

Tensor Network Computations That Capture Strict Variationality, Volume Law Behavior, and the Efficient Representation of Neural Network States

Wen-Yuan Liu,^{1,*} Si-Jing Du,^{2,*} Ruojing Peng,¹ Johnnie Gray,¹ and Garnet Kin-Lic Chan^{1,†}

¹*Division of Chemistry and Chemical Engineering, California Institute of Technology, Pasadena, California 91125, USA*

²*Division of Engineering and Applied Science, California Institute of Technology, Pasadena, California 91125, USA*

(Dated: May 24, 2024)

We introduce a change of perspective on tensor network states that is defined by the computational graph of the contraction of an amplitude. The resulting class of states, which we refer to as tensor network functions, inherit the conceptual advantages of tensor network states while removing computational restrictions arising from the need to converge approximate contractions. We use tensor network functions to compute strict variational estimates of the energy on loopy graphs, analyze their expressive power for ground-states, show that we can capture aspects of volume law time evolution, and provide a mapping of general feed-forward neural nets onto efficient tensor network functions. Our work expands the realm of computable tensor networks to ones where accurate contraction methods are not available, and opens up new avenues to use tensor networks.

Introduction. Simulating and representing quantum many-body systems is a central task of physics. For this, tensor network states have emerged as a fundamental language and numerical tool [1–6], with applications across diverse areas such as condensed matter physics [7–12], statistical physics [13–16], quantum field theory [17–19], and quantum information theory [20–22], as well as adjacent disciplines such as quantum chemistry and machine learning [23–26].

Tensor networks represent the amplitude of a quantum state on an L site lattice, denoted $\langle \vec{n} | \Psi \rangle \equiv \Psi(\vec{n}) \equiv \Psi(n_1, n_2, \dots, n_L)$, as the contraction of a set of tensors. As a simple example, we consider L tensors connected in a graph (Fig. 1a), each associated with a lattice site Hilbert space; this is known as a projected entangled pair state (PEPS) [27]. Each tensor contains dD^m elements (m is the number of bonds which connect tensors to adjacent tensors): d is the local Hilbert space dimension and D controls the expressivity. For tree graphs, the resulting contraction over the shared tensor bonds can be computed exactly in $O(L)\text{poly}(D)$ time [28]. For general graphs, however, both $\langle \vec{n} | \Psi \rangle$ and $\langle \Psi | \hat{O} | \Psi \rangle$ cannot be exactly contracted efficiently: the cost is exponential in the tree-width of the graph [21]. Approximate contraction algorithms have thus been introduced whereby during the contraction process, the large intermediate tensors that are formed are compressed, via singular value decomposition (SVD), to a controlled size, $O(\text{poly}(\chi))$, where χ is termed the auxiliary bond dimension [2, 29]. Only in the limit of $\chi \rightarrow \infty$ is the contraction exact. The resulting approximate contraction has a cost $O(L)\text{poly}(D)\text{poly}(\chi)$, with an error with respect to the exact contraction controlled by χ .

From the perspective of approximate contraction, tensor network computations should seek to be converged with respect to χ , otherwise the results are not meaningful. For example, for a Hamiltonian $\hat{H}$, if $\langle \Psi | \hat{H} | \Psi \rangle / \langle \Psi | \Psi \rangle$ is computed by approximate contraction, with insufficient χ the energy may violate the variational theorem. The requirement that computations should use a sufficiently large χ (such that the evaluation of $\langle \Psi | \hat{O} | \Psi \rangle$ or $\langle \vec{n} | \Psi \rangle$ is no longer changing

significantly with χ) places restrictions on the types of graphs that can be used to define the tensor network. For example, in many-body physics simulations, it is believed that only highly structured graphs with limited entanglement, such as (near-)area law states, support an efficient contraction scheme. Because of the latter requirement, it has been argued that, in practical use, tensor networks cannot replicate the expressivity of other flexible types of many-body ansatz, such as neural network quantum states [30–36].

In the current work, we observe that the above limitations on tensor network computations arise due to an assumption, namely that the tensor network ansatz is only meaningfully defined in the exact contraction limit. We can, however, adopt an alternative perspective, namely, to consider the approximate contraction scheme for any χ as itself part of the tensor network ansatz for the amplitude $\Psi(\vec{n})$. This view is natural and in some ways obvious when using tensor networks in variational Monte Carlo (VMC) calculations [37–41], but we will argue that this is not a manifestation of tensor networks in a particular numerical algorithm, but rather a change of perspective that alters and broadens the definition and utility of tensor networks. We refer to the larger set of parameterized many-body states that arise from “consistent” contraction schemes (to be defined later) as *tensor network functions*. This wider class of states removes many restrictions that are assumed to apply to tensor networks. We demonstrate that tensor network functions provide strictly variational estimates of the energy, allow for practical tensor network computations on graphs that are highly challenging to contract accurately, including those that describe volume law physics, and expand the set of states that can be efficiently represented by tensor networks to cover other classes of ansatz, such as neural network quantum states.

Tensor network functions. A tensor network function (TNF) is a consistent contraction of the tensor network representation of the amplitude $\Psi(\vec{n})$. To illustrate the basic idea, we first consider a standard PEPS for a two-dimensional lattice, illustrated in Fig. 1b. Because the graph has loops, the PEPS amplitude cannot be computed efficiently according to

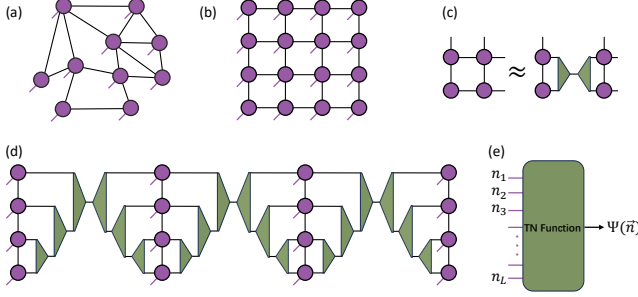


Figure 1. (a) PEPS on a general graph with L site tensors. (b) PEPS on a 4×4 square lattice. (c) Insertion of a pair of isometries for approximate tensor contractions; the isometries can be obtained with standard tensor network techniques by SVD [2, 3]. (d) Approximate PEPS contractions by inserting isometries for a given configuration $|\vec{n}\rangle$. (e) For a given input vector $|\vec{n}\rangle = (n_1, n_2, \dots, n_L)$, a tensor network function outputs a unique value $\Psi(\vec{n})$.

Fig. 1b. Instead, in practice, the amplitude is defined by a different tensor network illustrated in Fig. 1d. In this figure, isometries (triangles) with bond dimension χ have been inserted, where the values of the isometries are determined by an SVD of the contraction of neighbouring tensors (see Fig. 1c) [2, 3]. The final amplitude is obtained by contracting the tensors and isometries together. The pattern of isometries used here corresponds to a specific approximate contraction method (boundary contraction [2, 27]) and other choices can be made [29, 42], but the essential purpose is to make the contraction of the modified graph containing the original tensors and the isometries efficient.

Note, however, that although the structure of the tensor network with isometries in Fig. 1d now appears pseudo-one-dimensional, this does not mean that its entanglement structure (for $\chi > 1$) is reduced to that of a matrix product state (MPS). This is because the isometries are themselves functions of the configuration $|\vec{n}\rangle$, i.e. different values of the isometries are obtained, via SVD, for different configurations. However, we choose to keep the same *positions* of the isometries in the graph for any configuration (even though the entries are different); we refer to this as using *fixed* (position) isometries, which defines a consistent contraction. Because of the configuration dependence of the isometry entries, the graph with isometries no longer represents a multi-linear ansatz like the original tensor network state; it is only a multi-linear ansatz in the limit that $\chi \rightarrow \infty$ (and the isometries become identities). Instead, it defines a *tensor network function*: for any given (D, χ) , any $|\vec{n}\rangle$ is mapped to a single value, the amplitude, through an efficient contraction. The efficient TNF representation of $\Psi(\vec{n})$ does not guarantee deterministic efficiency of computing $\langle \Psi | \hat{O} | \Psi \rangle$; but such expectations may be evaluated stochastically by sampling with efficient sample complexity.

Tensor network functions and approximate contraction. VMC algorithms also compute $\Psi(\vec{n})$ via approximate contraction (for a brief review of tensor network VMC (TN-VMC)

algorithms, see Refs. [37–41]) and it is instructive to consider the differences between the standard TN-VMC approximation of $\Psi(\vec{n})$ and the definition of a TNF above. There are two seemingly small, but significant, differences (i) in TN-VMC, χ is viewed as a parameter (typically $\chi = \text{const} \times D$, with $\text{const} > 1$) where χ is increased to convergence for a fixed D , and (ii) the positions of the isometries are not fixed but dependent on the amplitude that is being computed (we refer to this as *dynamic* isometries), in order to allow for the reuse of computations when computing amplitudes of configurations which differ by only a few substitutions [41].

The most important distinction from TNF is that the use of dynamic isometries in TN-VMC *does not lead to a consistent function*. In other words, given the same configuration $|\vec{n}\rangle$, $\Psi(\vec{n})$ needs not be the same, as it depends on the prior configuration and which components of the tensor network contraction are being reused. Only in the limit of $\chi \rightarrow \infty$ are the amplitudes consistent and the function defined, which underlies the usual need to converge χ in TN-VMC calculations. In contrast, it is clear that with using fixed isometries, the corresponding TNF is consistent for any D and any χ .

Tensor network functions and variational energies. As an immediate consequence of the use of fixed isometries to define a TNF, we can always obtain a variational energy in Monte Carlo energy evaluation (up to statistical error).

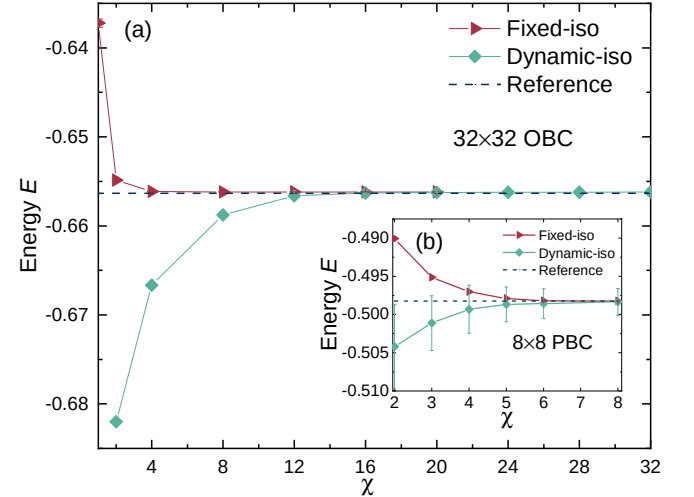


Figure 2. For a given PEPS, energies (per site) obtained from fixed (tensor network function) and dynamic (standard VMC) isometry contractions. (a) OBC square-lattice 32×32 spin-1/2 Heisenberg model with PEPS $D = 8$, and 15000 Monte Carlo sweeps are used for each data point. The dashed line denotes the numerically exact energy [41]. (b) PBC square-lattice 8×8 spin-1/2 frustrated $J_1 - J_2$ model at $J_2/J_1 = 0.5$ with PEPS $D = 5$, and 5000 Monte Carlo sweeps are used for each data point. The dashed line denotes the energy from fixed-isometry contraction with $\chi = 8$ using $D = 5$.

In Fig. 2(a) we show the energy of an optimized PEPS ($D = 8$, various χ) TNF for the ground-state of the open-boundary condition (OBC) spin-1/2 square-lattice 32×32

Heisenberg model, compared to that of a PEPS in standard TN-VMC using dynamic isometries, where χ controls the error of approximate contraction [41]. We see immediately that the PEPS TNF energy is variational for all χ . In contrast, viewing the PEPS amplitude as an (inconsistent) approximate contraction as in standard VMC leads to a non-variational energy, where the non-variationality is typically regarded as the “contraction error”. This behaviour is even more clear for a lattice that is usually considered more difficult to contract. Using the square-lattice $J_1 - J_2$ model with periodic boundary conditions (PBC) as an example, in Fig. 2(b) we show the same comparison between TNF and standard TN-VMC. In addition to the variational versus non-variational nature of the energies, we also see a large variance of the TN-VMC energy arising from the inconsistent definition of amplitudes. This leads to substantial difficulties in the stochastic optimization of the ansatz. Overall, we see that lack of variationality is not intrinsic to approximate contraction: it results from the perspective on the definition of the ansatz.

Expressivity of tensor network functions with small χ . In general, a TNF is parameterized by two bond dimension parameters, D and χ . We already know that as χ becomes large, the TNF reduces to the standard tensor network state from which it is derived, but TNFs are well defined also in the limit of small χ (e.g. $\chi < D$) and it is important to ask whether they are expressive in this limit.

In Figs. 3(a) and (b), we show results from approximating the ground-state of a spin-1/2 square-lattice 16×16 $J_1 - J_2$ Heisenberg model for $J_2/J_1 = 0.5$, and an 8×8 Hubbard model with $U/t = 8$ and $1/8$ hole doping, using bosonic and fermionic PEPS TNFs [27, 43–45], respectively. We minimize the energy for many different (D, χ) combinations using gradient descent, and we plot the computational cost of the amplitude against the obtained accuracy. For a specified relative accuracy, we can then search for the (D, χ) pair with the lowest computational cost. We see that contrary to the standard situation with tensor network states where one uses $\chi > D$, the most powerful TNFs for a given computational cost have $\chi < D$. For example, for ~ 0.01 relative accuracy in the Hubbard model, using $(D, \chi) = (12, 6)$ yields the same expressivity as $(D, \chi) = (10, 10)$, but with lower cost. In no case is it computationally more efficient to use $\chi > D$.

Tensor network functions and volume laws. As an example of a new kind of application that we can consider with TNFs, we consider the case where the tensor network graph corresponds to chaotic time evolution leading to a volume law state. Then the standard view is that approximate contraction should be performed with $\chi \sim \exp(\text{const} \times t)$ (where t is evolution time) to capture the volume law. However, if we think of the TNF formed by a random circuit, we might expect that the set of amplitudes, defined using fixed position but configuration dependent and essentially random isometries, are also scrambled as if by an approximate unitary design, regardless of χ , capturing (some) elements of the volume law behavior.

To illustrate this numerically, we consider the Floquet

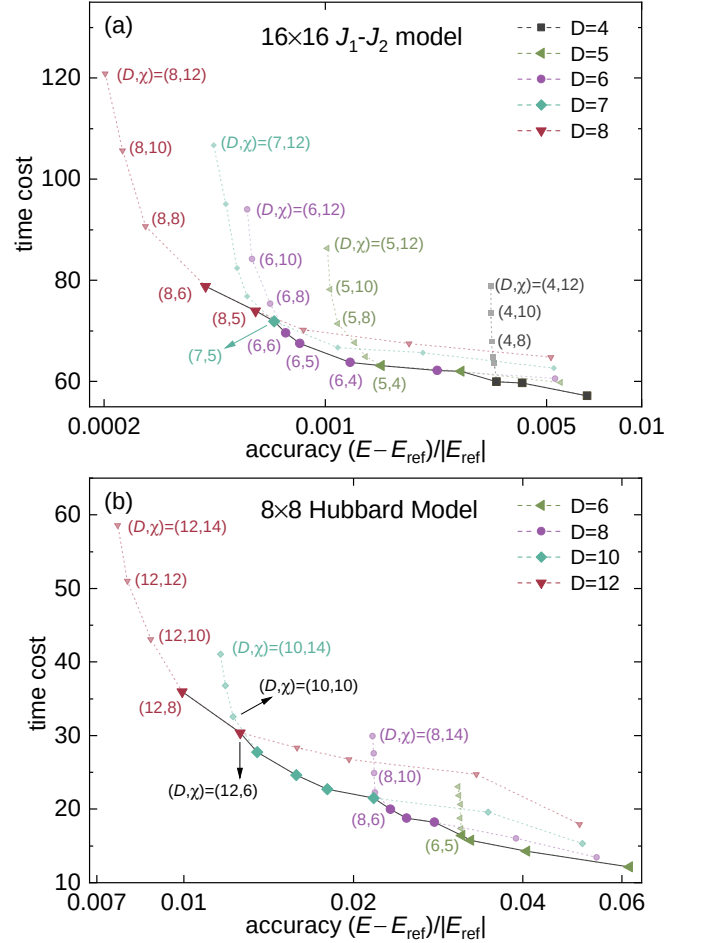


Figure 3. Time cost (in seconds) versus accuracy of the tensor network function $|\Psi(D, \chi)\rangle$, for the $J_1 - J_2$ model on the 16×16 lattice at $J_2/J_1 = 0.5$ (a) and Fermi-Hubbard model on the 8×8 lattice at $U/t = 8$ with hole doping $n_h = 1/8$ (b). Reference energies E_{ref} taken from PEPS $D = 10$ ($J_1 - J_2$, $E = -0.490478(4)$) and DMRG ($D = 10000$ with $SU(2)$ symmetry, Hubbard model, $E = -0.694391$). The time cost is for computing $\langle \vec{n} | \Psi(D, \chi) \rangle$ for a given $|\vec{n}\rangle$ on a single core. Highlighted symbols (connected by a black line) indicate the optimal (D, χ) pairs that have smallest time cost for a given accuracy. Here $D = 4 - 8$ and $\chi = 2, 3, 4, 5, 6, 8, 10, 12$ is used for the $J_1 - J_2$ model and $D = 6 - 12$ and $\chi = 2, 3, 4, 5, 6, 8, 10, 12, 14$ for the Hubbard model.

dynamics of an L -site 1D kicked Ising chain [47–49]. Expressing the Floquet evolution operator F as a matrix product operator (MPO) (see SM for details), the time evolution corresponds to a $(1+1)$ D tensor network, shown in Fig. 4(a). The tunable parameters in the model include the Ising interaction strength J , homogeneous transverse field g and longitudinal field h . We focus on two representative cases for $L = 14$ (i) the maximally chaotic case [50], $(J, g, h) = (\pi/4, \pi/4, 0.5)$, which exhibits volume law physics, and (ii) the less chaotic case, $(J, g, h) = (0.7, 0.5, 0.5)$, which shows sub-volume-law physics at short times. For each time step t , we compute the bipartite von Neumann entropy $S(t) = -\text{Tr}(\rho_A \log \rho_A)$ of the

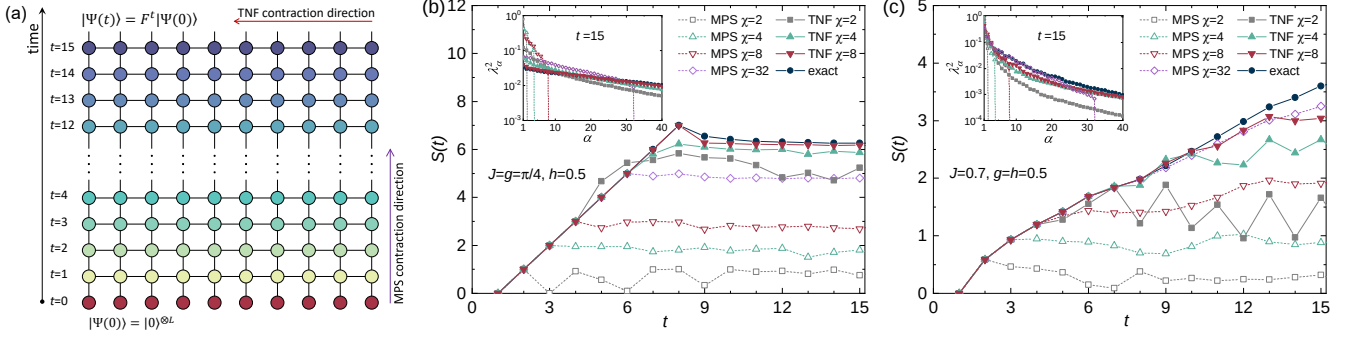


Figure 4. (a) (1+1)D TN representation of the Floquet dynamics of an L -site kicked Ising chain, starting from a product state. TNF contraction direction is either along the spatial direction or inverse time direction (see inverse direction results in SM [46]). Conventional MPS-MPO contraction is along the time direction. (b) and (c) show the dynamics of the half-system bipartite von Neumann entropy $S(t)$ of a 14-site chain at $(J, g, h) = (\pi/4, \pi/4, 0.5)$ and $(0.7, 0.5, 0.5)$, respectively. Insets of (b) and (c) present the entanglement spectrum for the 40 largest eigenvalues of the half-chain reduced density matrix at $t = 15$.

state $|\Psi(t)\rangle = F^t|0\rangle^{\otimes L}$, where ρ_A is the reduced density matrix of the half-chain. (Computing ρ_A is not an efficient operation, but we do so to illustrate the physics being captured). To define a useful TNF, we construct an contraction path that leads to configuration dependent isometries: we consider two natural ones, namely (i) contraction in the inverse time direction, and (ii) (transverse) contraction along the spatial direction. We then compare against the standard MPS-MPO compression to a bond dimension χ . Results from (i), as well as from amplitude-independent MPO-MPO compression in the inverse time direction, are shown in the SM [46].

The evolution of $S(t)$ is depicted in the main panel of Fig. 4(b) and (c), with exact results included. In the maximally chaotic regime (parameter set $(J, g, h) = (\pi/4, \pi/4, 0.5)$), which might be thought of as closest to the random circuit intuition above, the entanglement entropy grows linearly from the beginning until half-width saturation. The required conventional MPS bond dimension to describe the full dynamics grows exponentially with time (here for a finite $L = 14$ chain the maximum χ needed would be $2^{L/2} = 128$). However, volume law physics is already recovered using a TNF with $\chi = 2$. The $\chi = 8$ TNF accurately captures the features of the full evolution, including the maximum point in the entropy at $t = 8$. Similar results are seen for the TNF defined in the inverse time-direction (see SM [46]), although it is quantitatively less accurate.

In the less chaotic regime $(J, g, h) = (0.7, 0.5, 0.5)$, the entanglement entropy initially grows sublinearly up until $t = 7$, as shown in Fig. 4(c). The TNF for any χ captures the initial growth and more entanglement than the conventional MPS of the same bond dimension. However, at small χ the TNF entropy saturates at too small a value relative to the exact result (although the exact result also saturates below the maximum value, see SM). Nonetheless, the entanglement spectrum (insets of Fig. 4) for the TNF, even for small χ , has a long tail and captures the broad features of the spectrum, rather than exhibiting the sharp cutoff of a conventional truncated MPS.

Tensor network functions can efficiently represent neural network quantum states. As discussed TNFs only require a consistent and efficient contraction of $\Psi(\vec{r})$, and there exist TNFs for which this is the case without any (amplitude dependent) isometries, but for which the efficient contraction of $\langle\Psi|\hat{O}|\Psi\rangle$ is still not possible. (We note that tensor networks where computing $\langle\Psi|\hat{O}|\Psi\rangle$ is efficient but $\Psi(\vec{r})$ is not are used in the form of the multi-scale entanglement renormalization ansatz [51]). Here we describe one such application of TNFs without amplitude dependent isometries.

It has previously been observed that common tensor networks that can be efficiently exactly contracted for $\Psi(\vec{r})$ and $\langle\Psi|\hat{O}|\Psi\rangle$, i.e. MPS and tree tensor network states, can be mapped onto efficient neural networks by constructing polynomial size neural nets that emulate tensor network contraction, establishing that exactly contractible tensor networks are a formal subclass of efficient neural quantum states [36]. It has also been suggested that neural networks may not have an efficient tensor network representation, because the well known contractible tensor networks satisfy area laws, while neural networks need not do so [36].

However, if one allows for general tensor network geometries, it is straightforward to construct a TNF (without using isometries) that represents a neural network quantum state without the restriction of an area law. We consider the explicit example of a feed forward neural network (FNN), a popular universal architecture with k -layers of neurons, where the outputs of each layer $\{y_i^{(k-1)}\}$ form the inputs $\{x_i^{(k)}\}$ to the next layer, and the i -th neuron (activation) function in layer k is $f_i^{(k)}(\{x_j^{(k)}\}; \theta_i^{(k)})$ where $\theta_i^{(k)}$ are the neuron parameters. Assuming that all f are implemented without recursion, then the computational graph of a FNN is a directed acyclic graph. Then we observe that (i) if the FNN is efficient, the classical binary circuit implementation of the neural network function is of polynomial size in logic and copy gates, where a binary logic gate is a tensor with two binary input legs and one binary output leg, while copy takes one binary input and returns two binary outputs (nonlinearity of the activation functions enters

through the copy operator) (ii) given the input product state $|\vec{n}\rangle$ (i.e. a bitstring), the application of each gate G produces another product state $|\vec{n}'\rangle$ (i.e. another bitstring). Thus contracting the resulting tensor network in the appropriate order from inputs to the outputs (for example, using sparse tensor contraction, or by introducing SVDs to discover the product state structure of $|\vec{n}'\rangle$, see SM for more details [46]) can always be done efficiently.

The above can be seen as equating a TNF with a type of circuit computational model, in this case without a memory. Although the tensor network contraction is of polynomial cost, it is not very concise with respect to the number of tensors, due to representing floating point numbers in binary. In principle, we can introduce new elements to the computational model that improve the conciseness. In the Supplemental Material, we show how to use the arithmetic circuit tensor network construction of Ref. [52] to compute with tensor networks with floating point entries, and the tensor network contraction is efficient if we further introduce the ability to store results from contractions of tensor network subgraphs.

Conclusions. In summary, we have introduced tensor network functions (TNFs), a new perspective on the tensor network state, as a flexible representation for quantum many-body systems. The broadened perspective removes the conventional restrictions of tensor network computations on the form of the ansatz, and the computational need to converge the auxiliary bond dimension in an approximate contraction. On the other hand, TNFs inherit the conceptual advantages of tensor networks. Unlike generic function approximators, such as neural networks, the function architecture is directly suggested by the physics at hand and one can reuse standard tensor network algorithms. For example, in time evolution, the TNF structure can follow the evolution circuit and choice of contraction path; in variational optimization, existing TN algorithms provide a deterministic guess for the TNF tensors, which may then be further relaxed. TNFs derived from fermionic tensor networks naturally encode fermionic sign structure [44], avoiding the need for more complicated constructions [53–56]. Finally, TNFs can be further generalized into a wider class of states, where non-linearity other than the singular value decomposition is introduced into the tensor network computational graph. These possibilities open up many new avenues to investigate and to explore with tensor networks for the future.

Acknowledgments. This work was supported by the US National Science Foundation under grant no CHE-2102505. GKC acknowledges additional support from the Simons Investigator program and the Dreyfus Foundation under the program Machine Learning in the Chemical Sciences and Engineering.

[†] gkc1000@gmail.com

- [1] S. R. White, Density matrix formulation for quantum renormalization groups, *Phys. Rev. Lett.* **69**, 2863 (1992).
- [2] F. Verstraete, V. Murg, and J. I. Cirac, Matrix product states, projected entangled pair states, and variational renormalization group methods for quantum spin systems, *Advances in Physics* **57**, 143 (2008).
- [3] U. Schollwöck, The density-matrix renormalization group in the age of matrix product states, *Annals of Physics* **326**, 96 (2011).
- [4] G. K.-L. Chan and S. Sharma, The density matrix renormalization group in quantum chemistry, *Annual Review of Physical Chemistry* **62**, 465 (2011).
- [5] J. I. Cirac, D. Pérez-García, N. Schuch, and F. Verstraete, Matrix product states and projected entangled pair states: Concepts, symmetries, theorems, *Rev. Mod. Phys.* **93**, 045003 (2021).
- [6] T. Xiang, *Density Matrix and Tensor Network Renormalization* (Cambridge University Press, 2023).
- [7] Z.-C. Gu, M. Levin, and X.-G. Wen, Tensor-entanglement renormalization group approach as a unified method for symmetry breaking and topological phase transitions, *Phys. Rev. B* **78**, 205116 (2008).
- [8] N. Schuch, D. Pérez-García, and J. I. Cirac, Classifying quantum phases using matrix product states and projected entangled pair states, *Phys. Rev. B* **84**, 165139 (2011).
- [9] P. Corboz, T. M. Rice, and M. Troyer, Competing states in the t - J model: Uniform d -wave state versus stripe state, *Phys. Rev. Lett.* **113**, 046402 (2014).
- [10] H. J. Liao, Z. Y. Xie, J. Chen, Z. Y. Liu, H. D. Xie, R. Z. Huang, B. Normand, and T. Xiang, Gapless spin-liquid ground state in the $S = 1/2$ kagome antiferromagnet, *Phys. Rev. Lett.* **118**, 137202 (2017).
- [11] B.-X. Zheng, C.-M. Chung, P. Corboz, G. Ehlers, M.-P. Qin, R. M. Noack, H. Shi, S. R. White, S. Zhang, and G. K.-L. Chan, Stripe order in the underdoped region of the two-dimensional hubbard model, *Science* **358**, 1155 (2017).
- [12] W.-Y. Liu, J. Hasik, S.-S. Gong, D. Poilblanc, W.-Q. Chen, and Z.-C. Gu, Emergence of gapless quantum spin liquid from deconfined quantum critical point, *Phys. Rev. X* **12**, 031039 (2022).
- [13] M. Levin and C. P. Nave, Tensor renormalization group approach to two-dimensional classical lattice models, *Phys. Rev. Lett.* **99**, 120601 (2007).
- [14] Z. Y. Xie, H. C. Jiang, Q. N. Chen, Z. Y. Weng, and T. Xiang, Second renormalization of tensor-network states, *Phys. Rev. Lett.* **103**, 160601 (2009).
- [15] G. Evenbly and G. Vidal, Tensor network renormalization, *Phys. Rev. Lett.* **115**, 180405 (2015).
- [16] S. Yang, Z.-C. Gu, and X.-G. Wen, Loop optimization for tensor network renormalization, *Phys. Rev. Lett.* **118**, 110504 (2017).
- [17] F. Verstraete and J. I. Cirac, Continuous matrix product states for quantum fields, *Phys. Rev. Lett.* **104**, 190405 (2010).
- [18] J. Haegeman, T. J. Osborne, H. Verschelde, and F. Verstraete, Entanglement renormalization for quantum fields in real space, *Phys. Rev. Lett.* **110**, 100402 (2013).
- [19] L. Tagliacozzo, A. Celi, and M. Lewenstein, Tensor networks for lattice gauge theories with continuous groups, *Phys. Rev. X* **4**, 041024 (2014).
- [20] G. Vidal, Efficient classical simulation of slightly entangled quantum computations, *Phys. Rev. Lett.* **91**, 147902 (2003).
- [21] I. L. Markov and Y. Shi, Simulating quantum computation by contracting tensor networks, *SIAM Journal on Computing* **38**, 963 (2008).

* These authors contributed equally to this work.

- [22] I. Arad and Z. Landau, Quantum computation and the evaluation of tensor networks, *SIAM Journal on Computing* **39**, 3089 (2010).
- [23] K. Boguslawski, K. H. Marti, O. Legeza, and M. Reiher, Accurate ab initio spin densities, *Journal of Chemical Theory and Computation* **8**, 1970 (2012).
- [24] N. Nakatani and G. K.-L. Chan, Efficient tree tensor network states for quantum chemistry: Generalizations of the density matrix renormalization group algorithm, *The Journal of Chemical Physics* **138**, 134113 (2013).
- [25] E. M. Stoudenmire and D. J. Schwab, Supervised learning with tensor networks, *Advances in Neural Information Processing Systems* **29**, 4799 (2016).
- [26] Z.-Y. Han, J. Wang, H. Fan, L. Wang, and P. Zhang, Unsupervised generative modeling using matrix product states, *Phys. Rev. X* **8**, 031012 (2018).
- [27] F. Verstraete and J. I. Cirac, Renormalization algorithms for quantum-many body systems in two and higher dimensions, *arXiv:0407066* (2004).
- [28] Y.-Y. Shi, L.-M. Duan, and G. Vidal, Classical simulation of quantum many-body systems with a tree tensor network, *Phys. Rev. A* **74**, 022320 (2006).
- [29] S.-J. Ran, E. Tiritto, C. Peng, X. Chen, L. Tagliacozzo, G. Su, and M. Lewenstein, *Tensor network contractions: methods and applications to quantum many-body systems* (Springer Nature, 2020).
- [30] H. J. Changlani, J. M. Kinder, C. J. Umrigar, and G. K.-L. Chan, Approximating strongly correlated wave functions with correlator product states, *Phys. Rev. B* **80**, 245116 (2009).
- [31] D.-L. Deng, X. Li, and S. Das Sarma, Quantum entanglement in neural network states, *Phys. Rev. X* **7**, 021021 (2017).
- [32] J. Chen, S. Cheng, H. Xie, L. Wang, and T. Xiang, Equivalence of restricted boltzmann machines and tensor network states, *Phys. Rev. B* **97**, 085104 (2018).
- [33] I. Glasser, N. Pancotti, M. August, I. D. Rodriguez, and J. I. Cirac, Neural-network quantum states, string-bond states, and chiral topological states, *Phys. Rev. X* **8**, 011006 (2018).
- [34] Y. Levine, O. Sharir, N. Cohen, and A. Shashua, Quantum entanglement in deep learning architectures, *Phys. Rev. Lett.* **122**, 065301 (2019).
- [35] L. Pastori, R. Kaubruegger, and J. C. Budich, Generalized transfer matrix states from artificial neural networks, *Phys. Rev. B* **99**, 165123 (2019).
- [36] O. Sharir, A. Shashua, and G. Carleo, Neural tensor contractions and the expressive power of deep neural quantum states, *Phys. Rev. B* **106**, 205136 (2022).
- [37] A. W. Sandvik and G. Vidal, Variational quantum Monte Carlo simulations with tensor-network states, *Phys. Rev. Lett.* **99**, 220602 (2007).
- [38] N. Schuch, M. M. Wolf, F. Verstraete, and J. I. Cirac, Simulation of quantum many-body systems with strings of operators and Monte Carlo tensor contractions, *Phys. Rev. Lett.* **100**, 040501 (2008).
- [39] L. Wang, I. Pižorn, and F. Verstraete, Monte Carlo simulation with tensor network states, *Phys. Rev. B* **83**, 134421 (2011).
- [40] W.-Y. Liu, S.-J. Dong, Y.-J. Han, G.-C. Guo, and L. He, Gradient optimization of finite projected entangled pair states, *Phys. Rev. B* **95**, 195154 (2017).
- [41] W.-Y. Liu, Y.-Z. Huang, S.-S. Gong, and Z.-C. Gu, Accurate simulation for finite projected entangled pair states in two dimensions, *Phys. Rev. B* **103**, 235155 (2021).
- [42] J. Gray and G. K.-L. Chan, Hyperoptimized approximate contraction of tensor networks with arbitrary geometry, *Phys. Rev. X* **14**, 011009 (2024).
- [43] W.-Y. Liu, S.-S. Gong, Y.-B. Li, D. Poilblanc, W.-Q. Chen, and Z.-C. Gu, Gapless quantum spin liquid and global phase diagram of the spin-1/2 J_1 - J_2 square antiferromagnetic Heisenberg model, *Science Bulletin* **67**, 1034 (2022).
- [44] C. V. Kraus, N. Schuch, F. Verstraete, and J. I. Cirac, Fermionic projected entangled pair states, *Phys. Rev. A* **81**, 052338 (2010).
- [45] S. Lee, J. Lee, H. Zhai, Y. Tong, A. M. Dalzell, A. Kumar, P. Helms, J. Gray, Z.-H. Cui, W. Liu, *et al.*, Evaluating the evidence for exponential quantum advantage in ground-state quantum chemistry, *Nature Communications* **14**, 1952 (2023).
- [46] See Supplemental Material including Refs. [57–63], for additional information about the numerical simulations and tensor network function representation of neural network computational graphs.
- [47] T. Prosen, General relation between quantum ergodicity and fidelity of quantum dynamics, *Phys. Rev. E* **65**, 036208 (2002).
- [48] T. Prosen, Ruelle resonances in quantum many-body dynamics, *Journal of Physics A: Mathematical and General* **35**, L737 (2002).
- [49] T. Prosen, Chaos and complexity of quantum motion, *Journal of Physics A: Mathematical and Theoretical* **40**, 7881 (2007).
- [50] B. Bertini, P. Kos, and T. Prosen, Entanglement spreading in a minimal model of maximal many-body quantum chaos, *Phys. Rev. X* **9**, 021033 (2019).
- [51] G. Vidal, Entanglement renormalization, *Phys. Rev. Lett.* **99**, 220405 (2007).
- [52] R. Peng, J. Gray, and G. K.-L. Chan, Arithmetic circuit tensor networks, multivariable function representation, and high-dimensional integration, *Phys. Rev. Res.* **5**, 013156 (2023).
- [53] Y. Nomura, A. S. Darmawan, Y. Yamaji, and M. Imada, Restricted boltzmann machine learning for solving strongly correlated quantum systems, *Phys. Rev. B* **96**, 205152 (2017).
- [54] J. R. Moreno, G. Carleo, A. Georges, and J. Stokes, Fermionic wave functions from neural-network constrained hidden states, *Proceedings of the National Academy of Sciences* **119**, e2122059119 (2022).
- [55] J. Nys and G. Carleo, Variational solutions to fermion-to-qubit mappings in two spatial dimensions, *Quantum* **6**, 833 (2022).
- [56] Z. Liu and B. K. Clark, A unifying view of fermionic neural network quantum states: From neural network backflow to hidden fermion determinant states, *arXiv:2311.09450* (2023).
- [57] S. Kumar, *Digital Vlsi Design* (Prentice Hall India Pvt., Limited, 2010).
- [58] G. Lancaster, *Excel HSC Softw Design&Devel + Cards SG*, Excel HSC (Pascal Press, 2001).
- [59] M. Mano, *Digital Logic and Computer Design* (Pearson India, 2017).
- [60] T.-K. Liu, K. Hohulin, L.-E. Shiao, and S. Muroga, Optimal one-bit full adders with different types of gates, *IEEE Transactions on Computers* **C-23**, 63 (1974).
- [61] C. Baugh and B. Wooley, A two's complement parallel array multiplication algorithm, *IEEE Transactions on Computers* **C-22**, 1045 (1973).
- [62] M. Hatamian and G. Cash, A 70-mhz 8-bit/spl times/8-bit parallel pipelined multiplier in 2.5-/spl mu/m cmos, *IEEE Journal of Solid-State Circuits* **21**, 505 (1986).
- [63] Y. Leong, H. Lo, M. Drieberg, A. B. Sayuti, and P. Sebastian, Performance comparison review of 8–3 compressor on fpga, in *TENCON 2017 - 2017 IEEE Region 10 Conference* (2017) pp. 2462–2465.

Supplemental Material

S-1. 1D kicked Ising model dynamics

A. Definition of the model

Here we give the detailed definition of the 1D kicked Ising model and its corresponding unitary Floquet dynamics [47–49] discussed in the main article. Starting from a product state $|0\rangle^{\otimes L}$ in an L -site spin chain, we consider the time evolution of the kicked Ising Hamiltonian $H(t) = H_I + H_K \sum_{m=-\infty}^{\infty} \delta(t - m)$, where $\delta(t)$ is the Dirac δ function, with

$$H_I = -J \sum_{j=1}^{L-1} \sigma_j^z \sigma_{j+1}^z - h \sum_{j=1}^L \sigma_j^z, \quad (S1)$$

$$H_K = -g \sum_{j=1}^L \sigma_j^x. \quad (S2)$$

Setting the time interval between kicks to 1, the unitary Floquet operator generated by the above Hamiltonian reads as (set $\hbar = 1$):

$$F_{KI} = e^{-iH_K} e^{-iH_I} \quad (S3)$$

$$= \prod_j e^{ig\sigma_j^x} \prod_j e^{ih\sigma_j^z} \prod_{j \text{ even}} e^{iJ\sigma_j^z \sigma_{j+1}^z} \prod_{j \text{ odd}} e^{iJ\sigma_j^z \sigma_{j+1}^z}, \quad (S4)$$

which admits a straightforward quantum circuit representation with only single-qubit and two-qubit gates, as shown in Fig. S1(a). By combining the unitary gates within each Floquet time step into an MPO (where we use singular value decomposition to split the two-qubit gates into local MPOs and contract tensors on each site), we can represent the Floquet evolution as a $(1 + 1)$ D tensor network as illustrated in Fig. S1(b).

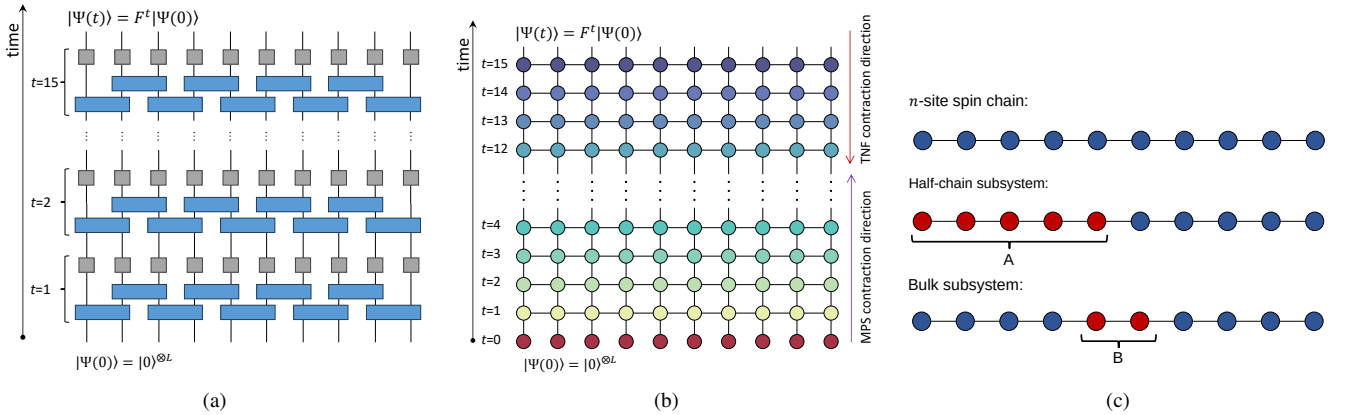


Figure S1. (a) The quantum circuit representation of the kicked Ising Floquet dynamics. (b) The $(1 + 1)$ D tensor network for Floquet dynamics, with contraction order either along the inverse-time direction for tensor network function or time-evolution direction for standard MPS-MPO compression. (c) The 1D spin chain and different subsystems to compute the entanglement entropy.

B. Numerical results on the 1D kicked Ising model dynamics

1. Choice of model parameters

To study the quantum (pure) state entanglement in the kicked Ising model dynamics, for each Floquet time step t , we compute the half-chain bipartite von Neumann entropy $S(t) = -\text{Tr}(\rho_A \log \rho_A)$ of the state $|\Psi(t)\rangle = F^t |0\rangle^{\otimes L}$ (where A indicates the subsystem of the half-chain (the second row in Fig. S1(c))) and record its value along the Floquet evolution. The reduced

density matrix is obtained as the partial trace of the whole density matrix $\rho = \sum_{ij} \langle \vec{n}_i | \Psi \rangle \langle \Psi | \vec{n}_j \rangle |\vec{n}_i\rangle \langle \vec{n}_j|$, which is constructed by enumerating all amplitudes $\langle \vec{n}_i | \Psi \rangle$ (assuming normalization) in the σ^z basis for the system sizes considered in this work. We choose two representative sets of parameters $(J, g, h) = (\pi/4, \pi/4, 0.5)$ and $(J, g, h) = (0.7, 0.5, 0.5)$. Both parameter sets correspond to non-integrable dynamics but have different entanglement growth speeds. For $(J, g, h) = (\pi/4, \pi/4, 0.5)$, the kicked Ising Floquet dynamics is maximally chaotic where the half-chain entanglement entropy grows at a maximal speed and linearly in time [50]. For $(J, g, h) = (0.7, 0.5, 0.5)$, the dynamics is less chaotic and has a slower half-chain entanglement growth speed. Results of the exact entanglement dynamics for both parameter sets are shown in Fig. 4 in the main article as well as in Fig. S5. Here we point out that the choice of the parameter values for the less chaotic case is not fine-tuned, but rather an arbitrary choice that is away from the maximally chaotic regime. We expect similar conclusions (about tensor network function) to hold for other general choices of parameters in the 1D kicked Ising dynamics away from the maximally chaotic regime.

2. Spatial entanglement structure of the evolved states

To demonstrate the spatial entanglement structure of the evolved states along the kicked Ising dynamics, we consider a bulk subsystem B in the middle of a spin chain (the third row in Fig.S1(c)) of size $L = 14$, and compute the entanglement entropy $S_B(L_B)$ of its reduced density matrix ρ_B as a function of the subsystem size L_B for Floquet time up to $t = 20$. For the maximally chaotic case $(J, g, h) = (\pi/4, \pi/4, 0.5)$, after a very short evolution time, the state starts to display volume-law physics where the bulk subsystem entanglement entropy grows linearly with the subsystem size L_B , as shown in Fig. S2(a). For the less chaotic case $(J, g, h) = (0.7, 0.5, 0.5)$, many early-evolved states obey an entanglement entropy's area law where the bulk subsystem entropy saturates as the subsystem size increases, as shown in Fig. S2(b).

In the maximally chaotic regime, an accurate conventional MPS simulation requires a bond dimension χ that grows exponentially with the system size, even for a short-time simulation. In the less chaotic case, however, a conventional MPS with a finite constant bond dimension χ can still describe the evolved state in the early-time evolution regime.

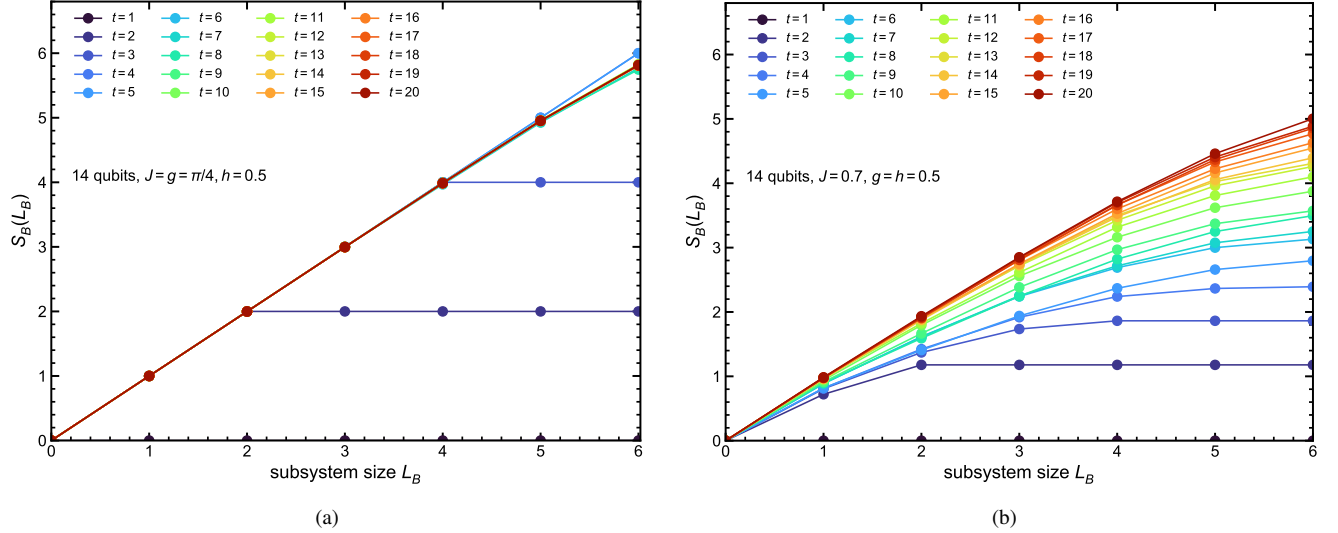
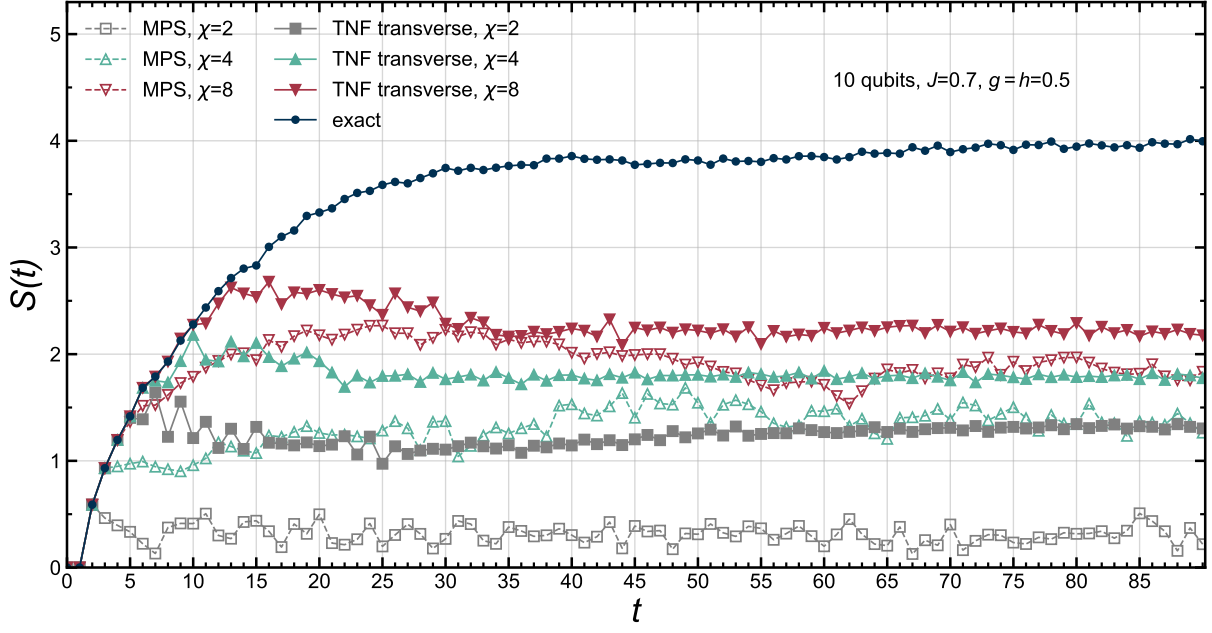


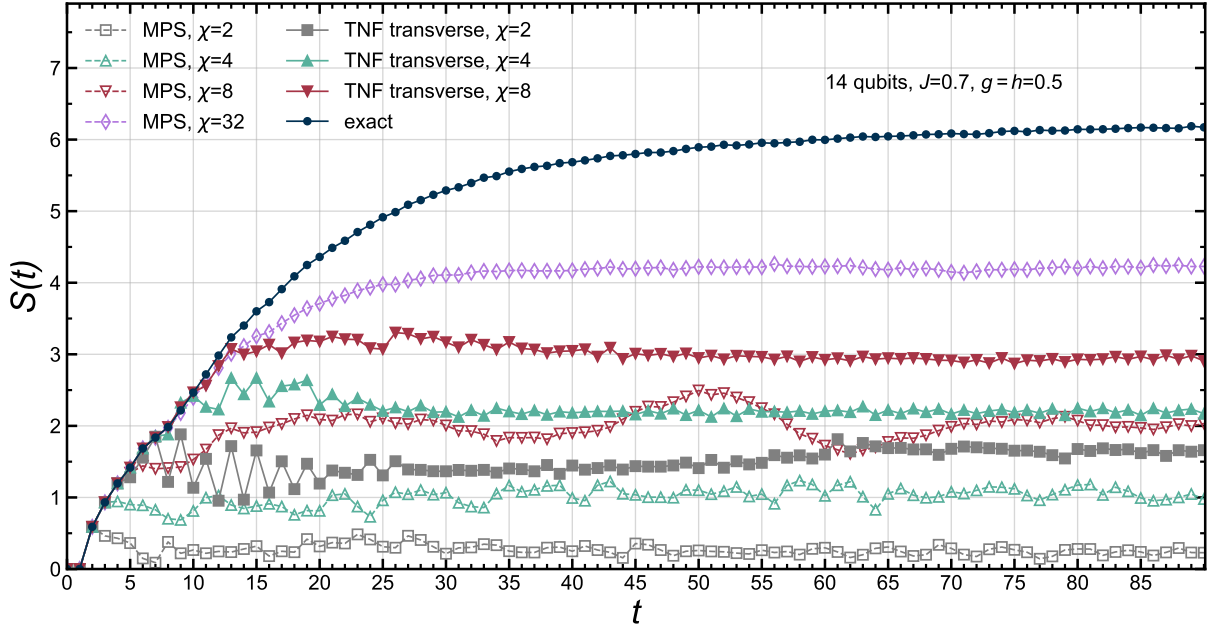
Figure S2. Bulk subsystem entanglement entropy scaling with respect to the subsystem size along the kicked Ising model dynamics, obtained from exact state vector simulation for 20 Floquet steps in a $L = 14$ spin chain. (a) Maximally chaotic regime ($J = g = \pi/4, h = 0.5$), where evolved states after $t = 3$ already display volume-law behavior. (b) Less chaotic regime ($J = 0.7, g = h = 0.5$), where many of the evolved states in early steps display area-law entanglement entropy behavior.

3. Longer time entanglement dynamics for the less chaotic regime

In the Fig. 4(c) in the main article where the dynamics is less chaotic at $(J, g, h) = (0.7, 0.5, 0.5)$, the half-chain entanglement entropy of the tensor network function appears to saturate at Floquet time $t = 15$. To show that the entanglement entropy grows very slowly at long times, we present more calculation results up to Floquet time $t = 90$ for $L = 10$ and $L = 14$ chains, as shown in Fig. S3.



(a)



(b)

Figure S3. Entanglement dynamics of the kicked Ising chain in the less chaotic regime for Floquet time up to $t = 90$ for (a) $L = 10$ chain and (b) $L = 14$ chain. Transverse contraction is used for the TN function, and conventional MPS-MPO contraction along the time direction is used for the MPS results.

4. Half-chain entanglement spectrum

To further analyze the entanglement structure captured by tensor network function, we compute the entanglement spectrum of the half-chain reduced density matrix at Floquet step $t = 15$ for the $L = 10$ and $L = 14$ chains, comparing the results of the conventional MPS and tensor network function. More specifically, we evolve a product state $|0\rangle^{\otimes L}$ under the kicked Ising dynamics (Eq. S4) for $t = 15$ steps, then compute the eigenvalues λ_α^2 of the half-chain reduced density matrix $\rho_A = \sum_\alpha \lambda_\alpha^2 |\alpha\rangle\langle\alpha|$

of the final state using the conventional MPS and tensor network function. Exact state vector simulation is also performed to provide benchmark. Results are shown in Fig. S4. The tensor network function is contracted transversely along the spatial direction. Figs. S4(c) and (d) are clearer versions of the insets of Figs. 4(b) and (c) in the main article. As shown, the tensor network function with small isometry bond dimension χ can already capture the broad features of the spectrum, while the conventional MPS with bond dimension χ exhibits a sharp cutoff at the χ -th largest eigenvalue.

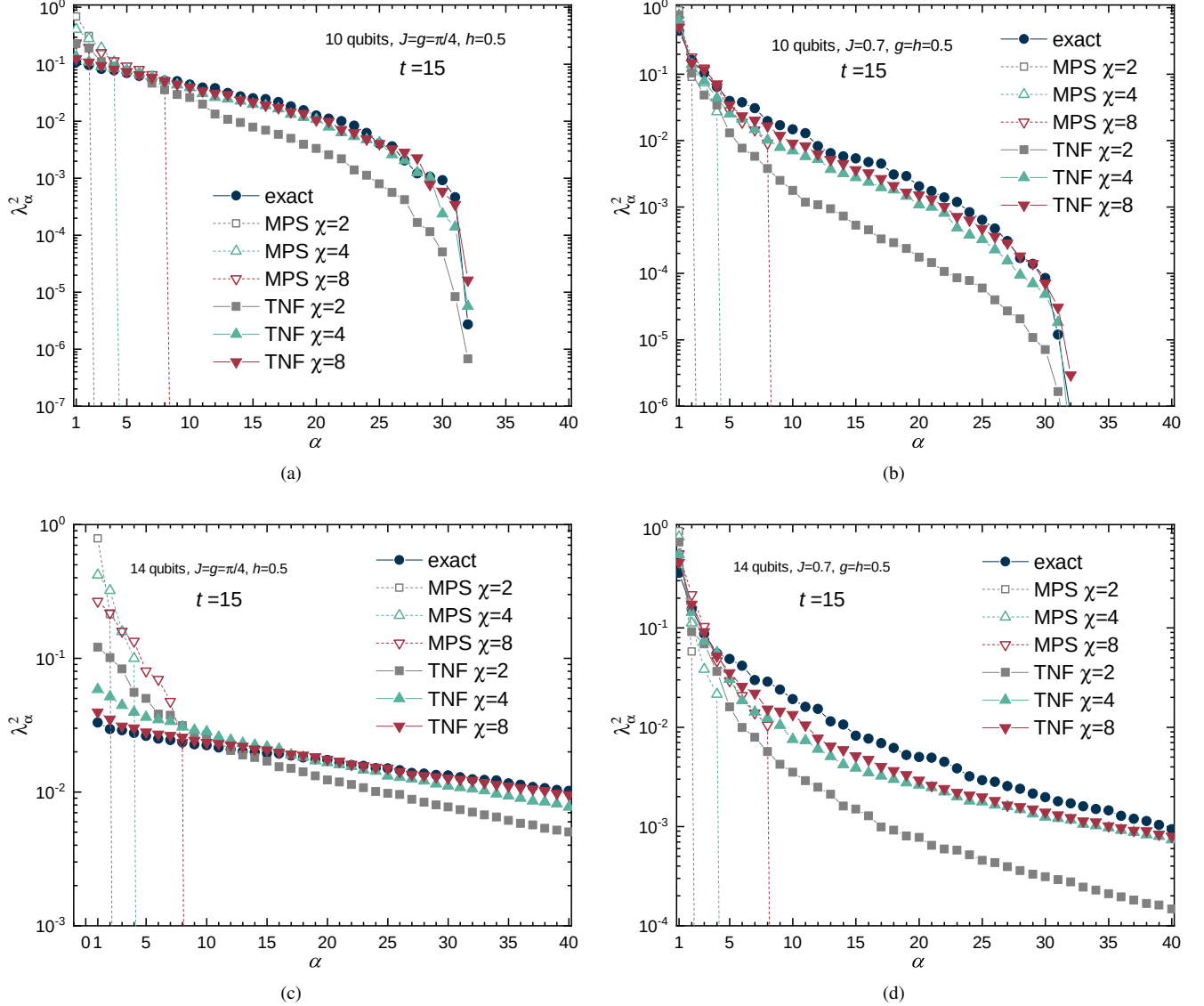


Figure S4. Half-chain entanglement spectrum for $L = 10$ (a)-(b) and $L = 14$ (c)-(d) kicked Ising chains at Floquet step $t = 15$. λ_α^2 is the α -th largest eigenvalue of the half-chain reduced density matrix, and 40 largest eigenvalues are shown here. (a) and (c): Maximally chaotic regime ($J = g = \pi/4$, $h = 0.5$); (b) and (d): Less chaotic regime ($J = 0.7$, $g = 0.5$, $h = 0.5$). Transverse contraction is used for the TN function, and conventional MPS-MPO contraction along the time direction is used for the MPS results.

5. Entanglement dynamics from inverse-time contraction

The tensor network function formalism allows for very flexible ways to define the amplitude from the tensor network graph, depending on how one inserts the isometries. Each choice of isometry insertion yields a different tensor network function. In the numerical data shown in the main article, we employed a transverse contraction along the spatial direction for the calculation of amplitudes. Here we perform the same half-chain entanglement dynamics calculation using a tensor network

function corresponding to contraction along the inverse-time direction, as illustrated in Fig. S1(b).

In Fig. S5, we show the numerical results for the entanglement entropy dynamics of the 1D kicked Ising model for two different parameter sets for $L = 10$ and $L = 14$ spin chains, where the tensor network functions are contracted along the inverse-time direction and the conventional MPS results are also shown for comparison. In addition, we show results from an inverse-time MPO-MPO contraction (i.e Heisenberg-style evolution) for comparison. In the latter case, the isometries are not amplitude dependent, as the amplitude is evaluated only after the entire tensor network is contracted (in the inverse time direction) into an MPS.

One can see that with the same χ , the tensor network function contracted in the inverse direction captures more entanglement than the conventional MPS, as well as the Heisenberg-style MPS from MPO-MPO contraction. However, compared to the transverse contraction, inverse-time contraction has a lower accuracy at the early steps of the Floquet dynamics, showing a convergence of the entanglement entropy ahead of the true saturation time for the maximally chaotic kicked Ising dynamics.

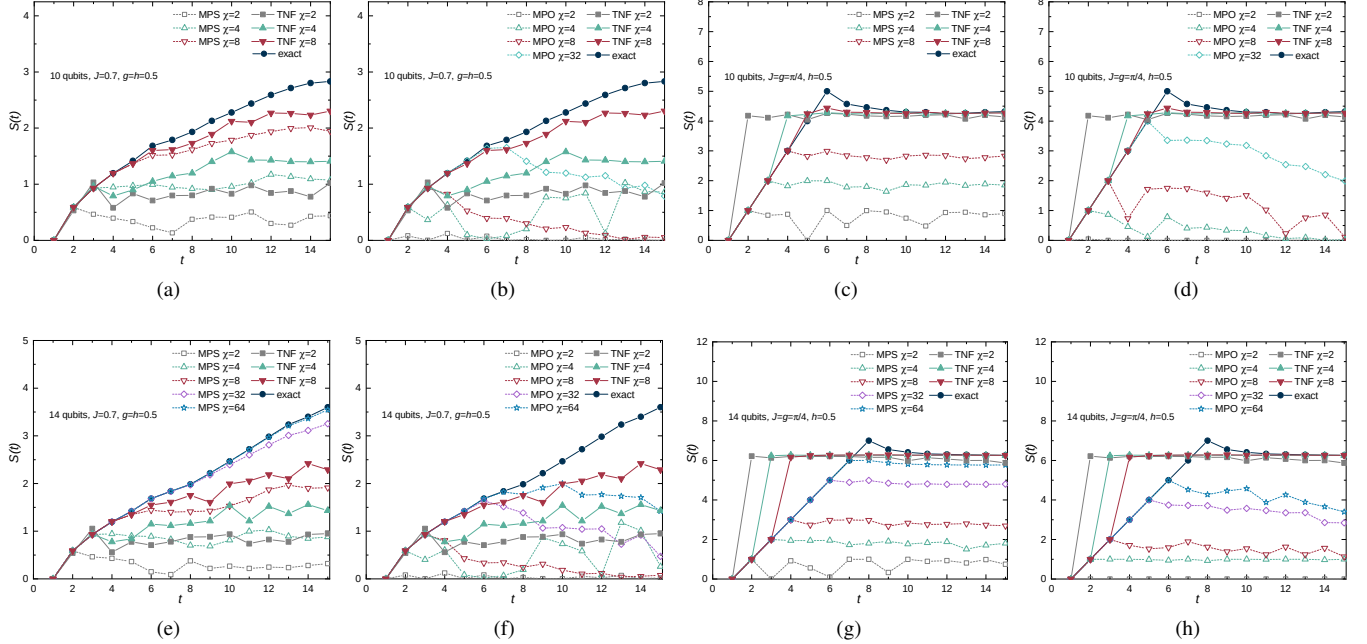


Figure S5. Half-chain entanglement entropy dynamics for $L = 10$ (a)-(d) and $L = 14$ (e)-(h) kicked Ising chain, obtained from different approaches. Inverse-time contraction is used for the TN function. For (a),(c),(e) and (g), conventional MPS-MPO contraction along the time direction is used (labeled as MPS in legends). For (b),(d),(f) and (h), MPO-MPO contraction along the inverse time direction is used (labeled as MPO in legends).

S-2. Additional discussion of the tensor network function representation of neural network computational graphs

A. Classical binary arithmetic circuits

In our first construction of the tensor network function representation of a feed forward neural network, we use classical binary arithmetic using logic and copy gates. Here we describe in more detail this binary arithmetic and the building blocks (in terms of the tensors) of polynomial functions. (We can use polynomial functions to efficiently represent activation functions in a neural network under mild conditions). Note that we have chosen the construction below for simplicity and this is not the only binary logic construction of arithmetic: for more examples, we refer to [57–60].

The n -bit binary representation of an integer

$$x = \sum_{i=0}^n x_i 2^i, \quad x_i \in \{0, 1\} \quad (\text{S5})$$

is a product of length-2 vectors as shown in Fig. S6a, where each vector

$$(x_i)_j = \delta_{x_i, j}; \quad (\text{S6})$$

floating point numbers can be represented similarly. Consider addition

$$z = x + y = \sum_{i=0}^n (x_i + y_i) 2^i. \quad (\text{S7})$$

The tensor network function to compute the 0-th digit

$$z_0 = (x_0 + y_0) \bmod 2 \quad (\text{S8})$$

is shown in the upper dotted box of Fig. S6b, where the XOR tensor of shape (2,2,2) has zero entries except at

$$(\text{XOR})_{0,0,0} = (\text{XOR})_{1,0,1} = (\text{XOR})_{0,1,1} = (\text{XOR})_{1,1,0} = 1. \quad (\text{S9})$$

The contribution carried to the 1st digit

$$c_1 = \text{floor}((x_0 + y_0)/2) \quad (\text{S10})$$

is shown in the lower dotted box of Fig. S6b where the AND tensor of shape (2,2,2) has zero entries except at

$$(\text{AND})_{0,0,0} = (\text{AND})_{1,0,0} = (\text{AND})_{0,1,0} = (\text{AND})_{1,1,1} = 1. \quad (\text{S11})$$

Fig. S6b is also known as a "half adder" in an electric circuit [57–59]. For any subsequent digit i , the diagram to compute

$$z_i = (c_i + x_i + y_i) \bmod 2 \quad (\text{S12})$$

is shown in the upper dotted box of Fig. S6c. The contribution carried over to the next digit

$$c_{i+1} = \text{floor}((c_i + x_i + y_i)/2) \quad (\text{S13})$$

is shown in the lower dotted box of Fig. S6c, where the OR tensor of shape (2,2,2) has zero entries except at

$$(\text{OR})_{0,0,0} = (\text{OR})_{1,0,1} = (\text{OR})_{0,1,1} = (\text{OR})_{1,1,1} = 1. \quad (\text{S14})$$

Fig. S6c is also known as a "full adder" in an electric circuit[57–60]. In Fig. S6d we show the full circuit representation for $x + y = z$ with each of x and y containing 3 digits, and where the boxes with 2 and 3 input legs represent the tensor network subgraph in Fig. S6b and Fig. S6c respectively.

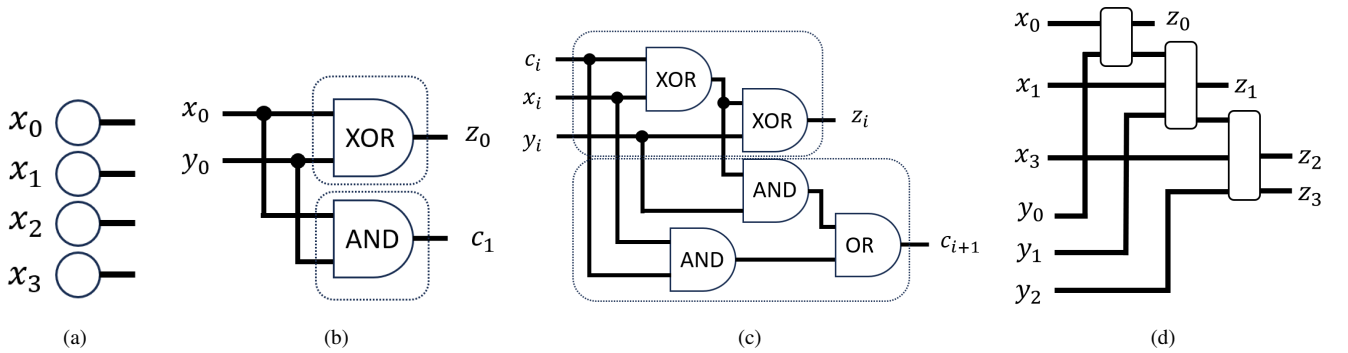


Figure S6. (a) Representing a bit string as a product of length-2 vectors. (b) Diagram for adding 2 bits. (c) Diagram for adding 3 bits. (d) Diagram for $x + y = z$ where x and y each has 3 digits.

Multiplication

$$z = x \times y = \left(\sum_{i=0}^m x_i 2^i \right) \left(\sum_{j=0}^n y_j 2^j \right) \quad (\text{S15})$$

in the binary arithmetic circuit representation can be built from binary long multiplication. Here we illustrate this with an example for $m = 3, n = 2$, where long multiplication

$$\begin{array}{r}
 \begin{array}{cccc}
 & x_3 & x_2 & x_1 & x_0 \\
 \times & & y_2 & y_1 & y_0 \\
 \hline
 & y_0 x_3 & y_0 x_2 & y_0 x_1 & y_0 x_0 \\
 + & y_1 x_3 & y_1 x_2 & y_1 x_1 & y_1 x_0 \\
 + & y_2 x_3 & y_2 x_2 & y_2 x_1 & y_2 x_0
 \end{array}
 \end{array} \quad (S16)$$

first forms partial products shown in each row under the line and then sums the shifted partial products. The corresponding tensor network is shown in Fig. S7, where all 3-legged tensors are the AND tensors given in Eq. S11, and partial products of y_0x , y_1x and y_2x are represented by blue, orange and green respectively. The summation of partial products is done sequentially using the tensor network subgraph shown in Fig. S6d represented as the dashed boxes in Fig. S7. In digital circuits, various more efficient alternatives exist [61–63].

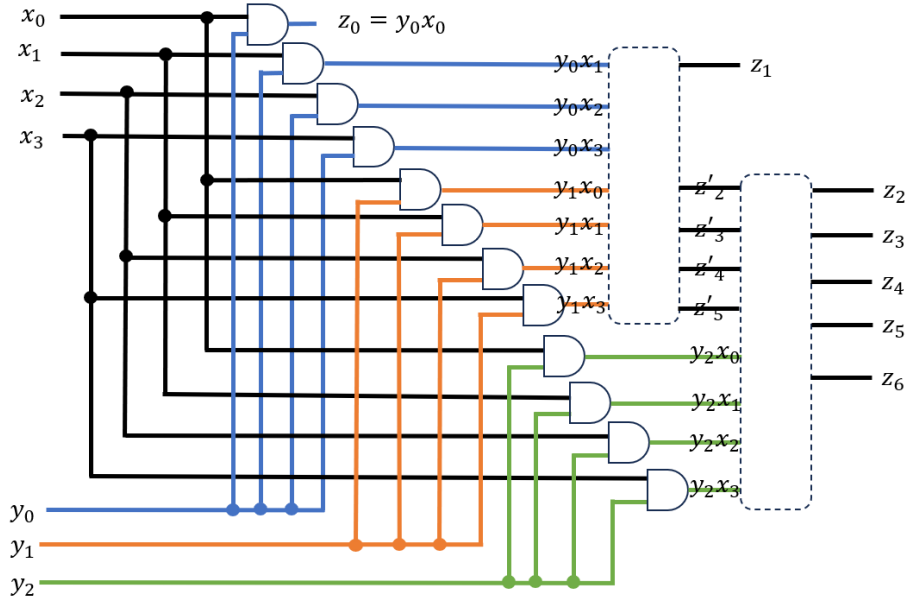


Figure S7. Tensor network for computing $z = xy$ with $x = 2^3x_3 + 2^2x_2 + 2x_1 + x_0$ and $y = 2^2y_2 + 2y_1 + y_0$. Partial products y_0x , y_1x and y_2x are represented in blue, orange and green. The dashed boxes represent the subgraph for adding 2 bit strings as shown in Fig. S6d.

Finally, we comment on the copying of data. In the binary arithmetic circuit, the information contained in a bit can be fully copied by the δ tensor which has zero entries except at

$$\delta_{0,0,0} = \delta_{1,1,1} = 1 \quad (S17)$$

and is represented as the black dot in Fig. S6b, Fig. S6c and Fig. S7. This is because a bit represents its value by the index associated with the nonzero amplitude, which is 1, but the nonzero amplitude itself contains no additional information. Therefore, let x represent a bit, then

$$\sum_i x_i \delta_{ijk} = \sum_i \delta_{x,i} \delta_{ijk} = \delta_{xjk} = \delta_{x,j} \delta_{x,k} = x_j x_k \quad (S18)$$

where we use the representation given in Eq. S6. To illustrate this point, we show the tensor network subgraph for computing $z = x^2$ in Fig. S8 upto the formation of all partial products, where only one set of x -bits is needed as input.

Furthermore, we note that contracting a bit with the δ tensor results in a product of 2 vectors, as in Eq. S18. Therefore, for a given pair of bit inputs, the output of Fig. S6b is a product; and similarly for Fig. S6c. Hence for a given bit string input, Fig. S6d and thus Fig. S7 can be easily contracted.

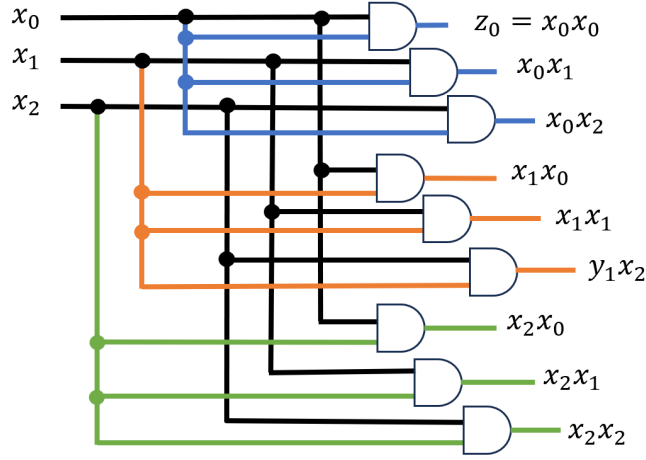


Figure S8. Tensor network for computing $z = x^2$ with $x = 2^2 x_2 + 2x_1 + x_0$. Partial products $x_0 x$, $x_1 x$ and $x_2 x$ are represented in blue, orange and green.

B. Amplitude arithmetic circuit construction of tensor network functions

The arithmetic circuit construction in Ref. [52] represents a floating point number x as a length-2 vector

$$X_i = x^i, \quad i = 0, 1 \quad (\text{S19})$$

where superscript i denotes raising to i the power. The diagrams for addition $z = x + y$ and for multiplication $w = xy$ are shown in Fig. S9a, where the tensors (+) and ($\times$) have shape $(2, 2, 2)$ and entries

$$(+)_{ijk} = \delta_{i+j,k}, \quad (\times)_{ijk} = \delta_{ijk} \quad (\text{S20})$$

for $i, j, k = 0, 1$ such that

$$Z_k = \sum_{ij} X_i Y_j (+)_{ijk} = \sum_{ij} x^i y^j \delta_{i+j,k} = (x + y)^k \quad (\text{S21})$$

$$W_k = \sum_{ij} X_i Y_j (\times)_{ijk} = \sum_{ij} x^i y^j \delta_{ijk} = (xy)^k. \quad (\text{S22})$$

The above representation allows us to efficiently compose more complicated functions using the addition and multiplication tensors and floating point numbers, but special consideration must be given to copying data (i.e. to construct a non-linear dependence). To illustrate the difference in copying data from the case of classical binary circuits, we consider the addition of single variable functions $z(x, y) = f(x) + g(y)$ shown in the upper diagram in Fig. S9b; multiplication is similar by replacing the (+) tensor with ($\times$). In the following we will use $p, q, r, \dots$ to denote tensor legs associated with discretized variables, and $i, j, k, \dots$ to denote the size-2 leg contracted with the (+) and ($\times$) tensors for performing arithmetic. The addition (multiplication) of single variable functions of the same variable $z(x) = f(x) + g(x)$ ($w(x) = f(x)g(x)$) is shown in Fig. S9c, where the blue dot represents the copy tensor on the variable leg δ_{pqr} . In particular, given tensor $F_{p,i}$ representing the discretized function $f(x)$, the tensor network representation of $w(x) = f(x)^2$ requires 2 copies of $F_{p,i}$ with the variable legs constrained to be the same by the δ_{pqr} tensor. This is to be contrasted with the binary representation, where a single copy of a bit string x is needed to construct the tensor network representation of x^2 . In general, as demonstrated in the previous section, copying of data in the binary representation is done using the shape $(2, 2, 2)$ δ_{ijk} tensor in Eq. S17 from a *single* copy of the bit string. In contrast, copying of data in the amplitude representation requires additional actual copies of the same tensor.

In principle, the requirement of actual copies of tensors in the amplitude representation can lead to exponential growth in the number of tensors in the full tensor network function representation of a computational graph, for instance, for computing an amplitude for a given input in a deep neural network. However, this formal exponential full tensor network representation is never necessary in actual computation. As in the case of the binary circuit representation, contraction of the amplitude circuit representation can proceed by trivially following the computational graph and storing intermediates as needed. So long as we augment the tensor network computational model with the ability to store intermediates in a “memory”, we can replace copied

subgraphs of the tensor network function by their evaluated value when it is available. Therefore, contraction of the full circuit representation can be done without forming the full tensor network, and the contraction has the same cost as the computational graph that it trivially follows.

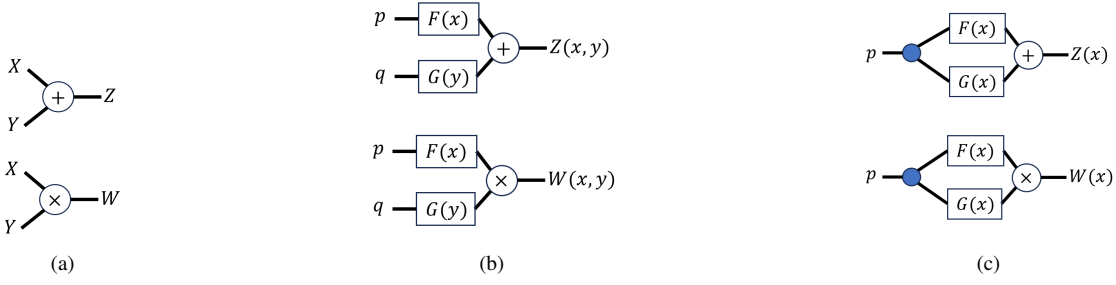


Figure S9. (a) Tensor diagrams for addition $z = x + y$ and for multiplication $w = xy$. (b) Tensor network diagrams for addition and multiplication of functions $z(x[p], y[q]) = f(x[p]) + g(y[q])$ and $z(x[p], y[q]) = f(x[p])g(y[q])$ where p, q index the discretized variables x, y respectively. (c) Tensor network diagrams for addition and multiplication of functions $z(x[p]) = f(x[p]) + g(x[p])$ and $z(x[p]) = f(x[p])g(x[p])$.