

# State Vector to Block Encoding: An Efficient Quantum Circuit Transformation

Giacomo Antonioli<sup>1,\*</sup>, Alessandro Berti<sup>1</sup>, Gianna M. Del Corso<sup>1</sup> and Alessandro Poggiali<sup>1</sup>

<sup>1</sup>Department of Computer Science, University of Pisa, Italy

## Abstract

Quantum state preparation and block encoding constitute fundamental building blocks in quantum computation. The state vector representation encodes data within the amplitudes of a quantum state. Conversely, block encoding is tailored for quantum linear algebra applications, enabling the manipulation of encoded operators via matrix functions. Because these frameworks are complementary, converting between them is essential for quantum linear algebra algorithms. In this paper, we provide a quantum subroutine for converting the state vector representation of  $A \in \mathbb{C}^{N \times N}$ , where  $\|A\|_F = 1$ , directly into a block encoding of  $A$ , using a constant-depth circuit of three layers of one- and two-qubit gates and only  $n = \log_2 N$  ancillary qubits. We show that the resulting normalization factor  $\sqrt{N}$  is theoretically optimal.

## Keywords

Quantum algorithms, block encoding, State preparation

## 1. Introduction

Quantum state preparation is a fundamental building block of quantum computation [1]. It consists in evolving a quantum system into a desired target state, so that subsequent unitary operations can process the data encoded in the amplitudes of the state vector. In many quantum algorithms [2, 3, 4, 5], the asymptotic advantage crucially depends on the availability of efficient state preparation routines [6, 7, 8, 9]. Block encoding (BE) [10, 11], on the other hand, is a fundamental tool for implementing quantum linear algebra algorithms. It allows embedding a general non-unitary matrix into a larger unitary operator that can be realized by a quantum circuit. Efficient block encoding algorithms exist for sparse [12, 13] or structured matrix classes [14, 15, 16] — such as matrices with displacement structure, semiseparable, and hierarchical matrices — that explicitly exploit the underlying structure; block encoding of general dense matrices is also possible [17, 11].

These two representations, state vectors and block-encoded operators, are complementary and suited for different algorithmic settings. State vector encodings embed information into the amplitudes of a quantum state, enabling its preparation and subsequent manipulation via unitary evolution, making it suitable for vector-based linear algebra operations [18]; operator encodings are preferable when the goal is to compute matrix functions with QSVT [10], such as inversion of matrices or computation of the exponential matrix function necessary in Hamiltonian simulation. Given this complementarity, it is often necessary to convert between the two representations. In this direction, authors in [18] propose a conversion from exact block encoding to matrix state preparation, while in [19] a subroutine for the reverse transformation from matrix state to operator is proposed using  $3n$  qubits, where the matrix  $A \in \mathbb{C}^{N \times N}$  and  $n = \log_2 N$ . Although the cost in depth is not explicitly analysed, the construction involves one doubly-controlled  $|j\rangle\langle g|$  operation for each basis pair  $(j, g) \in \{0, \dots, N-1\}^2$ , suggesting gate count that scales as  $O(N^2)$  in the number of controlled operations, followed by a diffusion layer and  $2n$  Hadamard gates. For a more general conversion, the authors in [17] introduce a bidirectional

ICTCS 2026: Italian Conference on Theoretical Computer Science, September 7–9, 2026, Udine, Italy

\*Corresponding author.

✉ giacomo.antonioli@phd.unipi.it (G. Antonioli); alessandro.berti@df.unipi.it (A. Berti); gianna.delcorso@unipi.it (G. M. Del Corso); alessandro.poggiali@phd.unipi.it (A. Poggiali)

🆔 0009-0000-6687-0357 (G. Antonioli); 0000-0001-9144-9572 (A. Berti); 0000-0002-5651-9368 (G. M. Del Corso); 0000-0002-1591-7925 (A. Poggiali)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

transition between arbitrary block encodings and state preparation. Their algorithm for converting a matrix to a block encoding operator requires an intermediate step where first the matrix is decomposed as a linear combination of higher-order Pauli matrices, using a multiplexer which they assume to implement with constant depth. The resulting circuit has constant depth, uses additional  $2n$  qubits and the normalization factor is  $\alpha N$  if the initial state contains also some garbage, and  $\alpha$  is the normalization factor for the entries of  $A$ .

In this paper, we propose an approach for constructing a block encoding of a matrix  $A$  directly from its matrix state preparation operator  $M_A$ , using a circuit of constant depth consisting of 3 layers of single- and two-qubit gates, with normalization factor  $\sqrt{N}$ . We prove that the algorithm is optimal within the class of algorithms that query  $M_A$  exactly once. The code implementation is available at <https://github.com/Giacomo-Antonioli/State-Vector-to-Block-Encoding--An-Efficient-Quantum-Circuit-Transformation>.

## 2. Preliminary notation

For simplicity, all the results are presented for square matrices, but they can be easily extended to the rectangular case. Given a matrix  $A = (a_{ij}) \in \mathbb{C}^{N \times N}$ , in the following we denote by  $A_j = [a_{0j}, a_{1j}, \dots, a_{N-1,j}]^T \in \mathbb{C}^N$ , the  $j$ -th column of  $A$ , i.e.  $A = [A_0 | A_1 | \dots | A_{N-1}]$ , and by  $\text{vec}(A) = [a_{00}, a_{10}, \dots, a_{N-1,N-1}]$  the vector obtained concatenating the columns of  $A$  one after the other. The encoding of the matrix  $A$  into the amplitudes of a quantum state is achieved by the matrix state preparation circuit. We assume that  $\|A\|_F = 1$  since this is a parameter that should be precomputed classically and the entries of  $A$  always need to be scaled by this factor.

**Definition 1.** (*Matrix state preparation operator*) [18] A  $2n$ -qubit unitary operator  $M_A$  is a matrix state preparation operator for  $A \in \mathbb{C}^{N \times N}$ , with  $N = 2^n$ , and  $\|A\|_F = 1$ , if, applying  $M_A$  to  $|0\rangle_{2n}$ , results in

$$|A\rangle = \sum_{i=0}^{N-1} \sum_{j=0}^{N-1} a_{ij} |j\rangle_n |i\rangle_n.$$

Equivalently  $M_A |0\rangle_{2n} = |A\rangle = \text{vec}(A)$ .

Alternatively, the matrix  $A$  can be encoded, using block encoding, into a unitary operator  $U_A$  that encodes  $A/\alpha$  as its top left block, where  $\alpha \geq \|A\|_2$  is a scaling factor.

**Definition 2.** (*Block encoding*). Given a matrix  $A \in \mathbb{C}^{N \times N}$ , where  $N = 2^N$ , if there exist  $\alpha, \varepsilon \in \mathbb{R}_+$ , with  $\alpha \geq \|A\|_2$  and a  $(m+n)$ -qubit unitary matrix  $U_A$  so that

$$\|A - \alpha(\langle 0|_m \otimes I_N) U_A (|0\rangle_m \otimes I_N)\|_2 \leq \varepsilon$$

then  $U_A$  is called an  $(\alpha, m, \varepsilon)$ -block-encoding of  $A$ . In particular, when the block encoding is exact with  $\varepsilon = 0$ ,  $U_A$  is called an  $(\alpha, m)$ -block-encoding of  $A$ .

In order to implement the proposed method, we employ the elementary gates: Hadamard (H), CNOT and SWAP. Conventionally, the CNOT gate applied on two consecutive qubits, flips the state of the target qubit based on the state of the control qubit, and the corresponding operator is a  $4 \times 4$  block diagonal matrix where the first diagonal block is  $I_2$  and the second block is  $X$ . We now introduce the definition of the operator of CNOT( $a, b$ ), with  $b > a$ , where the parameters are the indices respectively of the control and target qubit, which may not be consecutive. The matrix representation is as follows

$$\text{CNOT}(a, b) = \begin{bmatrix} I_{2^{(b-a)}} & \\ & I_{2^{(b-a-1)}} \otimes X \end{bmatrix} \in \mathbb{R}^{2^{b-a+1} \times 2^{b-a+1}}. \quad (1)$$

Furthermore, the effect of the SWAP gate is to permute two consecutive registers. We denote by SWAP( $a, b$ ) the operator that permute two generic registers  $a$  and  $b$ , which may be not contiguous.

### 3. Main Result

In this section, we describe the algorithm for obtaining the block encoding operator  $U_A$  from the matrix state preparation circuit described by  $M_A$ . The quantum circuit is in Figure 1.

**Theorem 3.** Let  $A = (a_{ij}) \in \mathbb{C}^{N \times N}$ , with  $N = 2^n$  and  $\|A\|_F = 1$ , and let  $M_A$  be a matrix state preparation operator such that  $|A\rangle = M_A |0\rangle_{2n} = \sum_{i,j=0}^{N-1} a_{ij} |j\rangle_n |i\rangle_n$ . Then, there exists a quantum circuit  $U_A$  of constant depth that realises a  $(\sqrt{N}, 2n)$ -block encoding of  $A$ , introducing only  $n$  ancilla qubits.

*Proof.* Consider a  $2n$ -qubit initial state  $|A\rangle = M_A |0\rangle_{2n}$ . Operator  $M_A$  contains in the first column  $\text{vec}(A)$ .

In order to create the block-encoding, we append an ancillary register of  $n$  qubits, thus obtaining the state  $|\psi_1\rangle = |A\rangle \otimes |0\rangle_n = \sum_{i,j=0}^{N-1} a_{i,j} |j\rangle_n |i\rangle_n |0\rangle_n$ . This extended state corresponds to applying the operator  $U_{\text{init}} = M_A \otimes I_N$  to the  $3n$ -qubit initial state  $|0\rangle_{3n}$ . Each of the first  $N$  columns of  $U_{\text{init}}$  contain a copy of the columns of  $A$ . From now on we only visualize the first  $N$  columns of the operators since  $A$  will be encoded in the leading  $N \times N$  submatrix.

$$U_{\text{init}}(|0\rangle_{2n} \otimes I_N) = M_A \otimes I_N(|0\rangle_{2n} \otimes I_N) = \begin{bmatrix} A_0 \otimes I_N \\ A_1 \otimes I_N \\ \vdots \\ A_{N-1} \otimes I_N \end{bmatrix} \in \mathbb{C}^{N^3 \times N},$$

where the blocks  $A_j \otimes I_N \in \mathbb{C}^{N^2 \times N}$  have the following structure

$$A_j \otimes I_N = \begin{bmatrix} a_{0j} I_N \\ a_{1j} I_N \\ \vdots \\ a_{ij} I_N \\ \vdots \\ a_{N-1,j} I_N \end{bmatrix}, \quad \text{where} \quad a_{ij} I_N = \begin{bmatrix} a_{ij} & & & \\ & a_{ij} & & \\ & & \ddots & \\ & & & a_{ij} \end{bmatrix}.$$

To produce the desired block encoding, we first regroup each copy of the columns, and then move them in the top  $N$  positions of each column. We achieve this by defining a new operator  $U_S$ , whose action on the state vector is to swap the second and third register and can be described as:

$$U_S = \left( \prod_{a=0}^{n-1} I_{2^{a+n}} \otimes \text{SWAP}(a+n, a+2n) \otimes I_{2^{n-a-1}} \right).$$

Looking at the effect on the initial operator, we get that the first  $N$  columns become:

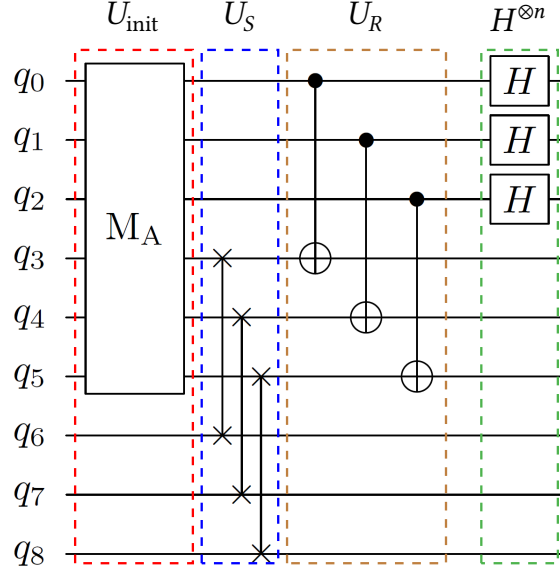
$$U_S U_{\text{init}}(|0\rangle_{2n} \otimes I_N) = \begin{bmatrix} I_N \otimes A_0 \\ I_N \otimes A_1 \\ \vdots \\ I_N \otimes A_{N-1} \end{bmatrix} \in \mathbb{C}^{N^3 \times N}.$$

The next step rearranges, inside each row-block, the  $N$  copies of the column  $A_j$  so that the useful copy is brought to the top. We achieve this with the operator  $U_R$ , which acts only on the first two registers (leaving the third untouched) as the bitwise controlled-NOT cascade

$$U_R = \left( \prod_{a=0}^{n-1} I_{2^a} \otimes \text{CNOT}(a, a+n) \otimes I_{2^{n-a-1}} \right) \otimes I_N,$$

that is,  $U_R : |i\rangle_n |k\rangle_n |\cdot\rangle_n \mapsto |i\rangle_n |k \oplus i\rangle_n |\cdot\rangle_n$ . Equivalently,  $U_R$  is block-diagonal in the first register: on the block selected by  $|i\rangle_n$  it applies

$$(U_R)_i = P_i \otimes I_N, \quad P_i := \bigotimes_{k=0}^{n-1} X^{i_{n-1-k}} = X^{i_{n-1}} \otimes \dots \otimes X^{i_0},$$



**Figure 1:** Block encoding circuit for the state vector produced by the matrix state preparation routine  $M_A$ , note that  $U_S$  and  $U_R$  are realized with depth 1 since the  $n$  operations act simultaneously on distinct qubits.

where  $i = \sum_{\ell=0}^{n-1} i_\ell 2^\ell$  with  $i_\ell \in \{0, 1\}$ , and we used  $X^0 = I$ ,  $X^1 = X$ . The factor  $P_i$  permutes the basis states of the second register according to  $|k\rangle \mapsto |k \oplus i\rangle$ , while  $I_N$  leaves the third register unchanged. Since  $(U_R)_i (I_N \otimes A_i) = P_i \otimes A_i$ , we have

$$U_R U_S U_{\text{init}} (|0\rangle_{2n} \otimes I_N) = \begin{bmatrix} P_0 \otimes A_0 \\ P_1 \otimes A_1 \\ \vdots \\ P_{N-1} \otimes A_{N-1} \end{bmatrix} \in \mathbb{C}^{N^3 \times N}.$$

In particular, in column  $m$  the copy of  $A_j$  lands in sub-block  $m \oplus j$ ; choosing  $m = j$  gives sub-block 0, so each  $A_j$  reaches the top row of its own row-block exactly in column  $j$ , meaning that  $(P_j \otimes A_j) |j\rangle = |0\rangle \otimes A_j$ .

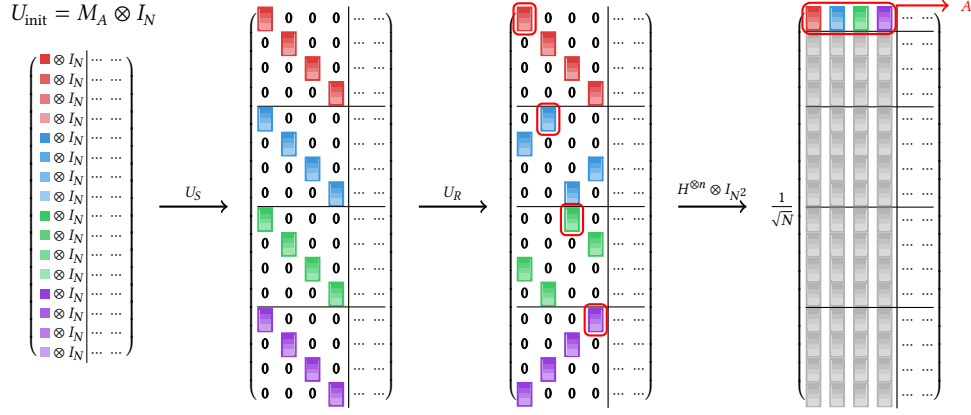
Finally, by applying  $H^{\otimes n} \otimes I_{N^2}$ , we sum the  $N$  row-blocks of size  $N^2 \times N^2$ , placing  $A$  in the leading  $N \times N$  principal sub-matrix.

$$(H^{\otimes n} \otimes I_{N^2}) U_R U_S U_{\text{init}} (|0\rangle_{2n} \otimes I_N) = \frac{1}{\sqrt{N}} \begin{bmatrix} A_0 & A_1 & \cdots & A_{N-1} \\ A_1 & A_0 & \cdots & A_{N-2} \\ \vdots & \vdots & \ddots & \vdots \\ A_{N-1} & A_{N-2} & \cdots & A_0 \\ \vdots & \vdots & \ddots & \vdots \\ \vdots & \vdots & \ddots & \vdots \end{bmatrix} = \frac{1}{\sqrt{N}} \begin{bmatrix} A \\ \star \end{bmatrix}$$

This gives a  $(\sqrt{N}, 2n)$ -block encoding of  $A$ . In Figure 2 we provide a pictorial representation of the steps of the algorithm, for a matrix of size  $N = 4$ .  $\square$

**Theorem 4. (Optimality)** Let  $A \in \mathbb{C}^{N \times N}$  with  $\|A\|_F = 1$ ,  $N = 2^n$ , and let  $M_A$  be a matrix state preparation operator for  $A$ . Any circuit that queries  $M_A$  once, where any gates applied before the query act only on ancilla qubits initialized independently of  $A$ , and that realizes an exact  $(\alpha, m)$ -block encoding of  $A$  for every  $M_A$ , satisfies  $\alpha \geq \sqrt{N}$ . The circuit of Theorem 3 attains  $\alpha = \sqrt{N}$  and is thus optimal in this class.

*Proof.* It suffices to prove this result for a subfamily of matrices. Take  $A = V/\sqrt{N}$  with  $V$  unitary, so that  $\|A\|_F = 1$  and  $|A\rangle = M_A |0\rangle_{2n} = \frac{1}{\sqrt{N}} \sum_{ij} v_{ij} |i\rangle |j\rangle$  is the (normalized)  $\text{vec}(V)$ . By Definition 2,



**Figure 2:** Construction of the  $(\sqrt{N}, 2n)$ -block encoding of  $A$ , illustrated for  $N = 4$  on the first  $N$  columns of each operator.  $U_S$  swaps the second and third qubit registers, producing a block-diagonal structure with  $N$  copies of each column of  $A$  along the diagonals.  $U_R$  applies a bitwise XOR on the second register, permuting the copies so that  $A_j$  reaches the top of the  $j$ -th row-block in column  $j$  (red boxes). The final Hadamard layer  $H^{\otimes n} \otimes I_{N^2}$  coherently sums the row-blocks, placing  $\frac{1}{\sqrt{N}} A$  in the leading submatrix (red box, right-most matrix).

measuring the ancillae of the block encoding on input  $|b\rangle$  yields the outcome  $0^{\otimes m}$  with probability  $p = \|\frac{A}{\alpha} |b\rangle\|^2 = \frac{1}{\alpha^2} \frac{1}{N} \|V|b\rangle\|^2 = \frac{1}{\alpha^2 N}$ , independent of  $V$  and  $|b\rangle$ . A single-query, input-independent circuit doing this is precisely a probabilistic storage-and-retrieval protocol for the unknown unitary  $V$  from a single use: the condition that gates preceding the query act independently of the input is exactly the technical hypothesis required by Theorem 2 of [20] to apply their optimality bound. By [20], any such protocol succeeds with probability at most  $1/N^2$ . Hence  $1/(\alpha^2 N) \leq 1/N^2$ , i.e.  $\alpha \geq \sqrt{N}$ . The full argument is given in the extended version.  $\square$

## 4. Conclusion

In this work, we introduced a direct method for converting a matrix state vector into a block encoding using only elementary one- and two-qubit gates, achieving constant circuit depth and the theoretically optimal normalization factor of  $\sqrt{N}$ . Compared to prior approaches [17, 18, 19], our method avoids any intermediate decomposition, for example through Pauli decomposition, and reduces circuit depth. A future direction is to investigate how the inverse transformation can be carried out, namely how a block-encoded operator can be converted back into a matrix state vector directly without intermediate representations.

## Funding

Partial support is provided by the INdAM - GNCS project “Algebra Lineare Quantistica, State Preparation e Compilazione di Circuiti Quantistici”, CUP E53C25002010001, and by the Italian Project Fondo Italiano per la Scienza FIS00001966 “MIMOSA”. G. Del Corso is also partially supported by the U.S. Department of Energy, Office of Science, National Quantum Information Science Research Centers, Superconducting Quantum Materials and Systems Center (SQMS) under contract number No. 89243024CSC000002.

## Declaration on Generative AI

During the preparation of this work, the authors used Claude Sonnet 4.6 in order to: Grammar and spelling check, help in proof of optimality of the result stated in Theorem 4 and help generating the pictorial representation of the circuit. After using these tool, the authors reviewed and edited the content as needed and take full responsibility for the publication’s content.

## References

- [1] M. Plesch, Č. Brukner, Quantum-state preparation with universal gate decompositions, *Physical Review A—Atomic, Molecular, and Optical Physics* 83 (2011) 032302.
- [2] A. W. Harrow, A. Hassidim, S. Lloyd, Quantum algorithm for linear systems of equations, *Physical review letters* 103 (2009) 150502.
- [3] I. Kerenidis, A. Prakash, Quantum recommendation systems, *arXiv preprint arXiv:1603.08675* (2016).
- [4] A. Bernasconi, A. Berti, G. M. Del Corso, A. Poggiali, Quantum subroutine for efficient matrix multiplication, *IEEE Access* 12 (2024) 116274–116284.
- [5] G. Antonioli, A. Berti, A. Poggiali, A. Bernasconi, G. M. Del Corso, Outlier detection and other applications of quantum matrix multiplication, in: *2025 IEEE International Parallel and Distributed Processing Symposium Workshops (IPDPSW)*, IEEE, 2025, pp. 509–518.
- [6] M. Möttönen, J. J. Vartiainen, V. Bergholm, M. M. Salomaa, Transformation of quantum states using uniformly controlled rotations, *Quantum Information & Computation* 5 (2005) 467–473. *arXiv:quant-ph/0407010*.
- [7] X.-M. Zhang, T. Li, X. Yuan, Quantum state preparation with optimal circuit depth: Implementations and applications, *Physical Review Letters* 129 (2022) 230504. doi:10.1103/PhysRevLett.129.230504. *arXiv:2201.11495*.
- [8] A. Berti, F. Ghisoni, Efficient quantum state preparation with bucket brigade qram, *arXiv preprint arXiv:2510.16149* (2025).
- [9] A. Berti, F. Ghisoni, Efficient complex-valued state preparation on bucket brigade qram, *arXiv preprint arXiv:2604.25644* (2026).
- [10] A. Gilyén, Y. Su, G. H. Low, N. Wiebe, Quantum singular value transformation and beyond: exponential improvements for quantum matrix arithmetics, in: *Proceedings of the 51st annual ACM SIGACT symposium on theory of computing*, 2019, pp. 193–204.
- [11] D. Camps, R. Van Beeumen, Fable: Fast approximate quantum circuits for block-encodings, in: *2022 IEEE International Conference on Quantum Computing and Engineering (QCE)*, IEEE, 2022, pp. 104–113.
- [12] D. Camps, L. Lin, R. Van Beeumen, C. Yang, Explicit quantum circuits for block encodings of certain sparse matrices, *SIAM Journal on Matrix Analysis and Applications* 45 (2024) 801–827.
- [13] C. Sünderhauf, E. Campbell, J. Camps, Block-encoding structured matrices for data input in quantum computing, *Quantum* 8 (2024) 1226.
- [14] L.-C. Wan, C.-H. Yu, S.-J. Pan, S.-J. Qin, F. Gao, Q.-Y. Wen, Block-encoding-based quantum algorithm for linear systems with displacement structures, *Physical Review A* 104 (2021) 062414.
- [15] G. Antonioli, P. Boito, G. M. D. Corso, M. Porcelli, Quantum block encoding for semiseparable matrices, 2026. URL: <https://arxiv.org/abs/2603.19130>. *arXiv:2603.19130*.
- [16] H. Li, H. Ni, L. Ying, On efficient quantum block encoding of pseudo-differential operators, *Quantum* 7 (2023) 1031.
- [17] L. M. Yosef, H. Avron, On encoding matrices using quantum circuits, *arXiv preprint arXiv:2510.20030* (2025).
- [18] L. Mor-Yosef, S. Ubaru, L. Horesh, H. Avron, Multivariate trace estimation using quantum state space linear algebra, *SIAM Journal on Matrix Analysis and Applications* 46 (2025) 172–209.
- [19] X. Li, P.-L. Zheng, C. Pan, F. Wang, C. Cui, X. Lu, Faster quantum subroutine for matrix chain multiplication via chebyshev approximation, *Scientific Reports* 15 (2025) 28559.
- [20] M. Sedlák, A. Bisio, M. Ziman, Optimal probabilistic storage and retrieval of unitary channels, *Physical Review Letters* 122 (2019) 170502. doi:10.1103/PhysRevLett.122.170502. *arXiv:1809.04552*.