulting in increased undesirable oscillations. We will propose a possible remedy in section 3.

2. Analysis of the WAN method. This section will frame and analyze the WAN method in an abstract framework. We aim to provide insight into choosing the approximation and adversarial networks to ensure the resulting method’s stability and optimality. For simplicity, we will perform the analysis based on (1.14) without considering the errors caused by the Monte Carlo and gradient descent methods.

We define the WAN method in the abstract framework as follows. Letting $V_\eta \subset V$ denote a function class parametrized by a parameter set $\mathcal{P}_\eta$, we introduce the discrete version of the $\|\cdot\|_{op}$ norm on W , defined as

$$(2.1) \quad \|w\|_{op,\eta} := \sup_{\substack{v_\eta \in V_\eta \\ \|v_\eta\|_V \neq 0}} \frac{\mathcal{A}(w, v_\eta)}{\|v_\eta\|_V}.$$

We observe that for all $w \in W$ we have that

$$(2.2) \quad \|w\|_{op,\eta} \leq \|w\|_{op} \leq M \|w\|_W.$$

The fully discrete problem then reads

$$(2.3) \quad u_\theta^* = \operatorname{argmin}_{w_\theta \in W_\theta} \|u - w_\theta\|_{op,\eta}.$$

In our analysis, a key role will be played by the function class of differences of elements of the approximation network W_θ :

$$(2.4) \quad S_\theta := \{w_{1,\theta} - w_{2,\theta}, w_{1,\theta}, w_{2,\theta} \in W_\theta\}.$$

We will first consider the case of coercive problems and then tackle problems only known to satisfy the stability (1.5).

2.1. Coercive problems. Let us at first consider the case $V = W$, and assume that the bilinear form $\mathcal{A}$ is coercive, i.e., there exist $\alpha > 0$ such that

$$(2.5) \quad \alpha \|\phi\|_W^2 \leq \mathcal{A}(\phi, \phi).$$

We make the following assumption on the networks W_θ and V_η :

$$(2.6) \quad W_\theta \cup S_\theta \subseteq V_\eta.$$

Observe that if $0 \in W_\theta$, we have that $W_\theta \cup S_\theta = S_\theta$. We start by remarking that, as the functional $w \rightarrow \|u - w\|_{op, \eta}$, with $u \in W$ given, is bounded from below by 0, we have that

$$\sigma^* := \inf_{w_\theta \in W_\theta} \|u - w_\theta\|_{op, \eta} \geq 0.$$

By the definition of infimum, there exist a sequence $\{w_\theta^n\}$ with $w_\theta^n \in W_\theta$ such that

$$(2.7) \quad \lim_{n \rightarrow \infty} \|u - w_\theta^n\|_{op, \eta} = \inf_{w_\theta \in W_\theta} \|u - w_\theta\|_{op, \eta}.$$

We call a sequence satisfying (2.7) a minimizing sequence for (2.3). We have the following lemma, where $\text{cl}_w^{seq}(W_\theta) \subseteq W$ denotes the weak sequential closure of W_θ in W (see [17]).

LEMMA 1. *Let $\{w_\theta^n\}$ be a minimizing sequence for (2.3). Then, under assumption (2.6), there exists a subsequence weakly converging to an element $u_\theta^* \in \text{cl}_w^{seq}(W_\theta)$ satisfying*

$$\|u - u_\theta^*\|_{op, \eta} \leq \inf_{w_\theta \in W_\theta} \|u - w_\theta\|_{op, \eta}.$$

Proof. Thanks to (2.5) and (2.6) it is not difficult to see that the sequence $\{w_\theta^n\}$ is bounded in W , and it therefore admits a weakly convergent subsequence $\{\tilde{w}_\theta^n\}$. We let $u_\theta^* \in W$ denote the weak limit of $\{\tilde{w}_\theta^n\}$. Let now $A^T : W \rightarrow W'$ be defined as $\langle A^T v, w \rangle = \mathcal{A}(w, v)$, with $\langle \cdot, \cdot \rangle$ denoting the duality pairing. We have, by the definition of weak limit,

$$(2.8) \quad \|u - u_\theta^*\|_{op, \eta} = \sup_{\substack{v_\eta \in V_\eta \\ v_\eta \neq 0}} \frac{\langle A^T v_\eta, u - u_\theta^* \rangle}{\|v_\eta\|_V} = \sup_{\substack{v_\eta \in V_\eta \\ v_\eta \neq 0}} \frac{\lim_{n \rightarrow \infty} \langle A^T v_\eta, u - \tilde{w}_\theta^n \rangle}{\|v_\eta\|_V}.$$

Now, for any $v_\eta \in V_\eta$, $v_\eta \neq 0$, we have

$$\frac{\lim_{n \rightarrow \infty} \langle A^T v_\eta, u - \tilde{w}_\theta^n \rangle}{\|v_\eta\|_V} \leq \lim_{n \rightarrow \infty} \sup_{\substack{v'_\eta \in V_\eta \\ v'_\eta \neq 0}} \frac{\mathcal{A}(u - \tilde{w}_\theta^n, v'_\eta)}{\|v'_\eta\|_V} = \lim_{n \rightarrow \infty} \|u - \tilde{w}_\theta^n\|_{op, \eta} = \sigma^*,$$

whence $\|u - u_\theta^*\|_{op, \eta} \leq \sigma^*$. $\square$

We now prove Cea's lemma of best approximation for WAN on coercive problems.

LEMMA 2. *Let assumption (2.6) hold, and let u be the solutions to (1.2) and $u_\theta^* \in \text{cl}_w^{seq}(W_\theta)$ be the weak limit of a weakly convergent minimizing sequence $\{\tilde{w}_\theta^n\}$ for (2.3). Then we have the following error bound:*

$$(2.9) \quad \|u - u_\theta^*\|_W \leq \left(1 + \frac{2M}{\alpha}\right) \inf_{w_\theta \in W_\theta} \|u - w_\theta\|_W.$$

Proof. We start by observing that (2.6) implies that for any two elements $w_{1,\theta}$ and $w_{2,\theta}$ of W_θ it holds that

$$(2.10) \quad \alpha \|w_{1,\theta} - w_{2,\theta}\|_W \leq \sup_{\substack{v_\eta \in V_\eta \\ v_\eta \neq 0}} \frac{\mathcal{A}(w_{1,\theta} - w_{2,\theta}, v_\eta)}{\|v_\eta\|_W} = \|w_{1,\theta} - w_{2,\theta}\|_{op,\eta}.$$

Let now $e^* = u - u_\theta^*$, and let w_θ be an arbitrary element in W_θ . Letting $(\cdot, \cdot)_W$ denote the scalar product in W and $\mathfrak{R}: W \rightarrow W'$ denote the Riesz isomorphism, we have

$$\begin{aligned} \|u_\theta^* - w_\theta\|_W &= \frac{(u_\theta^* - w_\theta, u_\theta^* - w_\theta)_W}{\|u_\theta^* - w_\theta\|_W} = \frac{\langle \mathfrak{R}(u_\theta^* - w_\theta), u_\theta^* - w_\theta \rangle}{\|u_\theta^* - w_\theta\|_W} \\ &= \lim_{n \rightarrow \infty} \frac{\langle \mathfrak{R}(u_\theta^* - w_\theta), \tilde{w}_\theta^n - w_\theta \rangle}{\|u_\theta^* - w_\theta\|_W} \leq \lim_{n \rightarrow \infty} \|\tilde{w}_\theta^n - w_\theta\|_W \leq \alpha^{-1} \lim_{n \rightarrow \infty} \|\tilde{w}_\theta^n - w_\theta\|_{op,\eta}. \end{aligned}$$

Note that we used (2.10) for the last bound. Adding and subtracting u in the right-hand side and using (2.3) and (1.1), we have

$$\begin{aligned} (2.11) \quad \|u_\theta^* - w_\theta\|_W &\leq \alpha^{-1} \lim_{n \rightarrow \infty} \|u - \tilde{w}_\theta^n\|_{op,\eta} + \alpha^{-1} \|u - w_\theta\|_{op,\eta} \\ &\leq \alpha^{-1} \inf_{w'_\theta \in W_\theta} \|u - w'_\theta\|_{op,\eta} + \alpha^{-1} \|u - w_\theta\|_{op,\eta} \leq \frac{2}{\alpha} \|u - w_\theta\|_{op,\eta}. \end{aligned}$$

Since $w_\theta \in W_\theta$ is arbitrary, using (2.2) and a triangle inequality we get (2.9). $\square$

Generally, the weak solutions to (1.14), defined as the weak limits of minimizing sequences for the right-hand side in W_θ , are not necessarily unique. Moreover, the solution of the minimization problem (2.3) itself might not lie in W_θ , but only in its weak sequential closure. In such a case, it can be proved (see [18]) that the sequence of parameters in $\mathcal{P}_\theta$ resulting from the minimization procedure is unbounded, which results in numerical instability. A possible remedy (see [1]) is to restrict both maximization in V_η and minimization in W_θ to subsets of V_η and W_θ corresponding to parameters in $\mathcal{P}_\eta$ and $\mathcal{P}_\theta$ with Euclidean norm bounded by a suitable constant B . In such a case, one can apply standard calculus results to prove the existence of a minimizer $w_\theta \in W_\theta$. However, finding an appropriate choice of B remains a challenging problem. A too-small value of B will result in poor approximation regardless of the network's approximation capability, and if B is very large, it ultimately serves no purpose. Lemma 2 does, instead, guarantee that even when multiple weak solutions exist, they all provide a quasi-best approximation of u in W . Moreover, we can obtain a quasi-best approximation to u within the approximation class W_θ by taking entries of any minimizing sequence sufficiently close to convergence. Indeed, for any minimizing sequence $\{\tilde{w}_\theta^n\}$, given $\varepsilon > 0$ we can choose k such that

$$\|u - \tilde{w}_\theta^k\|_{op,\eta} \leq \inf_{w_\theta \in W_\theta} \|u - w_\theta\|_{op,\eta} + \varepsilon.$$

Then, by (2.9) and (2.11) we have

$$\|u - \tilde{w}_\theta^k\|_W \leq \|u - u_\theta^*\|_W + \|u_\theta^* - \tilde{w}_\theta^k\|_W \lesssim \inf_{w_\theta \in W_\theta} \|u - w_\theta\|_W + \varepsilon,$$

meaning that any minimizing sequence does approximate the solution u in the norm $\|\cdot\|_W$ within the accuracy allowed by the chosen neural network class architecture in a finite number of steps. It is important to observe that, under proper assumptions, the cost functional is equivalent to the W' norm of the residual, thus providing a

reliable a posteriori error bound. Moreover, by (2.11), the cost functional evaluated on w_θ provides an upper bound for the discrepancy, in W , between w_θ and the weak limit w_θ^* , and can then be leveraged to devise a stopping criterion.

Remark 2.1. Since W_θ is a function class and not a function space, (2.6) implies that V_η should be a richer function class than W_θ . When $\mathcal{A}$ is coercive and symmetric, the Deep Ritz method can be interpreted as choosing, in our abstract formulation, $V_\eta = u - W_\theta$. It is not difficult to check that with such a definition of V_η , if $\mathcal{A}$ is coercive, both Lemma 1 and Lemma 2 still hold. However, in practice, numerical evidence suggests that using a separate and more comprehensive space for V_η than $u - W_\theta$ enhances both numerical efficiency (faster convergence) and accuracy.

Remark 2.2. To fully exploit (2.9), we combine it with approximation results on neural network classes. We refer to [10] for a survey of the different results available in the literature and to the references therein. In particular, we recall that when $W = H^1(\Omega)$ and W_θ is a function class of DNN network with ReLU activation function, it was shown in [9] that for any function $\varphi \in H^m(\Omega)$, $m > 1$ and Ω is Lipschitz,

$$(2.12) \quad \min_{\varphi_\theta \in V_\theta} \|\varphi - \varphi_\theta\|_{H^1(\Omega)} \leq C(m, d) N_\theta^{-(m-1)/d} \|\varphi\|_{H^m(\Omega)},$$

where $C(m, d) \geq 0$ is a function that depends on (m, d) and N_θ is the number of neurons in the DNN network. Combining such a bound with the quasi-best approximation estimates allows us to deduce a priori error estimates of the WAN schemes.

Remark 2.3. While we focused our analysis on linear problems, the WAN method can be, and is, applied also in the nonlinear framework. Indeed, under suitable assumptions on the operator A (for instance, if A is monotone and Lipschitz continuous) the existence of weakly converging minimizing sequences whose weak limit satisfies the estimate of Lemma 2 carries over to the nonlinear case. A proof in the case of monotone operators is given in the online supplementary material (SM_Algorithms.pdf [local/web 243KB], SM1.pdf [local/web 162KB]). Beyond that, also in cases where monotonicity does not hold, numerical results will show the effectiveness of our approach (see subsections 6.1 and 6.2).

2.2. PDE without coercivity. We now drop the assumption that $V = W$, and we assume instead that there exists an operator $\mathfrak{R}: V \rightarrow W'$ such that

$$(2.13) \quad \inf_{w \in W} \sup_{\substack{v \in V \\ v \neq 0}} \frac{\langle \mathfrak{R}v, w \rangle}{\|w\|_W \|v\|_V} \geq \alpha^* > 0, \quad \|\mathfrak{R}v\|_{W'} \leq M^* \|v\|_V.$$

Note that, as we assume that problem (1.2) is well-posed, a possible choice for $\mathfrak{R}$ is $\mathfrak{R} = A^T$, but choices with better stability constants α^* might exist. Moreover assume that $V_\eta \subset V$ can be chosen so that we have the discrete inf-sup condition:

$$(2.14) \quad \kappa \|w_\theta\|_W \leq \sup_{\substack{v_\eta \in V_\eta \\ v_\eta \neq 0}} \frac{\mathcal{A}(w_\theta, v_\eta)}{\|v_\eta\|_V} \quad \forall w_\theta \in W_\theta \cup S_\theta$$

with S_θ defined in (2.4). It is easy to see that Lemma 1 holds with proof unchanged also in this case, which gives us the existence of a (possibly not unique) element $u_\theta^* \in \text{cl}_w^{seq}(W_\theta)$, weak limit of a minimizing sequence $\{\tilde{w}_\theta^n\}$ of elements of W_θ , satisfying

$$\|u - u_\theta^*\|_{op, \eta} \leq \|u - w_\theta\|_{op, \eta} \quad \forall w_\theta \in W_\theta.$$

LEMMA 3. Let V_η be chosen in such a way that assumption (2.14) is satisfied for some constant $\kappa > 0$, possibly depending on V_η . Let u be the solutions to (1.2) and let $u_\theta^* \in \text{cl}_w^{\text{seq}}(W_\theta)$ be the weak limit of a weakly convergent minimizing sequence $\{\tilde{w}_\theta^n\}$ for (2.3). Then we have the following error bound:

$$(2.15) \quad \|u - u_\theta^*\|_W \leq \left(1 + 2 \frac{M^*}{\alpha_*} \frac{M}{\kappa}\right) \inf_{w_\theta \in W_\theta} \|u - w_\theta\|_W.$$

Proof. Let w_θ be an arbitrary element of W_θ . Thanks to (2.13) and (2.14) we can write

$$(2.16) \quad \begin{aligned} \alpha_* \|u_\theta^* - w_\theta\|_W &\leq \sup_{\substack{v \in V \\ v \neq 0}} \frac{\langle \Re v, u_\theta^* - w_\theta \rangle}{\|v\|_V} = \lim_{n \rightarrow \infty} \sup_{\substack{v \in V \\ v \neq 0}} \frac{\langle \Re v, \tilde{w}_\theta^n - w_\theta \rangle}{\|v\|_V} \\ &\leq M^* \lim_{n \rightarrow \infty} \|\tilde{w}_\theta^n - w_\theta\|_W \leq \frac{M^*}{\kappa} \lim_{n \rightarrow \infty} \|\tilde{w}_\theta^n - w_\theta\|_{op, \eta}. \end{aligned}$$

By the same argument used for the proof of Lemma 2, we then obtain that

$$(2.17) \quad \|u_\theta^* - w_\theta\|_W \leq 2 \frac{M^*}{\alpha_*} \frac{1}{\kappa} \|u - w_\theta\|_{op, \eta},$$

and, consequently, by the triangle inequality,

$$(2.18) \quad \|u - u_\theta^*\|_W \leq \left(1 + 2 \frac{M^*}{\alpha_*} \frac{M}{\kappa}\right) \|u - w_\theta\|_W,$$

which, thanks to the arbitrariness of w_θ , gives (2.15). $\square$

Like the coercive case, we can have an almost best approximation in a finite number of steps of any weakly converging minimizing sequence $\{\tilde{w}_\theta^n\}$. More precisely, by the same argument as for the coercive case, for all $\varepsilon > 0$ there exists a k such that

$$\|u - \tilde{w}_\theta^k\|_W \lesssim \inf_{w_\theta \in W_\theta} \|u - w_\theta\|_W + \varepsilon.$$

We conclude this section by the following observation: let $\mathcal{J}(\cdot)$ denote any functional on W equivalent to the $\|\cdot\|_{op, \eta}$ norm,

$$(2.19) \quad c_* \|w\|_{op, \eta} \leq \mathcal{J}(w) \leq C^* \|w\|_{op, \eta} \quad \forall w \in W,$$

and consider the problem

$$(2.20) \quad u_\theta^* = \operatorname{argmin}_{w_\theta \in W_\theta} \mathcal{J}(u - w_\theta).$$

Then there exists a possibly not unique $w_\theta^b \in \text{cl}_w^{\text{seq}}(W_\theta)$ such that

$$\mathcal{J}(u - w_\theta^b) \leq \inf_{w_\theta \in W_\theta} \mathcal{J}(u - w_\theta).$$

Moreover for all w_θ^b such that u_θ^b is the weak limit of a minimizing sequence $\{\tilde{w}_\theta^n\}$ for (2.20), it holds that

$$\|u - u_\theta^b\|_W \leq \left(1 + 2 \frac{C^*}{c_*} \frac{M^*}{\alpha_*} \frac{M}{\kappa}\right) \inf_{w_\theta \in W_\theta} \|u - w_\theta\|_W.$$

Indeed, by (2.19), all minimizing sequences are bounded with respect to the $\|\cdot\|_{op, \eta}$ norm and, hence, with respect to the $\|\cdot\|_W$ norm. Any minimizing sequence does

then weakly converge to an element u_{θ}^b . Moreover, initially proceeding as in (2.16), thanks to (2.19), we have, for w_{θ} arbitrary,

$$\begin{aligned} \alpha_* \|u_{\theta}^* - w_{\theta}\|_W &\leq \frac{M^*}{\kappa} \lim_{n \rightarrow \infty} \|\tilde{w}_{\theta}^n - w_{\theta}\|_{op, \eta} \\ &\leq \frac{M^*}{\kappa} \left(\lim_{n \rightarrow \infty} \|\tilde{w}_{\theta}^n - u\|_{op, \eta} + \|u - w_{\theta}\|_{op, \eta} \right) \\ &\leq \frac{M^*}{\kappa C_*} \left(\lim_{n \rightarrow \infty} \mathcal{J}(\tilde{w}_{\theta}^n - u) + \mathcal{J}(u - w_{\theta}) \right) \leq \frac{2MM^*C^*}{\kappa C_*} \|u - w_{\theta}\|_W. \end{aligned}$$

3. Two novel stabilized loss functions. To mitigate undesirable oscillations during the optimization procedure, resulting from the inexact solution of the maximization problem involved in the definition of the operator norm, we propose two alternative definitions of the cost functional that yields the same minimum, while avoiding direct normalization. More precisely, we introduce two new functionals on the product space $W \times V$, such that the supremum over $v \in V$, for $w \in W$ fixed, also gives, up to a possible rescaling and translation, the operator norm of w , while yielding a more favorable optimization problem under discretization.

3.1. Stabilized WAN method. Define

$$(3.1) \quad |||w|||_{op}^2 = \sup_{v \in V} \left(\mathcal{A}(w, v) - \frac{\gamma_d}{2} \|v\|_V^2 \right), \quad |||w|||_{op, \eta}^2 = \sup_{v_{\eta} \in V_{\eta}} \left(\mathcal{A}(w, v_{\eta}) - \frac{\gamma_d}{2} \|v_{\eta}\|_V^2 \right),$$

where $\gamma_d > 0$ is a constant, and consider the following problem:

$$(3.2) \quad u_{\theta}^{\sharp} = \operatorname{argmin}_{w_{\theta} \in W_{\theta}} |||u - w_{\theta}|||_{op, \eta}.$$

The following lemma shows that the norms defined in (3.1) coincide with the operator norms defined in (1.4) and (2.1), up to a constant dependent on γ_d .

LEMMA 4. Assume that $v_{\eta} \in V_{\eta}$ implies $\lambda v_{\eta} \in V_{\eta}$ for all $\lambda \in \mathbb{R}^+$. Then, for any $w \in W$, there holds

$$(3.3) \quad \frac{1}{2\gamma_d} \|w\|_{op}^2 = |||w|||_{op}^2, \quad \frac{1}{2\gamma_d} \|w\|_{op, \eta}^2 = |||w|||_{op, \eta}^2.$$

Remark 3.1. From this lemma, it seems reasonable to choose, e.g., $\gamma_d = 1$. However, experiments have shown that adjusting the value of γ can effectively control the oscillations in experiments.

Proof. We prove the second of the two equalities. The first can be proved by the same argument. For any fixed $w \in W_{\theta}$ with $w \neq 0$, and for all $\varepsilon > 0$, there exists a $\bar{\varphi}_w^{\varepsilon} \in V_{\eta}$ (depending on ε) with $\|\bar{\varphi}_w^{\varepsilon}\|_V = 1$, such that

$$\mathcal{A}(w, \bar{\varphi}_w^{\varepsilon}) \geq (1 - \varepsilon) \|w\|_{op, \eta},$$

which, setting $\varphi_w = \gamma_d^{-1} \|w\|_{op, \eta} \bar{\varphi}_w^{\varepsilon}$, yields

$$\begin{aligned} (3.4) \quad \sup_{\varphi_{\eta} \in V_{\eta}} \left(\mathcal{A}(w, \varphi_{\eta}) - \frac{\gamma_d}{2} \|\varphi_{\eta}\|_V^2 \right) &\geq \mathcal{A}(w, \varphi_w) - \frac{\gamma_d}{2} \|\varphi_w\|_V^2 \\ &= \gamma_d^{-1} \|w\|_{op, \eta} \mathcal{A}(w, \bar{\varphi}_w^{\varepsilon}) - \frac{1}{2\gamma_d} \|w\|_{op, \eta}^2 \geq \left(\frac{1}{2} - \varepsilon \right) \frac{1}{\gamma_d} \|w\|_{op, \eta}^2. \end{aligned}$$

By the arbitrariness of ε we then obtain that $|||w|||_{op,\eta}^2 \geq \|w\|_{op,\eta}^2/(2\gamma_d)$. To prove the converse inequality, using Young's equality gives

$$\sup_{\varphi_\eta \in V_\eta} \left(\mathcal{A}(w, \varphi_\eta) - \frac{\gamma_d}{2} \|\varphi_\eta\|_V^2 \right) \leq \sup_{\varphi_\eta \in V_\eta} \left(\|w\|_{op,\eta} \|\varphi_\eta\|_V - \frac{\gamma_d}{2} \|\varphi_\eta\|_V^2 \right) \leq \frac{1}{2\gamma_d} \|w\|_{op,\eta}^2.$$

The first part of (3.3) can be proved in a similar way. $\square$

The analysis of the minimization problem (2.3) then carries over to the minimization problem (3.2). Then, there exists at least a minimizing sequence in W_θ weakly converging to a limit $u_\theta^\sharp \in \text{cl}_w^{seq}(W_\theta)$, and all weak limits of minimizing sequences satisfy either bound (2.9) or bound (2.15), depending on whether $\mathcal{A}$ is coercive or not.

Remark 3.2. When w approaches the true solution, we have that

$$v^*(w) = \operatorname{argmax}_{v \in V} \left(\mathcal{A}(u - w, v) - \frac{\gamma_d}{2} \|v\|_V^2 \right) \rightarrow 0.$$

As a consequence, depending on how large the space W_θ is, the problem

$$u_\theta^* = \operatorname{argmin}_{w_\theta \in W_\theta} \left(\mathcal{A}(u - w_\theta, v^*(w_\theta)) - \frac{\gamma_d}{2} \|v^*(w_\theta)\|_V^2 \right)$$

might be close to the problem $u_\theta^* = \operatorname{argmin}_{w_\theta \in W_\theta} \mathcal{A}(u - w_\theta, v^*(w_\theta))$, which, in turn, if $\|u - w_\theta\|_W$ is small, might be ill-posed and too sensitive to the errors in evaluating $v^*(w_\theta)$. The following subsection introduces a further stabilized loss function that mitigates this issue.

3.2. A further stabilized WAN method. We now define the alternative operator norm $||| \cdot |||^+$ as

$$(3.5) \quad \begin{aligned} (|||w|||_{op}^+)^2 &:= \sup_{v \in V} \left(\mathcal{A}(w, v) - \frac{\gamma_d}{2} \|v\|_V^2 + \|v\|_V \right), \\ (|||w|||_{op,\eta}^+)^2 &:= \sup_{v_\eta \in V_\eta} \left(\mathcal{A}(w, v_\eta) - \frac{\gamma_d}{2} \|v_\eta\|_V^2 + \|v_\eta\|_V \right), \end{aligned}$$

where $\gamma_d > 0$ is a constant, and we define the following minimization problem:

$$(3.6) \quad u_\theta^\dagger = \operatorname{argmin}_{w_\theta \in W_\theta} |||u - w_\theta|||_{op,\eta}^+.$$

The following lemma states the relation between the norms defined in (3.5) and the operator norms.

LEMMA 5. Assume that $v_\eta \in V_\eta$ implies $\lambda v_\eta \in V_\eta$ for all $\lambda \in \mathbb{R}^+$. For any $w \in W$, there holds

$$(3.7) \quad \frac{1}{\sqrt{2\gamma_d}} (\|w\|_{op} + 1) = |||w|||_{op}^+, \quad \frac{1}{\sqrt{2\gamma_d}} (\|w\|_{op,\eta} + 1) = |||w|||_{op,\eta}^+.$$

Proof. We prove the second of the two equalities; the first can be proved by the same argument. For any fixed $w \in W$ with $w \neq 0$, and for all $\varepsilon > 0$, there exists $\bar{\varphi}_w^\varepsilon \in V_\eta$ with $\|\bar{\varphi}_w^\varepsilon\|_V = 1$ that satisfies

$$(3.8) \quad \mathcal{A}(w, \bar{\varphi}_w^\varepsilon) \geq (1 - \varepsilon) \|w\|_{op,\eta},$$

which, setting $\varphi_w = \gamma_d^{-1} s \|w\|_{op, \eta} \bar{\varphi}_w^\varepsilon$, for some $s > 0$ to be chosen, yields

$$\begin{aligned} \sup_{\varphi_\eta \in V_\eta} \left(\mathcal{A}(w, \varphi_\eta) - \frac{\gamma_d}{2} \|\varphi_\eta\|_V^2 + \|\varphi_\eta\|_V \right) &\geq \mathcal{A}(w, \varphi_w) - \frac{\gamma_d}{2} \|\varphi_w\|_V^2 + \|\varphi_w\|_V \\ &= \gamma_d^{-1} s \|w\|_{op, \eta} \mathcal{A}(w, \bar{\varphi}_w^\varepsilon) - \frac{s^2}{2\gamma_d} \|w\|_{op, \eta}^2 + \gamma_d^{-1} s \|w\|_{op, \eta} \\ &\geq (1 - \varepsilon) \frac{s}{\gamma_d} \|w\|_{op, \eta}^2 - \frac{s^2}{2\gamma_d} \|w\|_{op, \eta}^2 + \frac{s}{\gamma_d} \|w\|_{op, \eta}. \end{aligned}$$

We can choose s that maximizes the term on the right-hand side. By direct computations, we have

$$\max_{s \in \mathbb{R}^+} \left((1 - \varepsilon) \frac{s}{\gamma_d} \|w\|_{op, \eta}^2 - \frac{s^2}{2\gamma_d} \|w\|_{op, \eta}^2 + \frac{s}{\gamma_d} \|w\|_{op, \eta} \right) = \frac{1}{2\gamma_d} ((1 - \varepsilon) \|w\|_{op, \eta} + 1)^2.$$

Above we have used the fact that $\max_{s \in \mathbb{R}^+} (-as^2 + bs) = b^2/4a$. By the arbitrariness of ε we then obtain $(\|w\|_{op, \eta}^+)^2 \geq (\|w\|_{op, \eta} + 1)^2/(2\gamma_d)$. The converse inequality can be proved by once again using Young's inequality. $\square$

Again, the analysis of the minimization problem (2.3), including the existence of a weakly converging minimization sequence and a best approximation bound for all weak limits of minimizing sequences, carries over to the minimization problem (3.2).

4. Imposition of Dirichlet boundary conditions. Herein we will adapt to the case of WANs the technique introduced in [7] for dealing with Dirichlet boundary conditions. The idea is to weigh the elements of the test adversarial network by multiplying a cutoff function ϕ (see [21] and references therein), so that the resulting test functions are forced to be zero on the boundary. The Dirichlet boundary condition can then be imposed on the primal network using a penalty without violating the consistency of the equation. For simplicity, we assume that the problem is a symmetric second-order static elliptic PDE. We also assume that the boundary $\partial\Omega$ is smooth (C^3 to be precise).

We first present the ideas in the simple framework of the Deep Ritz method for the case of homogeneous boundary conditions. We assume the operator $\mathcal{A}$ of (1.2) to be symmetric under homogeneous Dirichlet boundary conditions. Then the continuous Ritz method may be written as

$$u = \operatorname{argmin}_{v \in H_0^1(\Omega)} (0.5\mathcal{A}(v, v) - (f, v)_\Omega).$$

Under the smoothness assumptions on $\partial\Omega$ we know that, provided f is sufficiently smooth, $u \in H^m(\Omega)$ for some $m \geq 3$ and $\|u\|_{H^3(\Omega)} \lesssim \|f\|_{H^1(\Omega)}$. Assuming that $w_\theta \in H_0^1(\Omega)$ for all $w_\theta \in W_\theta$, the Deep Ritz method takes the form

$$u_\theta^* = \operatorname{argmin}_{w_\theta \in W_\theta} (0.5\mathcal{A}(w_\theta, w_\theta) - f(w_\theta)).$$

As we already mentioned, the problem with this formulation is that it appears to be very difficult to design networks that satisfy boundary conditions by construction. Instead, typically, a penalty term of the form $\lambda \|Tw_\theta\|_{\partial\Omega}$ is added to the functional on the right-hand side [6]. The convergence to the solution $u \in H_0^1(\Omega)$ of the continuous problem is obtained by letting $\lambda \rightarrow \infty$ and enriching the network space. In the classical numerical methods, e.g., the finite element method, λ is proportional to h^{-s} , with h being the mesh size and $s > 0$ a carefully chosen exponent. With neural network

methods, it is, however, not obvious how to match the dimension of the space to the rate by which λ grows. In general, either the accuracy or the conditioning of the nonlinear system suffers.

Our idea is to build the boundary conditions into the formulation by weighting the network functions with the level set function ϕ , where $\phi|_{\partial\Omega} = 0$, $\phi|_{\Omega} > 0$, and ϕ behaves as a distance function in the vicinity of $\partial\Omega$. The solution we look for then takes the form ϕw_{θ} with $w_{\theta} \in W_{\theta}$. The Cut Deep Ritz method reads

$$\nu_{\theta}^* = \operatorname{argmin}_{w_{\theta} \in W_{\theta}} (0.5\mathcal{A}(\phi w_{\theta}, \phi w_{\theta}) - f(\phi w_{\theta})).$$

It is straightforward to show that this is equivalent to

(4.1)

$$\nu_{\theta}^* = \operatorname{argmin}_{w_{\theta} \in W_{\theta}} (\mathcal{A}(u - \phi w_{\theta}, u - \phi w_{\theta})) = \operatorname{argmin}_{w_{\theta} \in W_{\theta}} \left(\sup_{v_{\theta} \in W_{\theta}} \frac{\mathcal{A}(u - \phi w_{\theta}, u - \phi v_{\theta})}{\sqrt{\mathcal{A}(u - \phi v_{\theta}, u - \phi v_{\theta})}} \right).$$

Following Remark 2.1, and assuming, for the sake of simplicity, that the minimization problem (4.1) has a unique solution $\nu_{\theta}^* \in W_{\theta}$, we then have that

$$\|u - \phi \nu_{\theta}^*\|_{H^1(\Omega)} \leq C \inf_{w_{\theta} \in W_{\theta}} \|u - \phi w_{\theta}\|_{H^1(\Omega)}.$$

It remains to show that ϕw_{θ} is capable of approximating u in $H_0^1(\Omega)$. To this end let $\mathcal{O}$ be some domain such that $\Omega \subset \mathcal{O}$, where $\mathcal{O}$ is a box in $\mathbb{R}^d$ and let $\tilde{u}$ denote a stable extension of u to $\mathcal{O}$ [20]. We assume the following on the boundary $\partial\Omega$ and ϕ .

Assumption 4.1. Let Ω be a bounded domain in $\mathbb{R}^d$. The boundary $\partial\Omega$ can be covered by open sets $\mathcal{O}_i, i = 1, \dots, I$, and one can introduce on every $\mathcal{O}_i$ local coordinates $\xi_1, \dots, \xi_d$ with $\xi_d = \phi$ such that all the partial derivatives $\partial \xi_i^{\alpha} / \partial^{\alpha} x$ and $\partial x^{\alpha} / \partial^{\alpha} \xi$ up to order $k+1$ are bounded by some $C_0 > 0$. Moreover, ϕ is of class C^{k+1} on $\mathcal{O}$, where $k+1 \geq 3$ is the smoothness of the domain, and $|\phi| \geq M_0$ on $\mathcal{O} \setminus \cup \mathcal{O}_i$ with some $m > 0$, and in $\cup \mathcal{O}_i$, ϕ is a signed distance function to $\partial\Omega$.

We further need the following Hardy type inequality (see [7, Lemma 3.1]).

LEMMA 6. *We assume that the domain Ω is defined by the zero level set of the smooth function ϕ and that Assumption 4.1 is satisfied. Then for any $v \in H^{k+1}(\mathcal{O})$ such that $v|_{\partial\Omega} = 0$, there holds*

$$\|v/\phi\|_{H^k(\Omega)} \leq C \|v\|_{H^{k+1}(\mathcal{O})}.$$

Then, as by assumption $\|\phi\|_{W^{1,\infty}(\Omega)} < C$, combining the quasi-best approximation bound given by Lemma 2 with (2.12) we obtain the following estimate for the Cut Deep Ritz method, with ReLU activation function:

$$\|u - \phi \nu_{\theta}^*\|_{H^1(\Omega)} \lesssim N_{\theta}^{-(m-2)/d} \|u\|_{H^m}.$$

In particular, when $m = 3$, using elliptic regularity, we can bound the $H^1(\Omega)$ norm of the error with $N_{\theta}^{-1/d} \|f\|_{H^1(\Omega)}$. This shows that the method typically requires one order more regularity of the data than typically expected.

We now introduce the cut weak adversarial network (CutWAN) method for problems not necessarily coercive and with nonhomogeneous Dirichlet boundary conditions. We let

$$(4.2) \quad |||w|||_{op,\phi,\eta} := \sup_{\varphi_{\eta} \in V_{\eta}} \left(\mathcal{A}(w, \phi \varphi_{\eta}) - \frac{\gamma_d}{2} \|\phi \varphi_{\eta}\|_{H^1(\Omega)}^2 \right)$$

and set

$$(4.3) \quad u_{\theta}^{\diamond} = \operatorname{argmin}_{w_{\theta} \in W_{\theta}} (|||u - w_{\theta}|||_{op, \phi, \eta} + \|w_{\theta} - g\|_{H^{1/2}(\partial\Omega)}).$$

Note that the cutoff function ϕ only multiplies the test functions in the CutWAN network method. The boundary condition is weakly imposed by adding a penalty term on the primal network $w_{\theta}|_{\partial\Omega}$. It is not difficult to check that this problem falls in the abstract formulation considered at the end of subsection 2.2. Here, the role of adversarial test network is played by the product space $\phi V_{\theta} \times H^{-1/2}(\partial\Omega)$, and the inf-sup condition (2.14) becomes

$$(4.4) \quad \|w\|_W \lesssim |||w|||_{op, \phi, \eta} + \|w\|_{H^{1/2}(\partial\Omega)} \quad \forall w \in S_{\theta}.$$

In particular, under such an assumption, we have the following best approximation results for the CutWAN method.

LEMMA 7. *Assume that the inf-sup condition (4.4) holds. Let $u_{\theta}^{\diamond}$ be the weak limit of a minimizing sequence for problem (4.3). Then there holds*

$$(4.5) \quad \|u - u_{\theta}^{\diamond}\|_W \lesssim \inf_{w_{\theta} \in W_{\theta}} \|u - w_{\theta}\|_W.$$

Note that the CutWAN method achieves optimal convergence rates even though the test function class is multiplied by ϕ . So the difficulty handled by the Hardy inequality in the Cut Deep Ritz method does not appear. Indeed the difficulty of controlling the level set weighted test function is hidden in the inf-sup assumption (4.4). A study of this condition will be the topic of future work.

Similarly, we can also define the following algorithm. Define

$$(4.6) \quad |||w|||_{op, \phi, \eta}^+ := \sup_{\varphi_{\eta} \in V_{\eta}} \left(\mathcal{A}(w, \phi \varphi_{\eta}) - \frac{\gamma_d}{2} \|\phi \varphi_{\eta}\|_V^2 + \|\phi \varphi_{\eta}\|_V \right),$$

and let

$$(4.7) \quad u_{\theta}^{\diamond} = \operatorname{argmin}_{w_{\theta} \in W_{\theta}} (|||u - w_{\theta}|||_{op, \phi, \eta}^+ + \|w_{\theta} - g\|_{H^{1/2}(\partial\Omega)}).$$

We refer to the above method as the shifted CutWAN method. One can also prove the best approximation results for the shifted CutWAN method similarly as in Lemma 7.

Remark 4.2. For computational convenience, the $H^{1/2}(\partial\Omega)$ norm in (4.3) and (4.7) can be replaced by a suitable combination of the L^2 norm of the function and of its tangential derivative. Indeed, using the Gagliardo–Nirenberg inequality we have that

$$(4.8) \quad \|w_{\theta} - g\|_{H^{1/2}(\partial\Omega)} \leq \|w_{\theta} - g\|_{L^2(\partial\Omega)}^{1/2} \|w_{\theta} - g\|_{H^1(\partial\Omega)}^{1/2}.$$

It is not difficult to ascertain that if we replace the $H^{1/2}$ norm in (4.3) and (4.7) with the right-hand side of (4.8), the analysis of section 2 holds with minor changes, resulting in an error bound of the form

$$\|u - u_{\theta}^{\diamond}\| \lesssim \inf_{w_{\theta} \in W_{\theta}} \left(\|u - w_{\theta}\|_W + \|u - w_{\theta}\|_{L^2(\partial\Omega)}^{1/2} \|u - w_{\theta}\|_{H^1(\partial\Omega)}^{1/2} \right).$$

We observe that an analogous result would hold for the plain $L^2(\partial\Omega)$ penalization as originally proposed by [22, 1] if an inverse estimate were to hold for W_{θ} , allowing

to bound the $H^1(\partial\Omega)$ norm of the boundary residual with its $L^2(\partial\Omega)$ norm times a constant depending on W_{θ} . Unfortunately this is generally not true, and Lemma 2 does not hold when using such a stabilization, which is, however, computationally convenient, and which we will test extensively in the forthcoming sections. We point out that, thanks to the combination of (4.8) with a Cauchy–Schwarz inequality, simply adding $\|g - w_{\theta}\|_{H^1(\partial\Omega)}$ to the $L^2(\partial\Omega)$ penalized functional yields an a posteriori error estimator. This can be evaluated upon convergence of the optimization procedure, to check if the solution obtained with the cheaper L^2 penalized functional is satisfactory. If not, it can serve, in a two-stage strategy, as the starting point for an additional optimization procedure relying on the more expensive functionals for which our theoretical error analysis applies.

5. Neural network structures.

5.1. Deep neural network structure. A DNN structure is the composition of multiple linear functions and nonlinear activation functions. We will use the DNN structure for V_{η} . Specifically, the first component of DNN is a linear transformation $T^l: \mathbb{R}^{n_l} \rightarrow \mathbb{R}^{n_{l+1}}$, $l = 1, \dots, L$, defined as follows:

$$T^l(\mathbf{x}^l) = \mathbf{W}^l \mathbf{x}^l + \mathbf{b}^l \text{ for } \mathbf{x}^l \in \mathbb{R}^{n_l},$$

where $\mathbf{W}^l = (w_{i,j}^l) \in \mathbb{R}^{n_{l+1} \times n_l}$ and $\mathbf{b}^l \in \mathbb{R}^{n_{l+1}}$ are parameters in the DNN. The second component is an activation function $\psi: \mathbb{R} \rightarrow \mathbb{R}$ to be chosen, and typical examples of the activation functions are tanh, Sigmoid, and ReLU. Application of ψ to a vector $\mathbf{x} \in \mathbb{R}^n$ is defined componentwise, i.e., $\psi(\mathbf{x}) = (\psi(x_i))$, $i = 1, 2, \dots, n$. The l th layer of the DNN is defined as the composition of the linear transform T^l and the nonlinear activation function ψ , i.e.,

$$\mathcal{N}^l(\mathbf{x}^l) := \psi(T^l(\mathbf{x}^l)), \quad l = 1, \dots, L-1.$$

For an input $\mathbf{x} \in \mathbb{R}^{n_1}$, a general L -layer DNN is defined as follows:

$$(5.1) \quad \mathcal{NN}(\mathbf{x}; \boldsymbol{\theta}) := T^L \circ \mathcal{N}^{L-1} \circ \dots \circ \mathcal{N}^2 \circ \mathcal{N}^1(\mathbf{x}),$$

where $\boldsymbol{\theta} \in \mathbb{R}^N$ stands for all the parameters in the DNN, i.e., $\boldsymbol{\theta} = \{\mathbf{W}^l, \mathbf{b}^l\}_{l=1}^L$. For a fully connected DNN, the number of parameters corresponding to $\boldsymbol{\theta}$ is $N_{\boldsymbol{\theta}} := \sum_{l=1}^L n_{l+1}(n_l + 1)$. We will refer to $\mathcal{N}^1$ as the input layer, $\mathcal{N}^i$, $1 < i < L$, as the hidden layers, and T^L as the output layer. We assume that every DNN neural network has an input, an output, and at least one hidden layer. Note that for the outer layer, there is no followed activation function. Figure 1 shows an example of a DNN model with five hidden layers with $[n_1, n_2, \dots, n_7] = [6, 20, 10, 10, 10, 20, 1]$.

5.1.1. The recursive DNN model. In the case of a DNN model with consecutive hidden layers having an equal number of neurons, the weights and biases for those hidden layers can be easily shared due to the same data structure. We define the recursive DNN model as DNN models that share the parameters for all consecutive hidden layers with the same number of neurons. Therefore, A recursive DNN model could have significantly fewer total parameters than the corresponding nonrecursive DNN model. For instance, the nonrecursive DNN model described in Figure 1 has 931 parameters, while its corresponding recursive model has only 711 parameters. The contrast will become more pronounced when the number of hidden layers and hidden neurons increases. Our numerical results show that a recursive DNN model can benefit PDE solving.

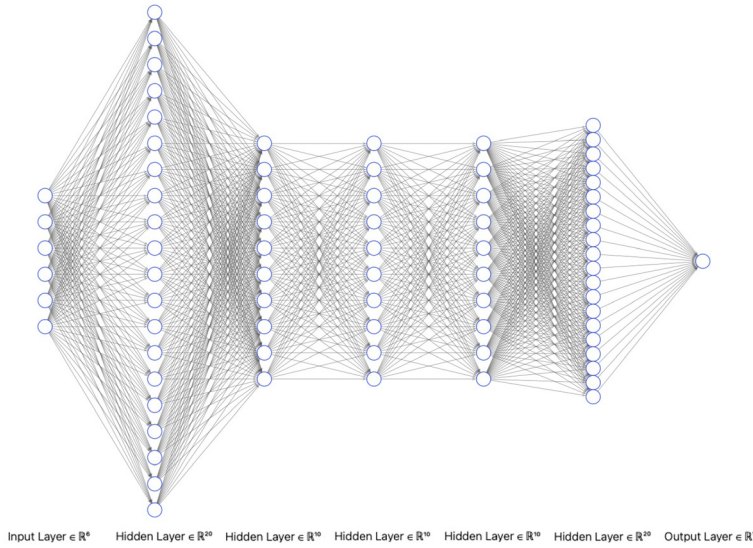


FIG. 1. A DNN network structure with five hidden layers.

5.1.2. Comments about DNN. Although DNNs have been widely used as the primary neural network for solving PDE problems, their performance often falls short of expectations. When using DNNs within the PINN and Deep Ritz methods, achieving the desired accuracy typically requires thousands of iterations due to oscillations and stagnation. The method of WAN helps the algorithm escape local minima. However, despite this improvement, the number of iterations remains in the range of several thousand, as reported in [22] and demonstrated in our numerical results in section 6. To enhance convergence, we explore different neural structures that approximate the trial functions with more efficacy.

5.2. XNODE model for parabolic PDE. It has been demonstrated in [16] that for time-dependent parabolic problems, the XNODE model achieves much faster convergence than traditional DNNs. We believe this rapid convergence is attributed to the structure of the XNODE model, which emulates the residual network, and the direct embedding of the initial condition in the model.

Consider the following parabolic PDE defined on an arbitrary bounded domain $\mathcal{D} \subset [0, T] \times \mathbb{R}^d$, possibly representing a time-dependent spatial domain:

$$(5.2) \quad \begin{cases} \partial_t u - \nabla \cdot A(t, \mathbf{x}) \nabla u + \mathbf{b}(t, \mathbf{x}) \nabla u + c(u, \mathbf{x})u - f(\mathbf{x}) = 0 & \text{for } (t, \mathbf{x}) \in \mathcal{D}, \\ u(t, \mathbf{x}) = g(t, \mathbf{x}) & \text{on } \partial\mathcal{D}, \\ u(0, \mathbf{x}) - h(\mathbf{x}) = 0 & \text{on } \Omega(0), \end{cases}$$

where $A = \{a_{ij}\}$, $\mathbf{b} = \{b_1, b_2, \dots, b_n\}$, $f : \mathcal{D} \rightarrow \mathbb{R}$, $c : \mathbb{R} \times \mathcal{D} \rightarrow \mathbb{R}$, and $h : \Omega(0) \rightarrow \mathbb{R}$ are given, with $\Omega(t) := \{\mathbf{x} | (t, \mathbf{x}) \in \mathcal{D}\}$ denoting the spatial domain of $\mathcal{D}$ when restricting time to be t . Note that c can be a nonlinear function with respect to the first argument.

We now briefly introduce the XNODE model in [16]. For simplicity, we consider a time-independent domain in this paper, i.e., $\mathcal{D} = [0, T] \times \Omega$, where $\Omega \subset \mathbb{R}^d$ is bounded.

The XNODE model maps an arbitrary input $\mathbf{x} \in \mathbb{R}^d$ to the output $o_{\mathbf{x}}(t)_{t \in [0, T]} \in \mathcal{C}([0, T])$ by solving the following ODE problem:

$$(5.3) \quad \begin{cases} \frac{d\mathbf{h}(t)}{dt} = \mathcal{N}_{\theta_2}^{\text{vec}}(\mathbf{h}(t), t, \mathbf{x}), & \mathbf{h}(0) = \mathcal{N}_{\theta_1}^{\text{init}}(h(\mathbf{x})) \in \mathbb{R}^h, \\ o_{\mathbf{x}}(t) = \mathcal{L}_{\theta_3}(\mathbf{h}(t)), \end{cases}$$

where $\mathcal{N}_{\theta_2}^{\text{vec}}$ and $\mathcal{N}_{\theta_1}^{\text{vec}}$ are DNN neural networks fully parameterized by $\mathcal{P}_{\theta_2}$ and $\mathcal{P}_{\theta_1}$ for the vector fields and the initial condition $\mathbf{h}(0)$, respectively. $\mathcal{L}_{\theta_3}$ is a single linear layer parameterized by $\mathcal{P}_{\theta_3}$. By $\Theta = (\theta_1, \theta_2, \theta_3)$ we denote the set of all trainable model parameters of the proposed XNODE model. Finally define

$$(5.4) \quad u_{\Theta}(t, \mathbf{x}) := o_{\mathbf{x}}(t) \approx u(t, \mathbf{x}) \quad \forall \mathbf{x} \in \Omega.$$

5.3. Pseudotime XNODE model for static PDEs. In this subsection, we expand the XNODE model to handle stationary PDE problems. To simplify matters, we will focus on the following form of stationary PDE problem:

$$(5.5) \quad \begin{cases} -\nabla \cdot A(\mathbf{x}) \nabla u(\mathbf{x}) + \mathbf{b}(\mathbf{x}) \cdot \nabla u(\mathbf{x}) + c(\mathbf{x})u - f(\mathbf{x}) = 0, & \mathbf{x} \in \Omega = [0, 1]^d, \\ u(\mathbf{x}) = g(\mathbf{x}) & \text{on } \partial\Omega. \end{cases}$$

The idea is to introduce a pseudotime variable, which we choose from one of the spatial variables, x_i , to compensate for the absence of t , i.e., we let $t = x_i$ for some prefixed i . For simplicity, we choose $i = 1$ without loss of generality. The remaining variables $x_i, i = 2, \dots, d$, will form the spatial variables in the XNODE model. More precisely, the spatial input point for the pseudotime XNODE model should be modified as $\tilde{\mathbf{x}} = \{x_2, \dots, x_d\}$. Similar to (5.4), we now define

$$(5.6) \quad u_{\Theta}(\mathbf{x}) = u_{\theta}(x_1, \tilde{\mathbf{x}}) := o_{\tilde{\mathbf{x}}}(x_1) \approx u(x_1, \tilde{\mathbf{x}}),$$

where $o_{\tilde{\mathbf{x}}}(x_1)$ is the numerical solution of (5.3).

5.4. Loss functions. We first recall the classical WAN loss function used in [16]:

$$(5.7) \quad L_{\text{wan}}(\theta, \eta) := \log \left(\frac{|\mathcal{A}(u_{\theta}) - f, \phi v_{\eta}|^2}{\|\phi v_{\eta}\|_{L^2(\mathcal{D})}^2} \right) + \alpha L_{\text{init}}^2(\theta) + \beta L_{\text{bdry}}^2(\theta),$$

where α, γ are hyperparameters as penalty terms and

$$L_{\text{init}}(\theta) = \|u_{\theta}(0, \mathbf{x}) - h(\mathbf{x})\|_{L^2(\Omega)}, \quad L_{\text{bdry}}(\theta) = \|u_{\theta}(t, \mathbf{x}) - g(t, \mathbf{x})\|_{L^2([0, T] \times \partial\Omega)},$$

and $\phi(\mathbf{x})|_{\partial\Omega} = 0$. Here $u_{\theta} \in W_{\theta}$ and $v_{\eta} \in V_{\eta}$, where W_{θ} and V_{η} are neural network function classes parameterized by θ and η , respectively. In this paper, we use the classical DNN function class for V_{η} . For W_{θ} , we will utilize and compare different neural network structures, which will be specified in each experiment. When the PDE problem is static, α is set to 0.

We also define the loss functions for the respective CutWAN and shifted CutWAN methods,

$$(5.8) \quad \begin{aligned} L_{\text{cwan}}(\theta, \eta) &= |\mathcal{A}(u_{\theta}) - f, \phi v_{\eta}| - \gamma_d \|\phi v_{\eta}\|_V^2 + \alpha L_{\text{init}}^2(\theta) + \beta L_{\text{bdry}}^2(\theta), \\ L_{\text{scwan}}(\theta, \eta) &= L_{\text{cwan}}(\theta, \eta) + \|\phi v_{\eta}\|_{H_0^1(\Omega)}. \end{aligned}$$

During computation, the integrals are estimated using the Monte Carlo sampling.

6. Numerical results. The authors carried out the numerical results on a personal CPU device (Apple M1 Max chip with 32 GB memory and 10 total cores). The Adam optimization method is used for all presented numerical experiments.

6.1. Parabolic equations.

Example 1. Following the numerical example in [22, 16], we consider the following nonlinear PDE problem in the form of a d -dimensional nonlinear diffusion-reaction equation (equation (6.1)) defined on a bounded domain $\mathcal{D} \subset [0, 1] \times [-1, 1]^d$:

$$(6.1) \quad \begin{cases} \partial_t u - \Delta u - u^2 - f = 0 & \text{for } (t, \mathbf{x}) \in \mathcal{D}, \\ u - g = 0 & \text{on } \partial\mathcal{D}, \\ u(0, \mathbf{x}) - h(\mathbf{x}) = 0 & \text{on } \Omega(0), \end{cases}$$

where the exact solution is given by

$$(6.2) \quad u(t, \mathbf{x}) = 2 \sin\left(\frac{\pi}{2} x_1\right) \cos\left(\frac{\pi}{2} x_2\right) e^{-t}.$$

The hyperparameters for the XNODE model for u and the DNN model for v used in these experiments are listed in Table 1, and their meanings are explained in Appendix A. The same hyperparameters were maintained across all experiments in Example 1 for the XNODE model. The recursive (nonrecursive) XNODE model u_θ has 1501 (2161) trainable parameters, while the recursive (nonrecursive) model of V_η has 5902 (23351) trainable parameters.

From Table 1, large penalty constants for α and β are utilized. We hypothesize that larger penalty constants can help strongly enforce initial and boundary conditions, which is beneficial in PDE solving using neural networks.

When utilizing the XNODE model to compute u_θ , the training process ceases either when the relative training error drops below 1% or after a maximum of 300 iterations. Conversely, if the DNN model is used to compute u_θ , the maximum number of iterations is set to 3000.

For a comparison, we first train the models using the PINN type loss function defined as follows:

$$(6.3) \quad L_{\text{pinn}}(\theta) = \|\partial_t u_\theta - \Delta u_\theta - u_\theta^2 - f\|_{L^2(\Omega)} + \alpha L_{\text{init}}(\theta) + \beta L_{\text{bdry}}(\theta).$$

The L_{pinn} type loss function was initially introduced in the physics-informed neural network by Raissi, Perdikaris, and Karniadakis [19]. We have conducted experiments on L_{pinn} with the random initialization, and the results are displayed in Figure 2. We utilized the XNODE and DNN models for both the recursive and nonrecursive versions. We note that the PINN loss function requires computing higher-order derivatives, which poses potential challenges. First, the loss function becomes invalid when there is no strong solution, and second, computing these derivatives increases the computational time. In each step, the relative error in Figure 2 and subsequent figures is calculated using a randomly chosen test set, denoted as X_{test} , that is the same size as the training data sets. More precisely, the relative error is computed as

$$\sum_{\mathbf{x}_i \in X_{\text{test}}} \frac{\sum_i (u(\mathbf{x}_i) - u_\theta(\mathbf{x}_i))^2}{\sum_i u(\mathbf{x}_i)^2}.$$

TABLE 1
Hyperparameter setting for Example 1.

d	N_r	N_b	n_T	K_u	K_ϕ	α	β
5	4000	4000	20	2	1	10^7	10^5
ϵ	l_θ	l_η	u_{layers}	$u_{\text{hid-dim1}}$	$u_{\text{hid-dim2}}$	v_{layers}	$v_{\text{hid-dim}}$
10^{-2}	.015	.04	8	20	10	9	50

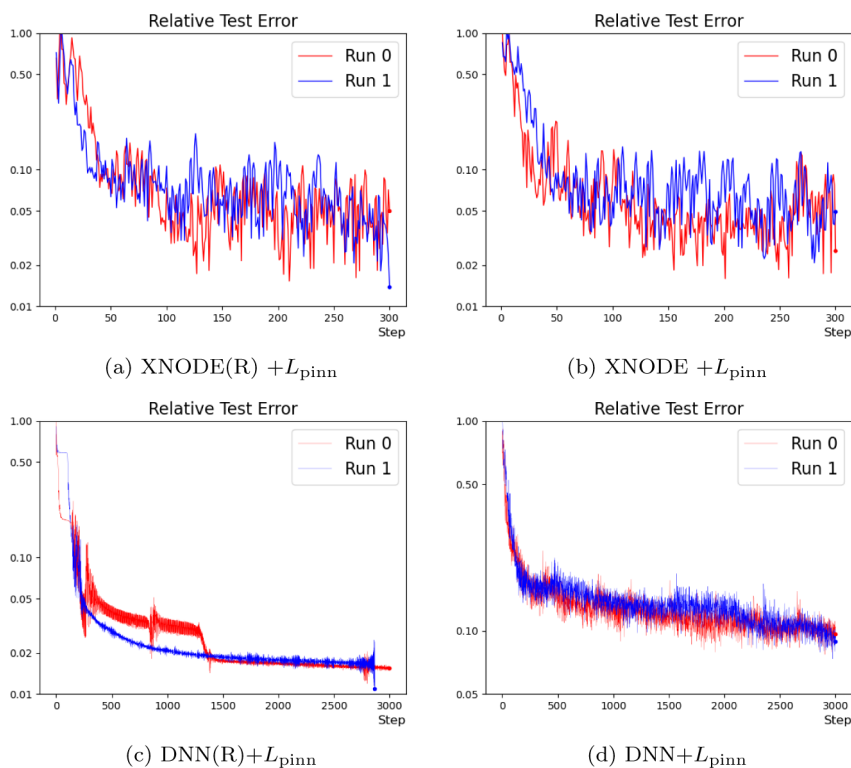


FIG. 2. Example 1: Relative L^2 error versus step for models using L_{pinn} .

It's worth noting that the test set is separate from the training set but has the same size.

In Figure 2, the training time for the DNN and XNODE models is about 2 and 8 seconds per step. The “(R)” after the model denotes the recursive model. In terms of L_{pinn} , it is evident that the DNN model exhibits slower convergence than the XNODE model, whether in recursive or nonrecursive scenarios. When we compare figures in the right column from the left column, it is apparent that the recursive DNN model produces comparable results.

When using the XNODE model, from Figures 2(a) and 2(b), in both cases, relative errors approach the 7% threshold within the first 50 iterations. However, the errors then oscillate with large amplitude, requiring various steps to achieve the next level of accuracy.

In Figure 3, we then train the models using L_{wan} using the same models as in Figure 2. The training time for the DNN and XNODE models is about 2 and 6 seconds per step. After analyzing both Figures 3(c) and 3(d), it is apparent that the utilization of $(L_{\text{wan}} + \text{DNN})$ produces less desirable results compared to $(L_{\text{pinn}} + \text{DNN})$ based on Figure 2. However, the combination of $(L_{\text{wan}} + \text{XNODE})$ produces comparable results with $(L_{\text{pinn}} + \text{XNODE})$. This indicates that XNODE is less sensitive to the chosen objective function.

Based on the observations from Figures 2 and 3, we can deduce that the utilization of the XNODE network for u_θ outperforms the DNN network in both the L_{wan} and L_{pinn} scenarios. However, it is noteworthy that when employing the XNODE network, the loss function during the training exhibits significant oscillation after

reaching a certain level of accuracy for both L_{wan} and L_{pinn} . Additionally, in every scenario presented, the recursive model delivers comparable outcomes to its nonrecursive counterpart.

We now evaluate the XNODE network using the loss functions L_{cwan} and L_{scwan} defined in (5.8) with the same models as shown in Figure 2. In each subfigure in Figure 4, we present the results of five out of six consecutive and randomly initialized experiments under the specific setting to show generality. Each training step takes about 6 seconds.

Overall, after comparing Figures 2 and 3 with Figure 4, we have noticed that L_{cwan} and L_{scwan} show uniformly faster and numerically more stable, i.e., less oscillations, convergence than L_{wan} . Moreover, we observe consistent/robust performance regardless of random initialization. In almost all experiments, the training relative error reaches the 1% relative error all within 200 steps.

When we compare the data in the right column to that of the left column, we notice that the recursive model performs just as well, if not better. Specifically, in the experiments depicted in Figure 4(a), the stopping criteria were met at an average of 144 steps, with individual results of 137, 85, 177, 132, and 190 for experiments 0 to 4, respectively. On the other hand, the nonrecursive counterpart met the stopping criteria at an average of 164 steps, with individual results of 168, 125, 210, 153, and 162 for experiments 0 to 4, respectively, as shown in Figure 4(b). We also note that for the L_{scwan} , the comparison results for $\gamma_d = 0.5$ and $\gamma_d = 0.001$ are similar in this example. However, with $\gamma_d = 0.001$, we notice slightly more oscillations than the case of $\gamma_d = 0.5$.

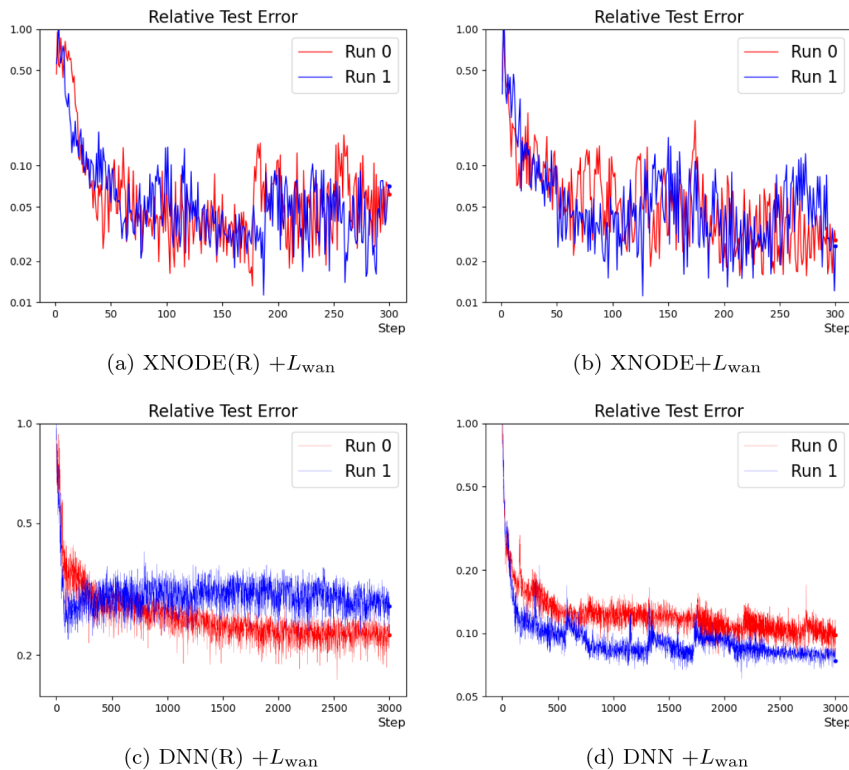


FIG. 3. Example 1: Relative L^2 error versus step for models using L_{wan} .

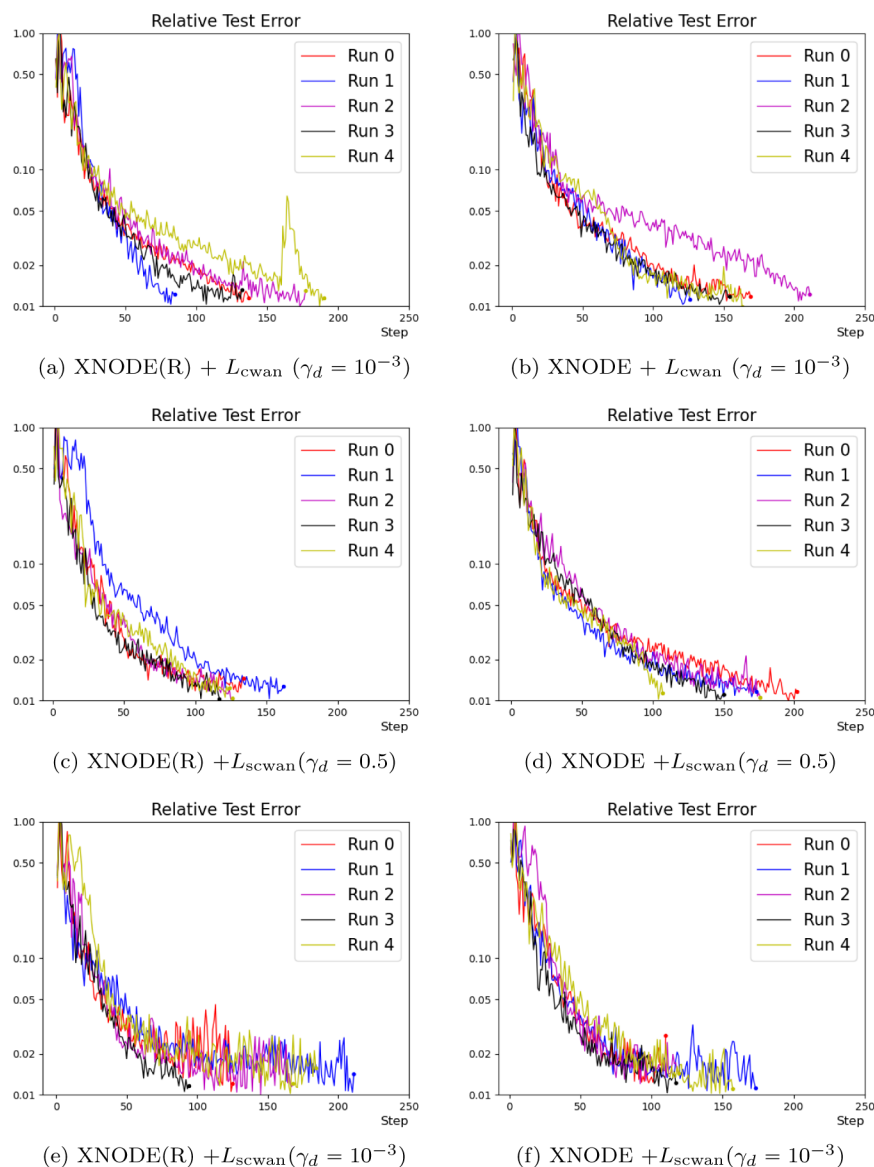
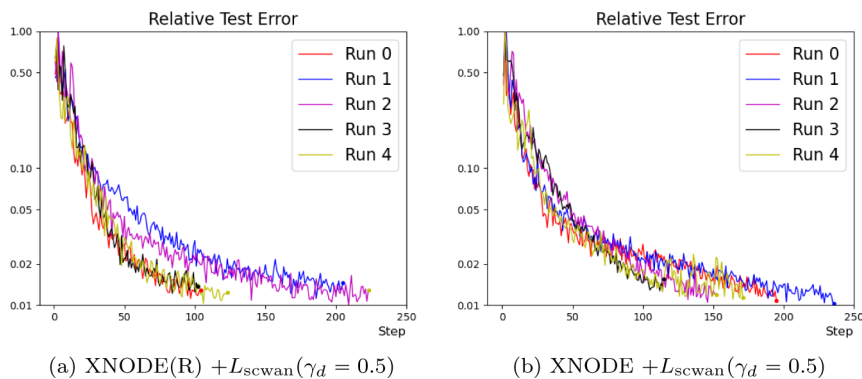


FIG. 4. Example 1: Relative L^2 error for XNODE models on L_{cwan} and L_{scwan} .

In summary, the utilization of the XNODE network and the CutWAN and shifted CutWAN loss functions, i.e., L_{scwan} and L_{cscwan} in (5.8), has demonstrated a highly competitive model for solving high-dimensional parabolic PDE problem. In particular, solving the five-dimensional nonlinear parabolic problem in (6.1) takes only about 15 minutes for the training to reach the 1% relative error on a personal computer. Furthermore, the recursive model necessitates fewer parameters in contrast to nonrecursive models. In comparison to classical numerical techniques such as the finite element method, which grows exponentially in the number of unknowns as the dimension expands, the potential benefit of our approach becomes more prominent as the disk space on a personal computer can rapidly become restricted with the classical approach.

FIG. 5. Example 1: The effect using $\tilde{L}_{\text{bdry}}(\theta)$.

We now consider the effect using the $H^{1/2}$ norm on the boundary based on (4.8). Define

$$\tilde{L}_{\text{bdry}}(\theta) = \|u\|_{L^2(0,T,\partial\Omega)}^{1/2} \|\nabla_{\Gamma} u\|_{L^2(0,T,\partial\Omega)}^{1/2},$$

where $\nabla_{\Gamma} u = \nabla_{\mathbf{x}} u - (\nabla_{\mathbf{x}} u \cdot \mathbf{n})\mathbf{n}$ is the tangential gradient of u and $\mathbf{n}$ is the unit outer normal of Ω . We test the results replacing L_{bdry} in (5.7) by $\tilde{L}_{\text{bdry}}$ using the loss function L_{scwan} with $\gamma_d = 1/2$ (see Figure 5). The results in Figure 5(a) (average iteration number = 151) and Figure 5(b) (average iteration number = 173) are comparable to Figure 4(c) (average iteration number = 132) and Figure 4(d) (average iteration number = 161). However, using $\tilde{L}_{\text{bdry}}$ resulted in an additional duration of approximately 1 s per iteration. For simplicity, we will use L_{bdry} for future experiments. It is worth noting that one can use L_{bdry} for the former iterations and switch to the more accurate $\tilde{L}_{\text{bdry}}$ for better accuracy and time efficiency. This can be necessary when g is of high frequency.

Remark 6.1 (How does V_{η} affect the method's performance?). In the proof, we require V_{η} to be rich enough to satisfy the stability condition. In this example, we tested multiple configurations for the V_{η} network, experimenting with different hidden layers and varying numbers of neurons. The results are all consistent with Figure 4. This indicates that the model is robust with V_{η} for this example.

6.2. Stationary PDE problems.

Example 2. We now test the following high-dimensional problem as in [22]:

$$\begin{cases} -\Delta u(\mathbf{x}) = f, & \mathbf{x} \in \Omega, \\ u(\mathbf{x}) = g(\mathbf{x}), & \mathbf{x} \text{ on } \partial\Omega, \end{cases}$$

where the true solution renders $u(\mathbf{x}) = \sum_{i=1}^d \sin(\frac{\pi}{2}x_i)$.

Observe that the boundary condition on the plane $x_1 = 0$ and $x_1 = 1$ now serves as the initial and terminal conditions in the pseudotime XNODE model. Subsequently, we adjust the initial loss and introduce the terminal loss as

$$(6.4) \quad \begin{aligned} L_{\text{init}}(\theta) &= \|u_{\theta}(x_1 = 0) - g(x_1 = 0)\|_{L^2(\Omega(0))}, \\ L_{\text{last}}(\theta) &:= \|u_{\theta}(x_1 = 1) - g(x_1 = 1)\|_{L^2(\Omega(1))}, \end{aligned}$$

where $\Omega(t) := \{\mathbf{x} \in \Omega, x_1 = t\}$. We shall utilize the following loss functions for the pseudotime XNODE model:

$$(6.5) \quad \tilde{L}_{\text{wan,cwan,scwan}}(\theta, \eta) = L_{\text{wan,cwan,scwan}}(\theta, \eta) + \gamma L_{\text{last}}(\theta).$$

We experimented with testing the pseudotime XNODE model with $d = 5$, using the same parameters as in Table 1 except for the penalty parameters. A grid search was performed to tune the hyperparameters α , β , and γ , which were restricted to the range $[10, 10^9]$. Optimal values found were $\alpha = \gamma = 10^5$ and $\beta = 10^7$.

We established stopping criteria that ensured the relative training error was below 1% or the maximum iteration number less than 300. Each training step takes about 8 seconds.

We have analyzed the loss functions $\tilde{L}_{\text{wan}}$, $\tilde{L}_{\text{scwan}}$ with $\gamma_d = 0.5$, and $\tilde{L}_{\text{cwan}}$ with $\gamma_d = 0.001$, as described in (6.5). We conducted tests on both the recursive and nonrecursive models for each setting, and the outcomes are displayed in Figure 6. Each subfigure in Figure 6 showcases the results of three consecutive experiments that were initialized randomly. The recursive and nonrecursive models are utilized in the left and right columns, respectively.

For $\tilde{L}_{\text{wan}}$, the recursive model in Figure 6(a) reached the stopping criteria at steps 169, 98, and 91 (with an average of 120). Meanwhile, the nonrecursive model in Figure 6(b) reached the stopping criteria at steps 277, 119, and 93 (with an average of 163), based on three experiments for each.

For $\tilde{L}_{\text{cwan}}$, the recursive model in Figure 6(c) reached the stopping criteria at steps 237, 174, and 200 (with an average of 203). Meanwhile, the nonrecursive model in Figure 6(d) reached the stopping criteria at steps 293, 149, and 153 (with an average of 198), based on three experiments for each.

For $\tilde{L}_{\text{scwan}}$, the recursive model in Figure 6(e) reached the stopping criteria at steps 172, 144, and 117 (with an average of 144). Meanwhile, the nonrecursive model in Figure 6(f) reached the stopping criteria at steps 151, 132, and 115 (with an average of 132), based on three experiments for each.

In all XNODE experiments, the relative error quickly reached the 2% threshold within the first 35 iterations. Although the relative error oscillations generated by $\tilde{L}_{\text{wan}}$ are still greater than those of $\tilde{L}_{\text{cwan}}$ and $\tilde{L}_{\text{scwan}}$, it is worth noting that, in this particular case, the stopping criterion was achieved with slightly fewer iterations on average. We believe this faster convergence takes place thanks to the Poisson type PDE used in this example. For the Poisson problem, it is easy to see that $\|\cdot\|_{op}$ is the most natural norm to minimize. It has also been observed that the recursive model's performance is almost comparable to that of the nonrecursive models in this example.

Example 3.

$$(6.6) \quad -\nabla \cdot (a(x)\nabla u) + \frac{1}{2}|\nabla u|^2 = f(x), \text{ in } \Omega = [0, 1]^d, \quad u(x) = g(x) \text{ on } \partial\Omega,$$

where $a(x) = 1 + \|x\|^2$. The true solution $u(x) = \sin(0.5\pi x_1^2 + 0.5x_2^2)$.

In this problem, the nonlinear term $\frac{1}{2}|\nabla u|^2$ presents a significant challenge. We will use the hyperparameter set from Table 2 in all numerical tests with the pseudotime XNODE model. Our objective in this example is to evaluate the performance of the pseudotime XNODE model using various loss functions. We established the stopping criteria for the maximum number of iterations to be less than 600. The duration of each iteration is approximately 8.5 seconds.

We first test the loss function $\tilde{L}_{\text{wan}}$ by conducting three consecutive experiments with random initialization. The results are presented in Figure 7. The left/right figure in Figure 7 shows the relationship between the number of steps and the L^2 relative error/minimal L^2 relative error based on test sets. After the stopping criteria have

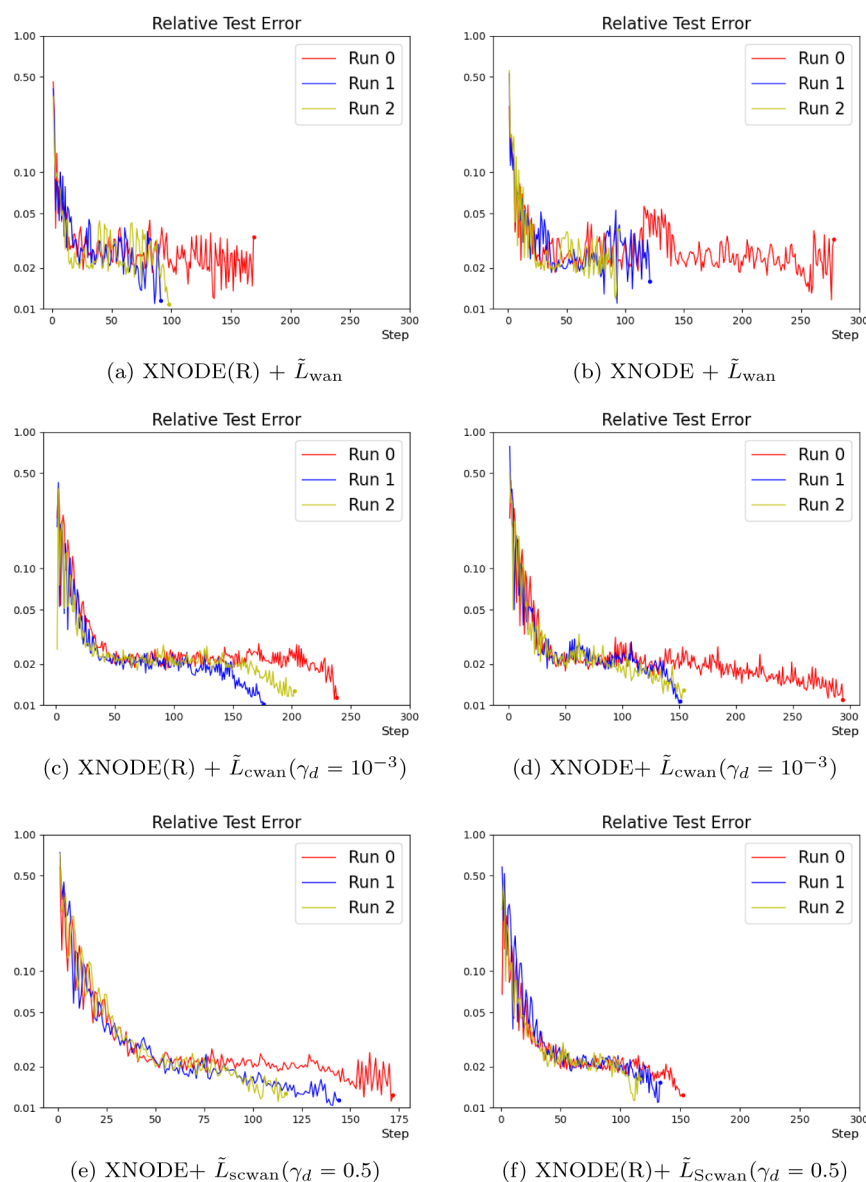


FIG. 6. Example 2. Relative L^2 error versus step using pseudotime XNODE.

TABLE 2
Hyperparameter setting for Example 3.

d	N_r	N_b	n_T	K_u	K_ϕ	α	β	γ
5	4000	4000	20	2	1	6×10^7	12×10^7	12×10^7
ϵ	l_θ	l_η	u_{layers}	$u_{\text{hid-dim1}}$	$u_{\text{hid-dim2}}$	v_{layers}	$v_{\text{hid-dim}}$	
10^{-2}	.015	.03	12	20	10	9	50	

been met, the minimal relative training L^2 error is 0.024, 0.056, and 0.023, respectively. We have observed a slower convergence rate compared to previous examples, which can be attributed to the more challenging nonlinear term.

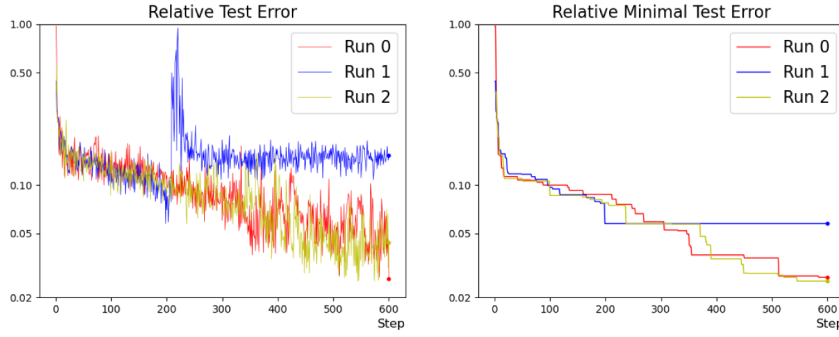


FIG. 7. Example 3. Pseudotime XNODE + $\tilde{L}_{wan}$.

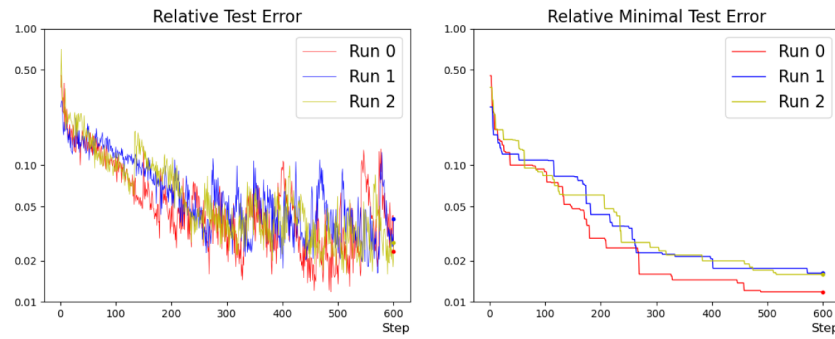


FIG. 8. Example 3. Pseudotime XNODE + $\tilde{L}_{cwan}$ ($\gamma_d = 0.5$).

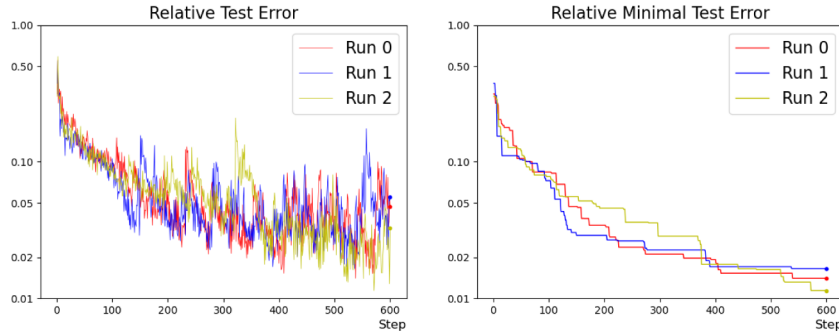


FIG. 9. Example 3. Pseudotime XNODE + $\tilde{L}_{scwan}$ ($\gamma_d = 0.5$).

The results for $\tilde{L}_{cwan}$ and $\tilde{L}_{scwan}$ both with $\gamma_d = 0.5$ are provided in Figures 8 and 9, respectively. For all three experiments, the minimal relative error calculated from $\tilde{L}_{cwan}$ was 0.013, 0.016, and 0.016. Meanwhile, the minimal relative error calculated from $\tilde{L}_{scwan}$ for the same experiments were 0.011, 0.012, and 0.016. Therefore, in comparison to Figure 7, using $\tilde{L}_{cwan}$ and $\tilde{L}_{scwan}$ provides a slight advantage over $\tilde{L}_{wan}$. Due to the highly nonlinear nature of this problem, we believe that a more refined approach to tuning the hyperparameters is necessary to achieve greater accuracy in the results. We will consider this as a future work.

Appendix A. Model setup for XNODE-WAN algorithm. The hyperparameters for the neural networks are explained in Table 3. For V_η , we use a classical DNN network. The activation is set to be Tanh for the last hidden layer and ReLU

TABLE 3
List of hyperparameters.

Notation	Meaning
d	Dimension for the physical domain (not including the time domain)
N_r	Number of sampled collocation points of the spatial domain
N_b	Number of sampled collocation points of the spatial domain boundary
n_T	Number of sampled time partitions
K_u	Inner iteration to update weak solution u_θ or u_θ
K_ϕ	Inner iteration to update test function ϕ_η
α	Weight parameter of boundary loss L_{bdry}
γ	Weight parameter of initial and terminal losses L_{init} and L_{last}
ϵ	Relative error tolerance
l_θ	Learning rate for the primal network
l_η	Learning rate for network parameter η of test function ϕ_η
u_{layers}	Number of hidden layers for $\mathcal{N}_{\theta_2}^{\text{vec}}$
$u_{\text{hid-dim1}}$	Intermediate and output dimension for $\mathcal{N}_{\theta_1}^{\text{init}}$, input dimension for $\mathcal{N}_{\theta_2}^{\text{vec}}$
$u_{\text{hid-dim2}}$	Intermediate dimension for $\mathcal{N}_{\theta_2}^{\text{vec}}$
v_{layers}	Number of hidden layers for V_η
$v_{\text{hid-dim}}$	Intermediate and output dimension for V_η .

for other hidden layers. Note that there is no activation function for the output layer. $\mathcal{N}_{\theta_1}^{\text{init}}$ has one input layer, one hidden layer, and one output layer. The activation function after both the input and the hidden layer is ReLU. $\mathcal{N}_{\theta_2}^{\text{vec}}$ has u_{layers} of hidden dimensions and Tanh as the activation function.

REFERENCES

- [1] G. BAO, X. YE, Y. ZANG, AND H. ZHOU, *Numerical solution of inverse problems by weak adversarial networks*, Inverse Problems, 36 (2020), 115003, <https://doi.org/10.1088/1361-6420/abb447>.
- [2] E. BURMAN, A. FEIZMOHAMMADI, A. MÜNCH, AND L. OKSANEN, *Space time stabilized finite element methods for a unique continuation problem subject to the wave equation*, ESAIM Math. Model. Numer. Anal., 55 (2021), pp. S969–S991, <https://doi.org/10.1051/m2an/2020062>.
- [3] E. BURMAN AND L. OKSANEN, *Data assimilation for the heat equation using stabilized finite element methods*, Numer. Math., 139 (2018), pp. 505–528, <https://doi.org/10.1007/s00211-018-0949-3>.
- [4] A. N. DESHMUKH, M. DEO, P. K. BHASKARAN, T. B. NAIR, AND K. SANDHYA, *Neural-network-based data assimilation to improve numerical ocean wave forecast*, IEEE J. Oceanic Eng., 41 (2016), pp. 944–953.
- [5] M. W. M. G. DISSANAYAKE AND N. PHAN-THIEN, *Neural-network-based approximations for solving partial differential equations*, Commun. Numer. Methods Eng., 10 (1994), pp. 195–201, <https://doi.org/10.1002/cnm.1640100303>.
- [6] C. DUAN, Y. JIAO, Y. LAI, X. LU, Q. QUAN, AND J. Z. YANG, *Deep Ritz methods for Laplace equations with Dirichlet boundary condition*, Commun. Comput. Phys., 31 (2022), pp. 1020–1048.
- [7] M. DUPREZ AND A. LOZINSKI, *ϕ -FEM: A finite element method on domains defined by level-sets*, SIAM J. Numer. Anal., 58 (2020), pp. 1008–1028, <https://doi.org/10.1137/19M1248947>.
- [8] A. ERN AND J.-L. GUERMOND, *Theory and Practice of Finite Elements*, Appl. Math. Sci. 159, Springer, New York, 2004.
- [9] I. GÜHRING, G. KUTYNIOK, AND P. PETERSEN, *Error bounds for approximations with deep ReLU neural networks in $W^{s,p}$ norms*, Anal. Appl., 18 (2020), pp. 803–859.
- [10] I. GÜHRING, M. RASLAN, AND G. KUTYNIOK, *Expressivity of Deep Neural Networks*, Cambridge University Press, Cambridge, UK, 2022.
- [11] K. HE, X. ZHANG, S. REN, AND J. SUN, *Deep residual learning for image recognition*, in Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2016, pp. 770–778.

- [12] Q. HONG, J. W. SIEGEL, AND J. XU, *A Priori Analysis of Stable Neural Network Solutions to Numerical PDEs*, preprint, arXiv:2104.02903, 2021.
- [13] S. KARUMURI, R. TRIPATHY, I. BILIONIS, AND J. PANCHAL, *Simulator-free solution of high-dimensional stochastic elliptic partial differential equations using deep neural networks*, J. Comput. Phys., 404 (2020), 109120.
- [14] S. MAHAN, E. J. KING, AND A. CLONINGER, *Nonclosedness of sets of neural networks in Sobolev spaces*, Neural Networks, 137 (2021), pp. 85–96.
- [15] M. H. MOEINI, A. ETEMAD-SHAHIDI, V. CHEGINI, AND I. RAHMANI, *Wave data assimilation using a hybrid approach in the Persian Gulf*, Ocean Dynam., 62 (2012), pp. 785–797.
- [16] P. V. OLIVA, Y. WU, C. HE, AND H. NI, *Towards fast weak adversarial training to solve high dimensional parabolic partial differential equations using XNODE-WAN*, J. Comput. Phys., 463 (2022), 111233.
- [17] M. OSTROVSKII, *Weak* Sequential Closures in Banach Space Theory and Their Applications*, <https://arxiv.org/abs/arXiv:math/0203139>, 2001.
- [18] P. PETERSEN, M. RASLAN, AND F. VOIGTLAENDER, *Topological properties of the set of functions generated by neural networks of fixed size*, Found. Comput. Math., 21 (2021), pp. 375–444.
- [19] 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*, J. Comput. Phys., 378 (2019), pp. 686–707.
- [20] E. M. STEIN, *Singular Integrals and Differentiability Properties of Functions*, Princeton Math. Ser. 30, Princeton University Press, Princeton, NJ, 1970.
- [21] N. SUKUMAR AND A. SRIVASTAVA, *Exact imposition of boundary conditions with distance functions in physics-informed deep neural networks*, Comput. Methods Appl. Mech. Engrg., 389 (2022), 114333.
- [22] Y. ZANG, G. BAO, X. YE, AND H. ZHOU, *Weak adversarial networks for high-dimensional partial differential equations*, J. Comput. Phys., 411 (2020), 109409, <https://doi.org/10.1016/j.jcp.2020.109409>.
- [23] S. ZENG, Z. ZHANG, AND Q. ZOU, *Adaptive deep neural networks methods for high-dimensional partial differential equations*, J. Comput. Phys., 463 (2022), 111232.