

Quantics Tensor Trains for Vibrational Eigenstates

Imaginary-time propagation, physical-core Fourier transforms, block Rayleigh–Ritz, and Morse-oscillator benchmarks

Victor S. Batista

Contents

1	Purpose and scope	3
2	Why tensor trains are useful	3
2.1	The exponential storage problem	3
2.2	Tensor-train representation	3
2.3	Quantics tensor trains	4
2.4	Tiny TT-SVD example	4
3	Core QTT algebra used by the notebook	4
3.1	Rounding and rank control	4
3.2	Inner products	5
4	Physical-core Fourier transform	5
4.1	Why a Fourier transform is needed	5
4.2	Round-trip validation	6
5	Morse Hamiltonian and numerical grid	6
5.1	One-dimensional Morse modes	6
6	Constructing the potential and kinetic operators in QTT form	7
6.1	Structured potential construction	7
6.2	Kinetic operator	7
7	Hamiltonian action without dense reconstruction	8
8	Block orthogonalization and Rayleigh–Ritz	8
9	Analytical Morse reference spectrum	9
9.1	Analytical wavefunctions	10
10	Local Fourier-grid eigenstates as numerical seeds	10
11	Imaginary-time propagation	11
11.1	Filtering principle	11
11.2	Second-order split operator	11
11.3	Executed propagation schedule	12

12 Validation and state assignment	12
12.1 Dense validation branch	12
12.2 Exact-state matching	12
12.3 Hamiltonian residual	13
13 QTT-native one-coordinate marginals	13
14 Convergence and compression diagnostics	16
15 What the results demonstrate	17
16 How the code generalizes to coupled systems	18
17 The dormant two-dimensional heat-map branch	18
18 Scaling and practical limitations	19
19 Code architecture: how the notebook pieces fit together	19
20 Recommended modifications for further study	20
21 Summary	20
A Complete code extracted from the notebook	21

1 Purpose and scope

This tutorial explains the complete numerical strategy implemented in the supplied notebook for computing multiple low-energy vibrational eigenstates without storing the full multidimensional wavefunction. The central idea is to represent functions on binary grids by *quantics tensor trains* (QTTs), apply coordinate- and momentum-space operations directly to those compressed objects, and extract a block of low-energy states by imaginary-time filtering combined with repeated orthogonalization and Rayleigh–Ritz diagonalization.

The Hamiltonian implemented in the notebook is

$$\hat{H} = \sum_{d=1}^n \left[\frac{\hat{p}_d^2}{2m_d} + D_{e,d} \left(1 - e^{-a_d(\hat{x}_d - x_{e,d})} \right)^2 \right] + V_{\text{coup}}(x_1, \dots, x_n). \quad (1)$$

The code supports uncoupled Morse modes, pairwise coupling in natural Morse coordinates, and an arbitrary custom many-body potential sampled by TT-cross.

Important note about the saved notebook run

The opening markdown in the notebook describes a two-mode default benchmark, but the actually executed parameter cell sets

$$\text{N_OSC} = 3, \quad \text{NSTATES} = 20, \quad \text{COUPLING_MODEL} = \text{"none"}.$$

Therefore, the numerical outputs and all saved figures in the supplied notebook correspond to **three uncoupled Morse oscillators**. The two-dimensional heat-map routine is present in the code, but its saved output is only the message that heat maps require `N_OSC == 2`; no heat-map image was generated in this run. This document includes **all eight figures that were actually produced** and all displayed numerical tables.

2 Why tensor trains are useful

2.1 The exponential storage problem

Suppose a wavefunction depends on n coordinates and coordinate d is sampled on N_d grid points. A dense representation requires

$$N_{\text{dense}} = \prod_{d=1}^n N_d \quad (2)$$

complex amplitudes. If $N_d = N$, this grows as N^n . The three-mode demonstration is still manageable as a dense array, but the code is written to avoid dependence on dense reconstruction because the same representation is intended to extend to much larger n .

2.2 Tensor-train representation

A tensor train factorizes a tensor into a chain of low-order cores,

$$\Psi_{i_1 \dots i_D} = \sum_{\alpha_1 \dots \alpha_{D-1}} G_{1,i_1,\alpha_1}^{(1)} G_{\alpha_1,i_2,\alpha_2}^{(2)} \dots G_{\alpha_{D-1},i_D,1}^{(D)}. \quad (3)$$

The intermediate dimensions r_k are the TT ranks (bond dimensions). If the physical dimension of each core is q and the ranks remain of order r , storage scales approximately as $O(Dqr^2)$ rather than exponentially in D .

2.3 Quantics tensor trains

QTT applies the TT idea to the binary digits of an ordinary grid index. For a one-dimensional grid of size $N = 2^L$,

$$i = \sum_{k=1}^L 2^{L-k} b_k, \quad b_k \in \{0, 1\}. \quad (4)$$

Thus a vector f_i is reinterpreted as an L -index binary tensor $f(b_1, \dots, b_L)$ and compressed as a TT with physical dimension 2 at every core. For multiple coordinates, the notebook orders the QTT cores as

$$[\text{bits of } x_1] [\text{bits of } x_2] \dots [\text{bits of } x_n].$$

With seven bits per coordinate and three coordinates, the saved calculation has 21 binary cores.

2.4 Tiny TT-SVD example

The notebook begins with the vector

$$v = (1, 2, 3, 4), \quad (5)$$

reshaped into the 2×2 matrix

$$M = \begin{pmatrix} 1 & 2 \\ 3 & 4 \end{pmatrix}. \quad (6)$$

Its singular values are 5.4649857 and 0.36596619, so the exact bond dimension is 2. The two TT cores have shapes $(1, 2, 2)$ and $(2, 2, 1)$, and contracting them reconstructs the original vector exactly. The full routine `qtt_from_dense()` recursively repeats this SVD splitting over all binary digits.

3 Core QTT algebra used by the notebook

The notebook deliberately implements the most important TT operations explicitly on `tntorch.Tensor` cores. This makes the handling of complex tensors transparent once Fourier transforms are introduced.

3.1 Rounding and rank control

After addition or elementwise multiplication, TT ranks increase. Exact TT addition forms block-structured cores and roughly adds ranks, whereas exact Hadamard products multiply ranks. The notebook therefore applies TT rounding after these operations. The truncation rule retains the smallest number of singular values whose discarded tail satisfies the requested Frobenius-norm tolerance, subject to a hard rank cap.

Table 1: Principal compression and TT-cross parameters in the executed notebook.

Parameter	Value	Role
QTT_EPS	1.0×10^{-12}	Rounding after QTT algebra
QFT_EPS	2.0×10^{-13}	Rounding after Fourier transforms
QTT_RMAX	96	Hard rank ceiling
CROSS_EPS	1.0×10^{-11}	TT-cross target for custom/coupled nonlinear maps
CROSS_RMAX	96	TT-cross rank ceiling
CROSS_MAX_ITER	40	TT-cross iterations
CROSS_VAL_SIZE	4000	TT-cross validation sample size

3.2 Inner products

For QTTs a and b , the algebraic inner product is contracted core by core,

$$\langle a|b \rangle = \sum_i a_i^* b_i. \quad (7)$$

The only intermediate object is a small environment matrix whose dimensions are TT ranks. Physical quadrature weights are inserted separately where needed. In the present uniform-grid problem,

$$\Delta V = \prod_{d=1}^n \Delta x_d. \quad (8)$$

4 Physical-core Fourier transform

4.1 Why a Fourier transform is needed

The potential is diagonal in coordinate space, but the kinetic energy is diagonal in momentum space. The notebook therefore needs a multidimensional Fourier transform of a QTT wavefunction.

Instead of implementing a gate-by-gate binary QFT circuit, the notebook uses a pedagogically transparent *unfold-FFT-refold* procedure for each physical coordinate:

1. Contract the L_d binary cores of coordinate d into one physical core of shape (R_L, N_d, R_R) .
2. Apply an ordinary 1D FFT along the physical index N_d , independently for every pair of bond indices.
3. Split the transformed physical core back into binary cores by exact recursive SVD, concatenate all coordinate blocks, and round the complete QTT once.

The FFT uses `norm="ortho"`, making the forward and inverse transforms unitary.

The high-level implementation is:

Listing 1: High-level QTT Fourier transform from the notebook.

```

1  def qtt_qft(
2      qtt,
3      bits_per_dim,
4      inverse=False,
5      eps=QTT_EPS,
6      rmax=QTT_RMAX,
7      norm="ortho"
8  ):
9      """
10     Multidimensional QFT/IQFT of a QTT wavefunction, tying the three steps
11     together:
12
13         QTT --(unfold)--> one physical TT core per coordinate
14         --(FFT each core's columns)--> transformed physical cores
15         --(fold + round)--> binary QTT
16
17     Set inverse=True for the inverse transform (used to come back from
18     momentum space to coordinate space after applying the kinetic
19     propagator/operator).
20     """
21     physical_cores = qtt_unfold_to_physical_tt(
22         qtt,
23         bits_per_dim

```

```

24 )
25
26     transformed = fft_physical_tt_cores(
27         physical_cores,
28         inverse=inverse,
29         norm=norm
30     )
31
32     return physical_tt_to_qtt(
33         transformed,
34         bits_per_dim,
35         eps=eps,
36         rmax=rmax
37     )

```

4.2 Round-trip validation

The notebook constructs a separable test function directly as a QTT, transforms it to momentum space and back, and computes

$$\epsilon_{\text{QFT}} = \frac{\|\text{QFT}^{-1}\text{QFT}f - f\|}{\|f\|}. \quad (9)$$

For the saved run,

$$\boxed{\epsilon_{\text{QFT}} = 7.721 \times 10^{-15}}, \quad (10)$$

which is essentially machine precision for double precision. The real-space test QTT and transformed test QTT both have maximum rank 8; the coefficient counts are 630 and 864, respectively.

5 Morse Hamiltonian and numerical grid

5.1 One-dimensional Morse modes

Each coordinate uses the unshifted Morse potential

$$V_d(x) = D_{e,d} \left(1 - e^{-a_d(x-x_{e,d})}\right)^2, \quad (11)$$

with minimum zero and dissociation limit $D_{e,d}$. The natural Morse coordinate used for pair coupling is

$$s_d(x_d) = 1 - e^{-a_d(x_d-x_{e,d})}. \quad (12)$$

For pairwise coupling,

$$V_{\text{coup}} = \sum_{d < e} \kappa_{de} s_d(x_d) s_e(x_e). \quad (13)$$

The saved run sets the coupling model to **none**.

Table 2: Executed three-mode Morse parameters and Fourier-grid discretization.

d	m	D_e	a	x_e	L	N	$x_{\min}$	$x_{\max}$	Δx
0	1.00	20.0	0.80	0.0	7	128	-6.0	18.0	0.1875
1	1.05	18.5	0.77	0.0	7	128	-6.0	18.0	0.1875
2	1.10	17.0	0.74	0.0	7	128	-6.0	18.0	0.1875

The run therefore uses

$$\text{BITS_PER_DIM} = [7, 7, 7], \quad (14)$$

$$L_{\text{total}} = 21, \quad (15)$$

$$N_{\text{dense}} = 128^3 = 2,097,152, \quad (16)$$

$$\Delta V = 6.591797 \times 10^{-3}. \quad (17)$$

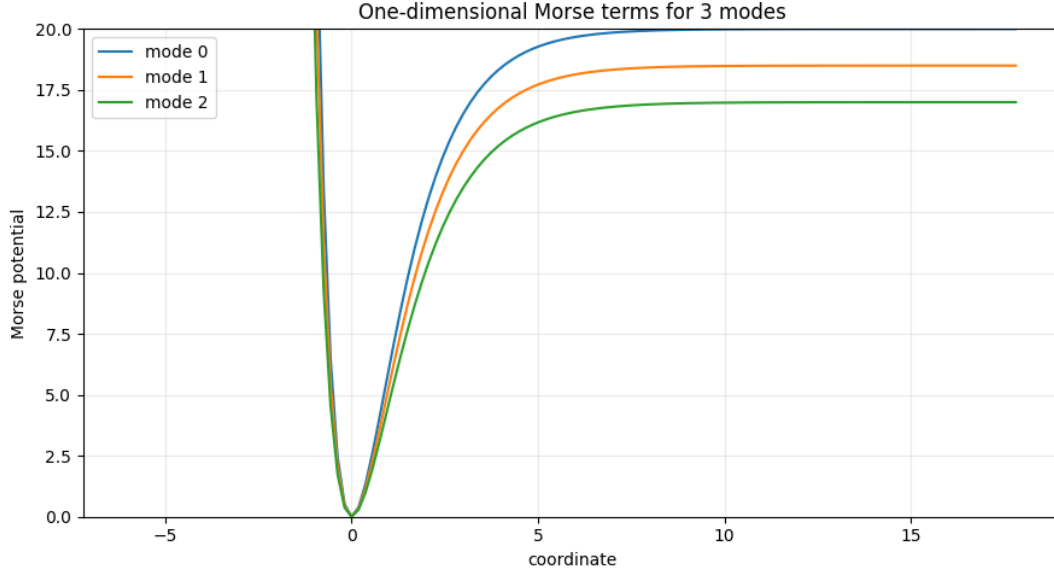


Figure 1: All three one-dimensional Morse potentials in the executed notebook. They share $x_e = 0$ but differ in mass, dissociation energy, and range parameter. The asymptotic plateaus are $D_e = 20.0$, 18.5 , and 17.0 .

6 Constructing the potential and kinetic operators in QTT form

6.1 Structured potential construction

For the uncoupled and pair-coupled Morse models, the notebook does not approximate the potential by TT-cross. Instead it exploits algebraic separability. For the uncoupled case,

$$V(\mathbf{x}) = \sum_d V_d(x_d), \quad (18)$$

and each V_d is embedded into the full QTT by tensoring it with all-ones QTTs on the other coordinates. For pair coupling, each $s_i(x_i)s_j(x_j)$ term is likewise a tensor product of one-dimensional factors. Only the custom many-body branch uses TT-cross to learn the full potential.

For the saved run, the structured potential has maximum QTT rank 4 and 368 core coefficients. Because no TT-cross approximation is used in this branch, its approximation error is controlled only by QTT rounding.

6.2 Kinetic operator

The momentum grid is generated from `torch.fft.fftfreq`, with

$$p = 2\pi\hbar f, \quad (19)$$

and

$$T(\mathbf{p}) = \sum_d \frac{p_d^2}{2m_d}. \quad (20)$$

The full momentum-space kinetic diagonal is again assembled as a sum of embedded one-dimensional QTTs. In the saved run, it has maximum rank 3 and 314 coefficients.

7 Hamiltonian action without dense reconstruction

For any QTT state ψ , the notebook evaluates

$$\hat{H}\psi = V\psi + \mathcal{F}^{-1}[T(\mathbf{p})\mathcal{F}\psi]. \quad (21)$$

The central routine is:

Listing 2: QTT-native Hamiltonian action.

```

1 def apply_H_qtt(psi_qtt, eps=QTT_EPS, qft_eps=QFT_EPS):
2     Vpsi = qtt_hadamard(Vtot_qtt, psi_qtt, eps=eps, rmax=QTT_RMAX)
3
4     psi_p = qtt_qft(
5         psi_qtt, BITS_PER_DIM, inverse=False,
6         eps=qft_eps, rmax=QTT_RMAX
7     )
8
9     Tpsi_p = qtt_hadamard(Ttot_qtt, psi_p, eps=eps, rmax=QTT_RMAX)
10
11     Tpsi = qtt_qft(
12         Tpsi_p, BITS_PER_DIM, inverse=True,
13         eps=qft_eps, rmax=QTT_RMAX
14     )
15
16     return qtt_add(Vpsi, Tpsi, eps=eps, rmax=QTT_RMAX)

```

The Rayleigh quotient is computed as

$$E[\psi] = \frac{\langle \psi | \hat{H} | \psi \rangle}{\langle \psi | \psi \rangle}, \quad (22)$$

where the common volume factor cancels. Normalization itself uses the physical quadrature norm $\Delta V \langle \psi | \psi \rangle = 1$.

8 Block orthogonalization and Rayleigh–Ritz

Independent imaginary-time propagation would drive every trial vector toward the ground state. To obtain many states simultaneously, the notebook propagates a block and repeatedly restores orthogonality.

For states $\{\psi_i\}_{i=1}^M$, the overlap matrix is

$$S_{ij} = \Delta V \langle \psi_i | \psi_j \rangle. \quad (23)$$

A symmetric Löwdin transformation constructs an orthonormal block,

$$\phi_j = \sum_i \psi_i [S^{-1/2}]_{ij}. \quad (24)$$

Only the $M \times M$ matrix S is dense; the states remain QTTs.

The projected Hamiltonian is

$$H_{ij}^{\text{sub}} = \Delta V \langle \phi_i | \hat{H} | \phi_j \rangle. \quad (25)$$

Diagonalizing H^{sub} gives the Ritz energies and the rotation of the QTT block:

Listing 3: Block Rayleigh–Ritz update.

```

1 def rayleigh_ritz_qtt(states):
2     """QTT block orthogonalization + projected-Hamiltonian diagonalization."""
3     states = orthonormalize_qtt_block(states, passes=2)
4     M = len(states)
5
6     Hstates = [apply_H_qtt(state) for state in states]
7     Hsub = torch.zeros((M, M), dtype=torch.complex128, device=device)
8
9     for i in range(M):
10        for j in range(M):
11            Hsub[i, j] = dvol * qtt_inner(states[i], Hstates[j])
12
13    Hsub = 0.5 * (Hsub + Hsub.conj().T)
14    evals, U = torch.linalg.eigh(Hsub)
15
16    rotated = rotate_qtt_block(states, U)
17    rotated = [qtt_normalize(state) for state in rotated]
18
19    return rotated, evals.real

```

This is the key step that converts imaginary-time filtering into a stable multi-eigenstate method.

9 Analytical Morse reference spectrum

For one Morse oscillator, the notebook defines

$$\lambda_d = \frac{\sqrt{2m_d D_{e,d}}}{\hbar a_d}, \quad (26)$$

and the bound-state energies

$$E_{n,d} = D_{e,d} - \frac{\hbar^2 a_d^2}{2m_d} \left(\lambda_d - n - \frac{1}{2} \right)^2, \quad \lambda_d - n - \frac{1}{2} > 0. \quad (27)$$

The three executed modes have $\lambda = (7.905694, 8.094775, 8.264262)$ and each supports eight bound levels according to this criterion.

For uncoupled modes, exact multidimensional states are products and energies add,

$$E_{n_1 n_2 n_3} = E_{n_1,1} + E_{n_2,2} + E_{n_3,3}. \quad (28)$$

The notebook uses a heap-based k -smallest-sums algorithm to generate only the requested lowest 20 product levels instead of enumerating all possible tuples.

Table 3: Lowest 20 analytical uncoupled product levels used as the reference spectrum.

Rank	Quantum numbers	E_{exact}	E_{QTT}	$ \Delta E $
0	(0, 0, 0)	6.659486316765	6.659486318053	1.288e-09
1	(0, 0, 1)	10.275768058601	10.275768060342	1.741e-09
2	(0, 1, 0)	10.665669315824	10.665669317948	2.124e-09
3	(1, 0, 0)	11.079130573034	11.079130575842	2.808e-09
4	(0, 0, 2)	13.394231618620	13.394231620733	2.113e-09
5	(0, 2, 0)	14.107185648216	14.107185651099	2.883e-09
6	(0, 1, 1)	14.281951057661	14.281951060154	2.493e-09

Rank	Quantum numbers	E_{exact}	E_{QTT}	$ \Delta E $
7	(1, 0, 1)	14.695412314871	14.695412318056	3.185e-09
8	(2, 0, 0)	14.858774829303	14.858774833587	4.284e-09
9	(1, 1, 0)	15.085313572093	15.085313575650	3.557e-09
10	(0, 0, 3)	16.014876996821	16.014876999249	2.428e-09
11	(0, 3, 0)	16.984035313942	16.984035317392	3.450e-09
12	(0, 1, 2)	17.400414617679	17.400414620559	2.880e-09
13	(0, 2, 1)	17.723467390053	17.723467393313	3.260e-09
14	(1, 0, 2)	17.813875874889	17.813875878483	3.594e-09
15	(3, 0, 0)	17.998419085573	17.998419090909	5.336e-09
16	(0, 0, 4)	18.137704193203	18.137704195775	2.572e-09
17	(2, 0, 1)	18.475056571140	18.475056575817	4.677e-09
18	(1, 2, 0)	18.526829904486	18.526829908837	4.351e-09
19	(1, 1, 1)	18.701595313930	18.701595317883	3.953e-09

9.1 Analytical wavefunctions

The one-dimensional reference wavefunctions are generated from

$$\psi_n(x) \propto z^s e^{-z/2} L_n^{(2s)}(z), \quad z = 2\lambda e^{-a(x-x_e)}, \quad s = \lambda - n - \frac{1}{2}, \quad (29)$$

then normalized on the numerical grid. Each one-dimensional factor is converted to QTT form, and the multidimensional analytical product state is formed by concatenating independent QTT core lists. Thus even the exact reference states need not be assembled as dense three-dimensional arrays.

10 Local Fourier-grid eigenstates as numerical seeds

The initial block is not built from the analytical Morse wavefunctions. Instead, the notebook solves a small one-dimensional Fourier-grid Hamiltonian for each mode,

$$H_d^{\text{FGH}} = F_d^\dagger \text{diag}[T_d(p)] F_d + \text{diag}[V_d(x)]. \quad (30)$$

Each F_d is the unitary discrete FFT matrix. These are only 128×128 eigenproblems, so they are inexpensive even when the full product grid becomes too large. Products of these local FGH eigenvectors provide the initial QTT block.

The notebook explicitly checks the local grid accuracy before the multidimensional calculation:

Table 4: One-dimensional Fourier-grid Hamiltonian energies versus analytical Morse energies.

Mode	n	E_{exact}	$E_{\text{1D FGH}}$	Absolute error
0	0	2.449822	2.449822	7.150e-14
0	1	6.869466	6.869466	3.464e-14
0	2	10.649111	10.649111	4.459e-13
0	3	13.788755	13.788755	2.689e-12
1	0	2.214841	2.214841	2.842e-14
1	1	6.221024	6.221024	5.063e-14
1	2	9.662541	9.662541	2.203e-13
1	3	12.539390	12.539390	6.786e-13
2	0	1.994823	1.994823	1.288e-14
2	1	5.611104	5.611104	2.842e-14
2	2	8.729568	8.729568	3.553e-15
2	3	11.350213	11.350213	4.796e-14
2	4	13.473041	13.473041	1.670e-13

The largest displayed local error is only a few 10^{-12} , so the grid discretization of the low-lying levels is already much more accurate than the final QTT residual norm. This cleanly separates ordinary grid error from errors introduced by repeated QTT transformations, truncations, and propagation.

The initial product seeds have maximum rank 8. Because the executed problem is uncoupled, these FGH product states are already extremely close to eigenstates of the multidimensional discretized Hamiltonian; the initial Rayleigh–Ritz energies already match the analytical product spectrum to many digits.

11 Imaginary-time propagation

11.1 Filtering principle

Expanding an initial state in exact eigenvectors,

$$|\psi(0)\rangle = \sum_n c_n |n\rangle, \quad (31)$$

its imaginary-time evolution is

$$|\psi(\tau)\rangle = e^{-\tau\hat{H}} |\psi(0)\rangle = \sum_n c_n e^{-\tau E_n} |n\rangle. \quad (32)$$

Higher-energy components are exponentially suppressed. For a block of states, orthogonalization and Rayleigh–Ritz after each step prevent all vectors from collapsing onto the same ground state.

11.2 Second-order split operator

The notebook uses Strang splitting,

$$e^{-\Delta\tau\hat{H}} = e^{-\Delta\tau V/2} e^{-\Delta\tau T} e^{-\Delta\tau V/2} + O(\Delta\tau^3). \quad (33)$$

For the uncoupled run, $V = \sum_d V_d$ and the potential propagator factorizes exactly,

$$e^{-\Delta\tau V/2} = \bigotimes_d e^{-\Delta\tau V_d/2}. \quad (34)$$

The kinetic propagator always factorizes in momentum space,

$$e^{-\Delta\tau T} = \bigotimes_d e^{-\Delta\tau p_d^2/(2m_d)}. \quad (35)$$

For coupled potentials, the potential propagator is instead obtained by applying the scalar map $v \mapsto e^{-\Delta\tau v/2}$ to the potential QTT using TT-cross.

The one-state step is:

Listing 4: One second-order imaginary-time step in QTT form.

```

1 def imaginary_time_step_qtt(psi_qtt, Uv_qtt, Ut_qtt):
2     # Second-order Strang splitting:
3     # exp(-dt H) = exp(-dt V/2) exp(-dt T) exp(-dt V/2) + O(dt^3)
4     psi = qtt_hadamard(Uv_qtt, psi_qtt, eps=QTT_EPS, rmax=QTT_RMAX)
5     psi = qtt_qft(psi, BITS_PER_DIM, inverse=False, eps=QFT_EPS, rmax=QTT_RMAX)
6     psi = qtt_hadamard(Ut_qtt, psi, eps=QTT_EPS, rmax=QTT_RMAX)
7     psi = qtt_qft(psi, BITS_PER_DIM, inverse=True, eps=QFT_EPS, rmax=QTT_RMAX)
8     psi = qtt_hadamard(Uv_qtt, psi, eps=QTT_EPS, rmax=QTT_RMAX)
9     return qtt_normalize(psi)

```

11.3 Executed propagation schedule

The uncoupled schedule contains 12 block iterations:

Table 5: Imaginary-time schedule in the saved uncoupled run.

Stage	$\Delta\tau$	Number of steps
1	5.0×10^{-3}	3
2	1.0×10^{-3}	4
3	2.0×10^{-4}	5

After every step, all 20 propagated states undergo block Rayleigh–Ritz. The maximum rank is 15 at every recorded iteration.

12 Validation and state assignment

12.1 Dense validation branch

The notebook contains an optional dense multidimensional FGH validation when the total grid contains at most 200,000 points. The executed grid contains 2,097,152 points, so this branch is skipped. This is intentional: the scalable QTT calculation is unaffected by the absence of dense validation.

12.2 Exact-state matching

For the uncoupled case the notebook constructs the squared-overlap matrix

$$O_{ij} = \frac{|\Delta V \langle \Psi_i^{\text{exact}} | \Psi_j^{\text{QTT}} \rangle|^2}{(\Delta V \langle \Psi_i^{\text{exact}} | \Psi_i^{\text{exact}} \rangle)(\Delta V \langle \Psi_j^{\text{QTT}} | \Psi_j^{\text{QTT}} \rangle)}. \quad (36)$$

The Hungarian algorithm finds the globally optimal one-to-one assignment between analytical and numerical states. This is more robust than matching only by sorted energy, especially near degeneracies.

12.3 Hamiltonian residual

For each state, the notebook also evaluates

$$r_j = \frac{\|\hat{H}\psi_j - E_j\psi_j\|}{\|\psi_j\|}. \quad (37)$$

Energy error, fidelity, and residual norm probe different aspects of accuracy. In particular, a Rayleigh quotient can be extremely accurate even when a small vector residual remains; the residual is therefore an important complementary convergence diagnostic.

The summary printed by the notebook is

$$\max_j |E_j^{\text{QTT}} - E_j^{\text{exact}}| = 5.336 \times 10^{-9}, \quad (38)$$

$$\min_j O_{jj} = 0.99999999755, \quad (39)$$

$$\max_j r_j = 3.798 \times 10^{-4}. \quad (40)$$

Every matched state is therefore essentially identical to its analytical product reference by overlap, while the QTT coefficient counts remain between roughly one and a few thousand for a dense grid containing more than two million amplitudes.

13 QTT-native one-coordinate marginals

For dimensions higher than two, full wavefunction heat maps become both visually unhelpful and potentially expensive. The notebook instead computes

$$P_d(x_d) = \int \prod_{e \neq d} dx_e |\Psi(x_1, \dots, x_n)|^2. \quad (41)$$

It first forms the QTT of $|\Psi|^2$, unfolds each physical coordinate, contracts every coordinate except d , multiplies by the missing volume factor, removes tiny negative roundoff, and renormalizes the resulting one-dimensional probability distribution.

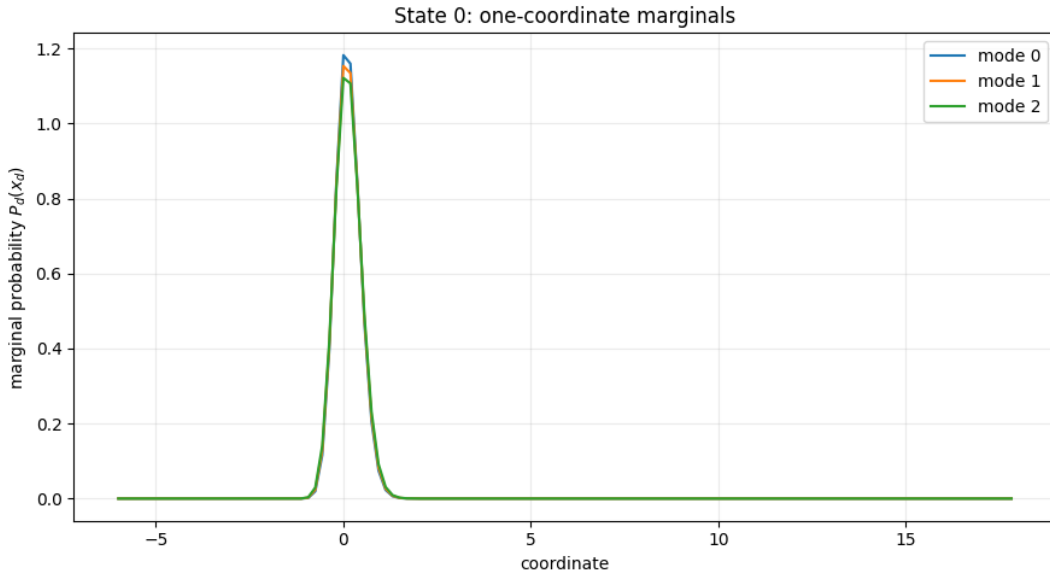


Figure 2: One-coordinate probability marginals of state 0, identified with $(n_0, n_1, n_2) = (0, 0, 0)$. All three modes show the single-peaked ground-state probability distribution near the Morse minimum.

Table 6: Complete final state-comparison table displayed by the notebook. The saved pandas display rounds the individual fidelities to 1.000000; the separately printed minimum fidelity is 0.9999999999755.

Exact rank	Num. index	Quantum numbers	E_{exact}	E_{QTT}	$ \Delta E $	Relative error	Fidelity	Residual	Max rank	QTT coeffs.
0	0	(0, 0, 0)	6.659486	6.659486	1.288e-09	1.935e-10	1.000000	0.000185	10	1128
1	1	(0, 0, 1)	10.275768	10.275768	1.740e-09	1.694e-10	1.000000	0.000207	10	1104
2	2	(0, 1, 0)	10.665669	10.665669	2.124e-09	1.991e-10	1.000000	0.000229	8	1024
3	3	(1, 0, 0)	11.079131	11.079131	2.808e-09	2.535e-10	1.000000	0.000270	10	1218
4	4	(0, 0, 2)	13.394232	13.394232	2.113e-09	1.578e-10	1.000000	0.000228	12	1678
5	5	(0, 2, 0)	14.107186	14.107186	2.883e-09	2.044e-10	1.000000	0.000268	10	1398
6	6	(0, 1, 1)	14.281951	14.281951	2.493e-09	1.746e-10	1.000000	0.000246	11	1538
7	7	(1, 0, 1)	14.695412	14.695412	3.185e-09	2.167e-10	1.000000	0.000284	13	1850
8	8	(2, 0, 0)	14.858775	14.858775	4.283e-09	2.883e-10	1.000000	0.000337	12	1660
9	9	(1, 1, 0)	15.085314	15.085314	3.556e-09	2.357e-10	1.000000	0.000300	11	1282
10	10	(0, 0, 3)	16.014877	16.014877	2.428e-09	1.516e-10	1.000000	0.000243	13	1948
11	11	(0, 3, 0)	16.984035	16.984035	3.450e-09	2.031e-10	1.000000	0.000295	11	1622
12	12	(0, 1, 2)	17.400415	17.400415	2.879e-09	1.655e-10	1.000000	0.000263	13	1876
13	13	(0, 2, 1)	17.723467	17.723467	3.260e-09	1.839e-10	1.000000	0.000283	12	1888
14	14	(1, 0, 2)	17.813876	17.813876	3.594e-09	2.017e-10	1.000000	0.000299	15	2272
15	15	(3, 0, 0)	17.998419	17.998419	5.336e-09	2.965e-10	1.000000	0.000380	15	2302
16	16	(0, 0, 4)	18.137704	18.137704	2.572e-09	1.418e-10	1.000000	0.000250	15	2420
17	17	(2, 0, 1)	18.475057	18.475057	4.677e-09	2.532e-10	1.000000	0.000349	14	2284
18	18	(1, 2, 0)	18.526830	18.526830	4.351e-09	2.348e-10	1.000000	0.000331	12	1748
19	19	(1, 1, 1)	18.701595	18.701595	3.953e-09	2.114e-10	1.000000	0.000313	13	2086

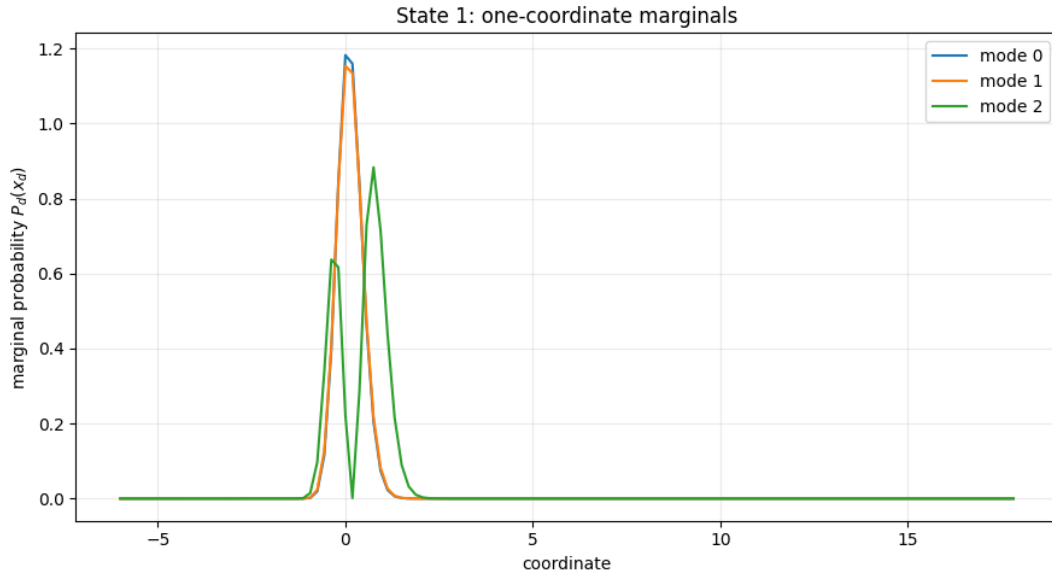


Figure 3: Marginals of state 1, identified with $(0, 0, 1)$. Modes 0 and 1 remain in their ground states, while mode 2 carries the first vibrational excitation and therefore shows the characteristic nodal separation in its probability density.

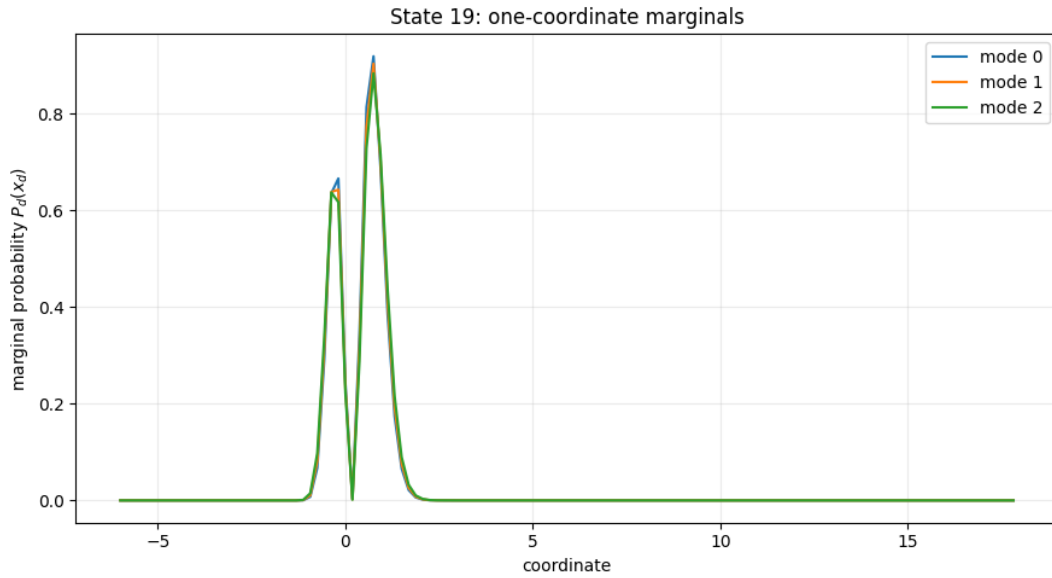


Figure 4: Marginals of state 19, identified with $(1, 1, 1)$. All three coordinates occupy first-excited Morse states, producing closely overlapping two-lobed marginal probability distributions.

14 Convergence and compression diagnostics

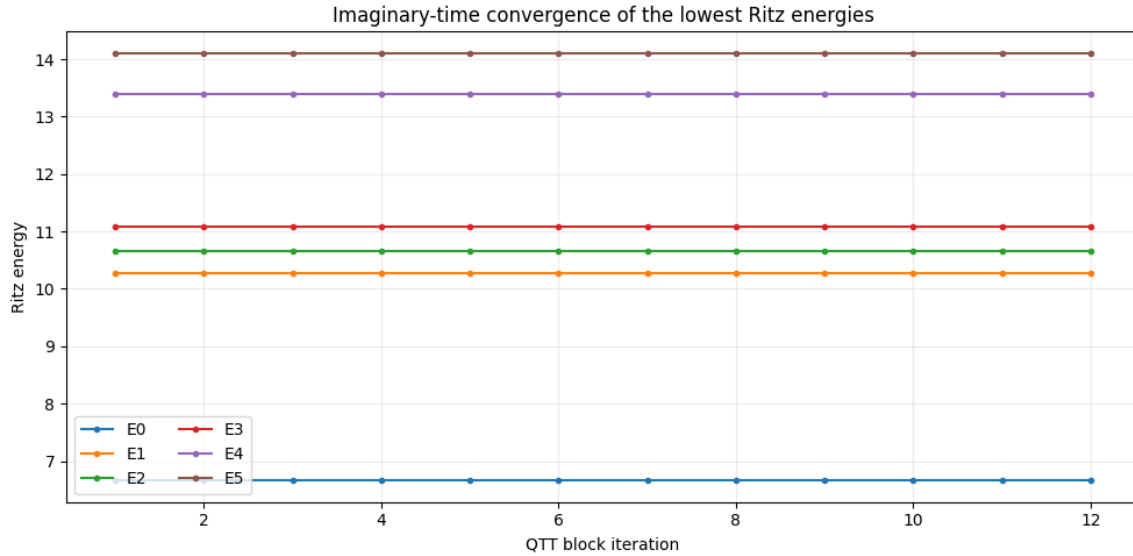


Figure 5: Imaginary-time evolution of the six lowest Ritz energies. The curves are nearly flat because the local-FGH product seeds are already extremely accurate eigenstates for the uncoupled Hamiltonian. In this run the propagation therefore serves mainly as a consistency test of the QTT split-operator and block-rotation machinery.

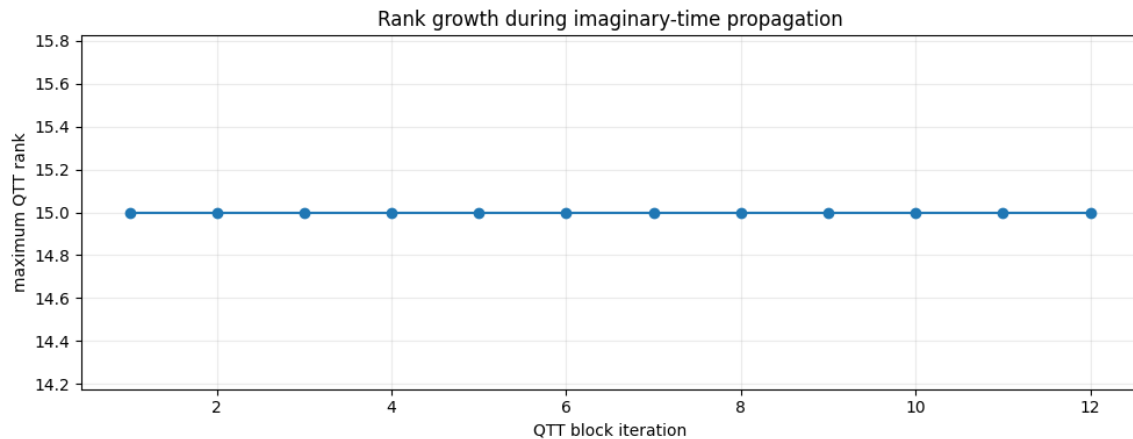


Figure 6: Maximum TT rank of the 20-state block during the 12 imaginary-time iterations. It remains fixed at 15, indicating stable compression for this uncoupled benchmark.

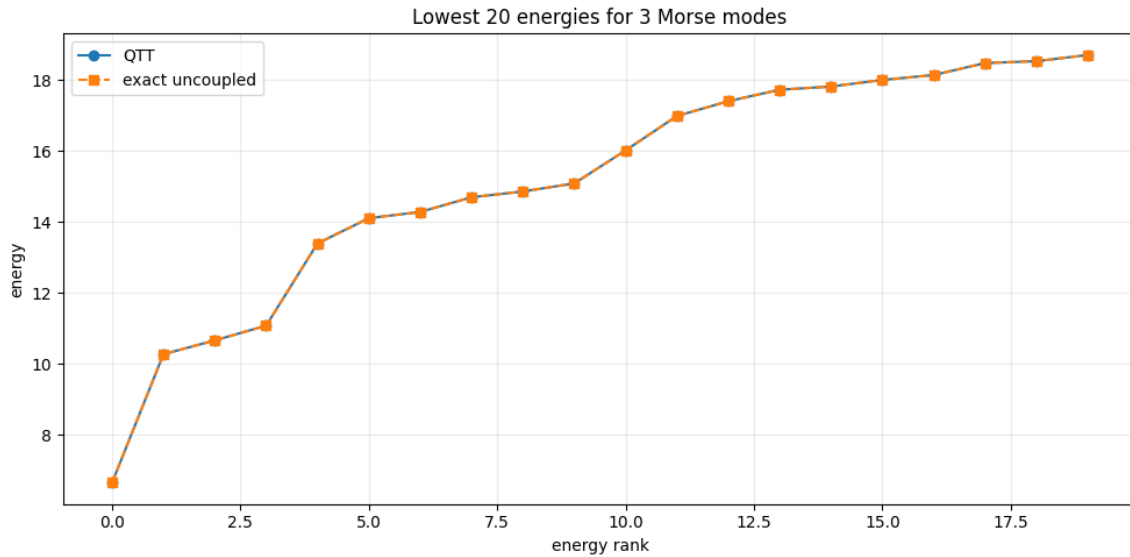


Figure 7: Lowest 20 QTT energies compared with the analytical uncoupled Morse product spectrum. The two curves overlap at plotting resolution, consistent with the maximum absolute error of 5.336×10^{-9} .

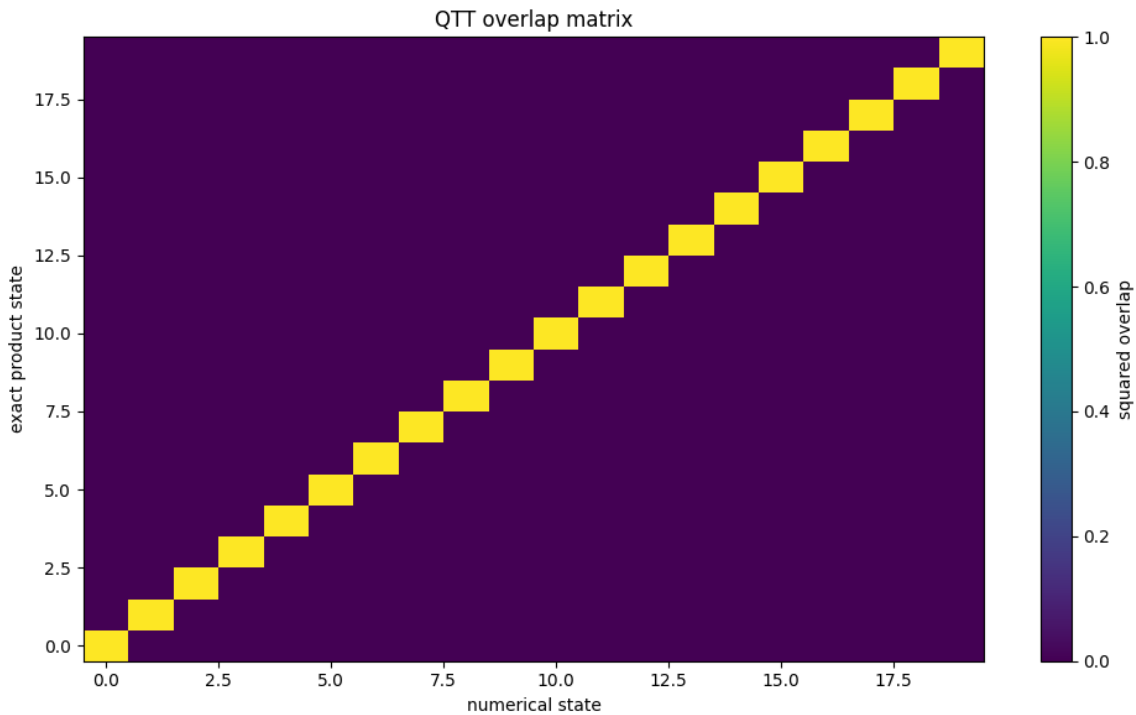


Figure 8: Squared-overlap matrix between exact product states and QTT eigenstates. The nearly perfect diagonal confirms one-to-one state recovery across the 20-state block.

15 What the results demonstrate

The benchmark isolates the QTT machinery in a setting where an independent analytical solution is available. Several conclusions follow directly from the saved outputs.

1. **The physical-core Fourier transform is numerically reversible.** A round-trip error of

- 7.721×10^{-15} shows that the unfold-FFT-refold construction is not introducing a measurable error at this scale beyond floating-point roundoff.
2. **The structured Hamiltonian is highly compressible.** The full three-dimensional potential and kinetic diagonals have maximum QTT ranks only 4 and 3, respectively.
 3. **The one-dimensional grid is already converged for the low levels.** The FGH/analytical errors in Table 4 are 10^{-14} – 10^{-12} .
 4. **Block imaginary-time propagation preserves the target subspace.** Because the uncoupled FGH seeds are already excellent, the Ritz energies change only in the ninth or later decimal places while the rank remains 15.
 5. **The final state identities are correct, not merely the energies.** The overlap matrix is essentially the identity and the minimum matched fidelity is 0.999999999755.
 6. **Energy agreement alone is not the whole convergence story.** The largest energy error is only 5.336×10^{-9} , while the largest residual norm is 3.798×10^{-4} . This illustrates why the notebook reports both scalar energy accuracy and a vector residual.

16 How the code generalizes to coupled systems

The notebook has three potential branches.

none	The potential and its imaginary-time propagator both factorize across coordinates. This is the clean analytical benchmark.
morse_pair	The potential itself is still assembled exactly from one-dimensional QTT factors using $\kappa_{ij}s_i s_j$, but $e^{-\Delta\tau V/2}$ does not factorize. The scalar exponential is therefore applied to VQTT using TT-cross.
custom	The many-body potential itself is constructed by TT-cross from selected binary-grid samples, and the potential propagator is likewise generated through TT-cross.

For coupled systems, analytical product states are no longer exact eigenvectors. The code consequently changes its interpretation step: it projects each numerical state onto a low-energy uncoupled product basis, reports the dominant product configuration, the captured weight, and an effective participation measure. Rank growth then becomes a physical/numerical indicator of increased intermode correlation.

17 The dormant two-dimensional heat-map branch

If the configuration is changed to

```
1 N_OSC = 2
2 COUPLING_MODEL = "none"
```

the final plotting cell reconstructs each selected two-dimensional QTT state and its analytical product reference. Because eigenvectors are defined only up to a global phase, it first computes

$$\phi = \arg\langle\Psi^{\text{exact}}|\Psi^{\text{QTT}}\rangle, \quad \Psi_{\text{aligned}}^{\text{QTT}} = e^{-i\phi}\Psi^{\text{QTT}}, \quad (42)$$

and then plots a 2×2 panel containing the real parts and probability densities of the numerical and exact states on common color scales.

No such figures are included here because the supplied executed notebook uses three modes and therefore did not generate them. This is why the one-coordinate marginals are the appropriate state visualization for the saved run.

18 Scaling and practical limitations

With L_d binary digits per coordinate, a dense state contains

$$N_{\text{dense}} = 2^{\sum_d L_d} \quad (43)$$

amplitudes. A binary QTT with characteristic rank r stores approximately

$$O\left(2r^2 \sum_d L_d\right) \quad (44)$$

core coefficients, although actual ranks vary along the train.

The compression is useful only while the wavefunctions, operators, and Fourier-transformed states remain low rank. Strong coupling, delocalized or near-dissociation states, and complicated many-body correlations can drive the ranks upward. Accordingly, a meaningful convergence study must monitor *both* physical observables and tensor ranks.

The notebook recommends the following progression:

1. Verify the uncoupled spectrum and state fidelities against the analytical Morse solution.
2. Increase L_d until the states of interest are insensitive to the grid.
3. Tighten QTT_EPS and QFT_EPS, and increase QTT_RMAX, until energies, residuals, and ranks stabilize.
4. Introduce weak coupling and inspect both energy shifts and rank growth.
5. Increase the coupling gradually and track product-basis weights and one-coordinate marginals.
6. Use dense validation only for small cases; it is not part of the scalable algorithm.

19 Code architecture: how the notebook pieces fit together

Table 7: Functional map of the notebook.

Cell(s)	Main objects/functions	Purpose
2	Manual TT-SVD example	Demonstrates binary tensorization, SVD, TT cores, and exact reconstruction.
7	qtt_from_dense, qtt_round, qtt_add, qtt_hadamard, qtt_inner	Implements the core low-rank tensor algebra, including complex-valued support and explicit rank truncation.
9	qtt_qft and helper routines	Unfolds each binary coordinate block, performs column-wise FFTs, refolds by SVD, and rounds.
11, 24	Block overlap and Rayleigh–Ritz helpers	Performs QTT-native Löwdin orthogonalization and diagonalization in the small state subspace.
13	Grid and physical parameters	Defines the number of oscillators, Morse parameters, binary grids, momenta, and quadrature volume.
16	VQTT	Builds the potential by structured QTT algebra for built-in Morse models or TT-cross for a custom PES.

Cell(s)	Main objects/functions	Purpose
18	<code>Ttot_qtt</code>	Builds the separable momentum-space kinetic diagonal.
20	QFT validation	Tests QFT followed by inverse QFT without dense reconstruction.
22	<code>apply_H_qtt</code> , <code>qtt_energy</code>	Applies the full Hamiltonian and evaluates Rayleigh quotients.
26–30	Analytical levels, exact QTTs, local FGH seeds	Creates independent reference data and high-quality numerical initial states.
32–34	Propagators and block imaginary time	Executes Strang splitting and Rayleigh–Ritz after every propagation step.
36	Dense validation	Optional small-grid cross-check; skipped in the saved 128^3 run.
38	Final state analysis	Computes overlaps, Hungarian matching, residuals, ranks, and coefficient counts.
40	Marginals	Contracts $ \Psi ^2$ to obtain one-coordinate distributions without dense reconstruction.
42	Convergence diagnostics	Plots energies, maximum rank, final spectrum, and overlap matrix.
44	2D heat maps	Runs only for uncoupled two-mode calculations; skipped in the saved three-mode run.

20 Recommended modifications for further study

1. **Grid convergence:** compare $L = 5, 6, 7$ and track the maximum energy error, minimum fidelity, and residual norm.
2. **Compression convergence:** repeat with `QTT_EPS` = $10^{-8}, 10^{-10}, 10^{-12}$ and monitor the accuracy/rank tradeoff.
3. **Rank-cap convergence:** vary `QTT_RMAX` through 32, 64, and 96.
4. **Propagation convergence:** increase the number of imaginary-time steps and reduce the smallest $\Delta\tau$ to test whether residuals decrease further.
5. **Coupling scan:** activate `morse_pair`, vary the coupling strength, and track energy shifts, product-state mixing, and rank growth.
6. **Dimensional scaling:** repeat for $n = 2, 3, 4$ at fixed L and compare nominal dense size against actual QTT coefficient counts.
7. **Custom PES:** replace `CUSTOM_COUPLING(X)` with a genuine three-mode interaction and inspect TT-cross ranks and validation errors.

21 Summary

The notebook implements a fully QTT-native eigensolver architecture in which the multidimensional wavefunctions remain tensor trains throughout the scalable path. Binary tensorization compresses each coordinate, a physical-core FFT moves efficiently between coordinate and momentum representations, structured QTT algebra builds the Morse Hamiltonian, and block imaginary-time propagation combined

with Löwdin orthogonalization and Rayleigh–Ritz extracts multiple low-energy states. The analytical Morse solution then supplies a stringent independent benchmark.

In the executed three-mode, 20-state run, the QFT is reversible to 7.721×10^{-15} relative error, the structured potential and kinetic terms remain at maximum ranks 4 and 3, propagation maintains a maximum state rank of 15, and the final 20-state spectrum agrees with the analytical product spectrum to a maximum absolute energy error of 5.336×10^{-9} with minimum fidelity 0.99999999755. The results therefore demonstrate both the numerical correctness of the implemented QTT machinery and the importance of monitoring residuals and rank growth in addition to energies.

A Complete code extracted from the notebook

The following listing is an exact concatenation of all code cells in the supplied notebook, with cell-number comments added only as separators. It is also provided as `code/notebook_code_all.py` in the accompanying bundle.

```

1
2 # ===== Notebook cell 2 =====
3
4 # A hands-on, from-scratch TT-SVD of a length-4 vector.
5 # This uses plain numpy so every step is visible, before we hand things
6 # off to tntorch's Tensor objects for the real calculation below.
7
8 import numpy as np
9
10 v = np.array([1.0, 2.0, 3.0, 4.0]) # the vector we want to compress
11 L = 2                               # number of bits, N = 2*L = 4
12
13 # Reshape (1, N) -> (2, 2): first axis = bit 1 (MSB), second axis = bit 2 (LSB)
14 M = v.reshape(2, 2)
15 print("Reshaped as a matrix M[b1, b2]:\n", M)
16
17 # Standard SVD. The rank of M (number of nonzero singular values) is the
18 # TT bond dimension r1 needed between core 1 and core 2.
19 U, S, Vt = np.linalg.svd(M, full_matrices=False)
20 print("\nSingular values:", S)
21
22 r1 = int(np.sum(S > 1e-12))
23 print("Bond dimension r1 =", r1)
24
25 # Core 1: shape (r0=1, n1=2, r1) -- here U already has shape (2, r1)
26 core1 = U[:, :r1].reshape(1, 2, r1)
27
28 # Core 2: shape (r1, n2=2, r2=1) -- fold the singular values into core 2
29 core2 = (np.diag(S[:r1]) @ Vt[:r1, :]).reshape(r1, 2, 1)
30
31 print("\ncore1 shape:", core1.shape)
32 print("core2 shape:", core2.shape)
33
34 # Reconstruct: contract core1 with core2 over the shared bond index.
35 rebuilt = np.tensordot(core1, core2, axes=([2], [0])) # -> shape (1, 2, 2, 1)
36 rebuilt = rebuilt.reshape(4)
37 print("\nReconstructed vector:", rebuilt)
38 print("Matches original:", np.allclose(rebuilt, v))
39
40
41
42
43 # ===== Notebook cell 4 =====

```

```

44
45 # Install tntorch, the tensor-train library used throughout this tutorial.
46 # In Google Colab this cell is normally all that is needed.
47 !pip -q install tntorch
48
49
50
51 # ===== Notebook cell 5 =====
52
53 # Imports
54
55 import math
56 import heapq
57 from pathlib import Path
58
59 import numpy as np
60 import pandas as pd
61 import torch
62 import tntorch as tn
63 import matplotlib.pyplot as plt
64
65 from scipy.special import eval_genlaguerre
66 from scipy.optimize import linear_sum_assignment
67 from IPython.display import display
68
69 torch.set_default_dtype(torch.float64)
70 device = torch.device("cpu")
71
72 OUTPUT_DIR = Path("/content") if Path("/content").exists() else Path(".")
73
74 print("PyTorch :", torch.__version__)
75 print("Device  :", device)
76 print("Output  :", OUTPUT_DIR.resolve())
77 print("tntorch : imported successfully")
78
79
80
81 # ===== Notebook cell 7 =====
82
83 # Cell 3 -- Generic TT/QTT routines
84 #
85 # Everything below operates on tntorch's tn.Tensor, which is just a
86 # Python list of cores (torch.Tensor objects), each of shape
87 # (r_left, 2, r_right). We reimplement TT-SVD, rounding, addition,
88 # Hadamard product, and inner product by hand (rather than calling
89 # tntorch's built-ins) so that complex dtypes -- needed once we start
90 # Fourier-transforming wavefunctions -- are handled correctly throughout.
91
92 QTT_EPS = 1.0e-12 # rounding tolerance for algebra results (add, hadamard, ...)
93 QFT_EPS = 2.0e-13 # rounding tolerance specifically after Fourier transforms
94 QTT_RMAX = 96     # hard ceiling on any TT bond dimension
95
96
97 def _rank_from_singular_values(S, delta=None, rmax=None):
98     """
99     Given a vector of singular values S (sorted descending, as returned by
100     torch.linalg.svd), decide how many to keep.
101
102     The rule mirrors the standard TT-SVD/TT-rounding truncation criterion:
103     keep the smallest r such that the "discarded tail" satisfies
104     || S[r:] ||_2 <= delta

```

```

105     i.e. we keep just enough singular values that the Frobenius-norm error
106     introduced by throwing away the rest is below the target tolerance.
107     'rmax' additionally hard-caps r regardless of delta.
108     """
109     m = int(S.numel())
110
111     if m == 0:
112         return 1
113
114     r = m
115
116     if delta is not None:
117         delta = float(delta)
118
119         for candidate in range(1, m + 1):
120
121             if candidate == m:
122                 tail = 0.0
123             else:
124                 # Norm of the singular values we would be discarding if we
125                 # kept only the first 'candidate' of them.
126                 tail = float(
127                     torch.linalg.vector_norm(
128                         S[candidate:]
129                     ).detach().cpu()
130                 )
131
132                 if tail <= delta:
133                     r = candidate
134                     break
135
136     if rmax is not None:
137         r = min(
138             r,
139             int(rmax)
140         )
141
142     return max(1, r)
143
144
145 def qtt_from_dense(
146     tensor,
147     bits_per_dim,
148     eps=QTT_EPS,
149     rmax=QTT_RMAX
150 ):
151     """
152     TT-SVD of a dense physical tensor into a binary QTT.
153
154     bits_per_dim = [L1, L2, ...]
155     means physical dimension d has size 2**Ld.
156
157     This is exactly the "reshape into a matrix, SVD, keep U, fold the
158     remainder into the next core, repeat" procedure you saw by hand in the
159     Primer's 4-element example -- just generalized to 'total_bits' bits and
160     wrapped with the delta-based truncation from '_rank_from_singular_values'.
161     The overall relative-error budget 'eps' is spread evenly across the
162     total_bits-1 SVD cuts via delta = eps * ||tensor|| / sqrt(total_bits-1),
163     which is the standard TT-SVD error-balancing heuristic (Oseledets 2011):
164     it guarantees the final relative Frobenius-norm error is <= eps.
165     """

```

```

166     bits_per_dim = list(
167         map(int, bits_per_dim)
168     )
169
170     total_bits = sum(
171         bits_per_dim
172     )
173
174     physical_shape = [
175         2**b
176         for b in bits_per_dim
177     ]
178
179     tensor = torch.as_tensor(
180         tensor,
181         device=device
182     )
183
184     if int(tensor.numel()) != int(
185         np.prod(physical_shape)
186     ):
187         raise ValueError(
188             "Dense tensor size does not match bits_per_dim."
189         )
190
191     # Reinterpret the dense array as a total_bits-index binary array: this
192     # is the quantics reshape described in the Primer (MSB-first per
193     # coordinate, with each coordinate's bits stored contiguously in order.
194     work = tensor.reshape(
195         [2] * total_bits
196     ).contiguous()
197
198     if total_bits == 1:
199         return tn.Tensor([
200             work.reshape(
201                 1,
202                 2,
203                 1
204             ).clone()
205         ])
206
207     norm0 = float(
208         torch.linalg.vector_norm(
209             work.reshape(-1)
210         ).detach().cpu()
211     )
212
213     delta = (
214         eps * norm0
215         / math.sqrt(total_bits - 1)
216         if eps is not None
217         else None
218     )
219
220     qcores = []
221     rleft = 1
222
223     # Sweep left to right: at each bit, fold the current bit into the
224     # left "row" index, SVD, keep U as the core, and push (S @ Vh) into
225     # the still-unprocessed remainder to be reshaped/truncated next.
226     for k in range(

```

```

227     total_bits - 1
228 ):
229
230     M = work.reshape(
231         rleft * 2,
232         -1
233     )
234
235     U, S, Vh = torch.linalg.svd(
236         M,
237         full_matrices=False
238     )
239
240     r = _rank_from_singular_values(
241         S,
242         delta=delta,
243         rmax=rmax
244     )
245
246     U = U[:, :r]
247     S = S[:r]
248     Vh = Vh[:, :r]
249
250     qcores.append(
251         U.reshape(
252             rleft,
253             2,
254             r
255         ).contiguous()
256     )
257
258     remaining_bits = (
259         total_bits
260         - k
261         - 1
262     )
263
264     work = (
265         S[:, None]
266         * Vh
267     ).reshape(
268         r,
269         *([2] * remaining_bits)
270     ).contiguous()
271
272     rleft = r
273
274     # Last core absorbs whatever remains; its right bond dimension is 1
275     # by construction (a TT always closes with boundary rank 1).
276     qcores.append(
277         work.reshape(
278             rleft,
279             2,
280             1
281         ).contiguous()
282     )
283
284     return tn.Tensor(
285         qcores
286     )
287

```

```

288
289 def qtt_to_dense(
290     qtt,
291     bits_per_dim
292 ):
293     """
294     The inverse of qtt_from_dense(): contract every core left-to-right
295     (summing over shared bond indices) to rebuild the dense physical
296     tensor. Complex-valued cores (which appear once we Fourier-transform
297     a wavefunction) are preserved -- no .real is taken here.
298     """
299     bits_per_dim = list(
300         map(int, bits_per_dim)
301     )
302
303     cores = qtt.cores
304
305     if len(cores) != sum(
306         bits_per_dim
307     ):
308         raise ValueError(
309             "Number of QTT cores does not match bits_per_dim."
310         )
311
312     for core in cores:
313         if (
314             core.dim() != 3
315             or core.shape[1] != 2
316         ):
317             raise ValueError(
318                 "Expected pure binary TT cores."
319             )
320
321     # Start from the (trivial) left boundary rank-1 index of the first core.
322     acc = cores[0][
323         0,
324         :,
325         :
326     ]
327
328     # Contract in the remaining cores one at a time over the shared bond.
329     for core in cores[1:]:
330
331         acc = torch.tensordot(
332             acc,
333             core,
334             dims=(
335                 [-1],
336                 [0]
337             )
338         )
339
340     if acc.shape[-1] != 1:
341         raise ValueError(
342             "Final TT rank is not one."
343         )
344
345     binary_tensor = (
346         acc[..., 0]
347     )
348

```

```

349     physical_shape = [
350         2**b
351         for b in bits_per_dim
352     ]
353
354     return binary_tensor.reshape(
355         physical_shape
356     )
357
358
359 def qtt_ranks(qtt):
360     """Return the full list of bond dimensions [r0=1, r1, r2, ..., rD=1]."""
361     return [
362         int(
363             qtt.cores[0].shape[0]
364         )
365     ] + [
366         int(core.shape[-1])
367         for core in qtt.cores
368     ]
369
370
371 def qtt_numcoef(qtt):
372     """Total number of stored floating-point numbers across all cores --
373     the actual memory footprint of the QTT, to compare against the full dense grid."""
374     return sum(
375         int(core.numel())
376         for core in qtt.cores
377     )
378
379
380 def qtt_summary(
381     name,
382     qtt
383 ):
384     """Pretty-print a one-line rank/size summary; used throughout the
385     notebook to keep an eye on how compressed each QTT is."""
386     ranks = qtt_ranks(
387         qtt
388     )
389
390     print(
391         f"{name:18s} "
392         f"max rank={max(ranks):3d}    "
393         f"coefficients={qtt_numcoef(qtt):8d}    "
394         f"ranks={ranks}"
395     )
396
397
398 def qtt_round(
399     qtt,
400     eps=QTT_EPS,
401     rmax=QTT_RMAX
402 ):
403     """
404     Complex-safe TT rounding using QR + SVD.
405
406     Rounding re-expresses the SAME tensor at (at most) the given error
407     tolerance, using the smallest possible ranks. It is a two-pass sweep:
408
409     1. Left-orthogonalization sweep (QR, left to right): reshapes each

```

```

410         core into "tall" matrices and QR-factorizes them so that every
411         core except the last becomes left-orthogonal (its columns are
412         orthonormal). The leftover R factor is absorbed into the next
413         core. This does not change the tensor or its rank -- it just
414         puts it into a canonical form that makes the next step's error
415         estimate meaningful.
416
417         2. Right-to-left SVD truncation sweep: now sweep from the right,
418         SVD each core, discard the small singular values (per the same
419         delta budget used in TT-SVD), and fold the U*S factor into the
420         core to its left. This step is where the actual compression
421         happens.
422
423     This exact two-sweep structure is the standard TT-rounding algorithm
424     (Oseledets 2011); reimplementing it here (rather than calling
425     tntorch's rounding) lets us keep everything correct for complex dtype.
426     """
427     cores = [
428         core.clone()
429         for core in qtt.cores
430     ]
431
432     d = len(
433         cores
434     )
435
436     if d <= 1:
437         return tn.Tensor(
438             cores
439         )
440
441     # --- Pass 1: left orthogonalization -----
442     for k in range(
443         d - 1
444     ):
445
446         r0, n, r1 = (
447             cores[k].shape
448         )
449
450         M = cores[k].reshape(
451             r0 * n,
452             r1
453         )
454
455         Q, R = torch.linalg.qr(
456             M,
457             mode="reduced"
458         )
459
460         rnew = Q.shape[1]
461
462         cores[k] = Q.reshape(
463             r0,
464             n,
465             rnew
466         ).contiguous()
467
468         cores[k + 1] = torch.tensordot(
469             R,
470             cores[k + 1],

```

```

471         dims=(
472             [1],
473             [0]
474         )
475     ).contiguous()
476
477     tensor_norm = float(
478         torch.linalg.vector_norm(
479             cores[-1].reshape(-1)
480         ).detach().cpu()
481     )
482
483     delta = (
484         eps * tensor_norm
485         / math.sqrt(d - 1)
486         if eps is not None
487         else None
488     )
489
490     # --- Pass 2: right-to-left SVD truncation -----
491     for k in range(
492         d - 1,
493         0,
494         -1
495     ):
496
497         r0, n, r1 = (
498             cores[k].shape
499         )
500
501         M = cores[k].reshape(
502             r0,
503             n * r1
504         )
505
506         U, S, Vh = torch.linalg.svd(
507             M,
508             full_matrices=False
509         )
510
511         r = _rank_from_singular_values(
512             S,
513             delta=delta,
514             rmax=rmax
515         )
516
517         U = U[:, :r]
518         S = S[:r]
519         Vh = Vh[:, :r]
520
521         cores[k] = Vh.reshape(
522             r,
523             n,
524             r1
525         ).contiguous()
526
527         US = (
528             U
529             * S[None, :]
530         )
531

```

```

532         cores[k - 1] = torch.tensordot(
533             cores[k - 1],
534             US,
535             dims=(
536                 [2],
537                 [0]
538             )
539         ).contiguous()
540
541     return tn.Tensor(
542         cores
543     )
544
545 def qtt_hadamard(
546     a,
547     b,
548     eps=QTT_EPS,
549     rmax=QTT_RMAX
550 ):
551     """
552     Pointwise (Hadamard/elementwise) product of two QTTs,  $c(i) = a(i)*b(i)$ .
553
554     Core-by-core, the elementwise product of two TTs is built by taking the
555     "outer product" of each pair of cores along both bond directions:
556     if A has shape (rA_left, 2, rA_right) and B has shape (rB_left, 2, rB_right),
557     the combined core C has shape (rA_left*rB_left, 2, rA_right*rB_right),
558     with  $C[(a,b), i, (a',b')] = A[a,i,a'] * B[b,i,b']$ . Physically this means
559     the Hadamard-product rank is (at most) the PRODUCT of the input ranks --
560     which is why every call here is immediately followed by qtt_round() to
561     bring the rank back down to what the actual accuracy tolerance needs.
562     """
563
564     if len(a.cores) != len(
565         b.cores
566     ):
567         raise ValueError(
568             "QTT dimensions do not match."
569         )
570
571     out = []
572
573     for A0, B0 in zip(
574         a.cores,
575         b.cores
576     ):
577
578         if A0.shape[1] != B0.shape[1]:
579             raise ValueError(
580                 "Physical mode sizes do not match."
581             )
582
583         dtype = torch.promote_types(
584             A0.dtype,
585             B0.dtype
586         )
587
588         A = A0.to(
589             dtype
590         )
591
592         B = B0.to(

```

```

593         dtype
594     )
595
596     # Outer product along both bond legs, elementwise along the
597     # shared physical (bit) index.
598     C = (
599         A[
600             :,
601             None,
602             :,
603             :,
604             None
605         ]
606         * B[
607             None,
608             :,
609             :,
610             None,
611             :
612         ]
613     )
614
615     C = C.reshape(
616         A.shape[0]
617         * B.shape[0],
618         A.shape[1],
619         A.shape[2]
620         * B.shape[2]
621     )
622
623     out.append(
624         C.contiguous()
625     )
626
627     return qtt_round(
628         tn.Tensor(out),
629         eps=eps,
630         rmax=rmax
631     )
632
633
634 def qtt_add(
635     a,
636     b,
637     eps=QTT_EPS,
638     rmax=QTT_RMAX
639 ):
640     """
641     Sum of two QTTs,  $c(i) = a(i) + b(i)$ , by TT block concatenation.
642
643     The standard exact-TT-addition trick: stack a's and b's cores
644     block-diagonally along the bond dimensions (except at the very first
645     and very last core, where you concatenate along the single free bond,
646     since the boundary ranks must stay 1). Like Hadamard product, exact
647     addition ADDS the ranks together (rather than multiplying), so this
648     is also always followed by qtt_round().
649     """
650     if len(a.cores) != len(
651         b.cores
652     ):
653         raise ValueError(

```

```

654         "QTT dimensions do not match."
655     )
656
657     d = len(
658         a.cores
659     )
660
661     if d == 1:
662
663         dtype = torch.promote_types(
664             a.cores[0].dtype,
665             b.cores[0].dtype
666         )
667
668         return tn.Tensor([
669             a.cores[0].to(dtype)
670             + b.cores[0].to(dtype)
671         ])
672
673     cores = []
674
675     for k, (A0, B0) in enumerate(
676         zip(
677             a.cores,
678             b.cores
679         )
680     ):
681
682         dtype = torch.promote_types(
683             A0.dtype,
684             B0.dtype
685         )
686
687         A = A0.to(
688             dtype
689         )
690
691         B = B0.to(
692             dtype
693         )
694
695         if k == 0:
696             # First core: concatenate along the right bond only
697             # (left bond must stay 1).
698             C = torch.cat(
699                 [A, B],
700                 dim=2
701             )
702
703         elif k == d - 1:
704             # Last core: concatenate along the left bond only
705             # (right bond must stay 1).
706             C = torch.cat(
707                 [A, B],
708                 dim=0
709             )
710
711         else:
712             # Interior cores: block-diagonal placement.
713             C = torch.zeros(
714                 A.shape[0]

```

```

715         + B.shape[0],
716         A.shape[1],
717         A.shape[2]
718         + B.shape[2],
719         dtype=dtype,
720         device=A.device
721     )
722
723     C[
724         :A.shape[0],
725         :,
726         :A.shape[2]
727     ] = A
728
729     C[
730         A.shape[0]:,
731         :,
732         A.shape[2]:
733     ] = B
734
735     cores.append(
736         C.contiguous()
737     )
738
739     return qtt_round(
740         tn.Tensor(cores),
741         eps=eps,
742         rmax=rmax
743     )
744
745
746 def qtt_inner(
747     a,
748     b
749 ):
750     """
751     Algebraic TT inner product  $\langle a|b \rangle = \sum_i \text{conj}(a_i) * b_i$ .
752
753     Contract the two TTs against each other core by core, carrying a
754     small (rA x rB) "environment" matrix between sites -- this is the
755     standard TT/MPS inner-product contraction, with cost linear in the
756     number of cores rather than exponential in the number of physical
757     indices, because we never form the dense tensors.
758     """
759     if len(a.cores) != len(
760         b.cores
761     ):
762         raise ValueError(
763             "QTT dimensions do not match."
764         )
765
766     dtype = torch.promote_types(
767         a.cores[0].dtype,
768         b.cores[0].dtype
769     )
770
771     env = torch.ones(
772         (1, 1),
773         dtype=dtype,
774         device=a.cores[0].device
775     )

```

```

776
777     for A0, B0 in zip(
778         a.cores,
779         b.cores
780     ):
781
782         A = A0.to(
783             dtype
784         )
785
786         B = B0.to(
787             dtype
788         )
789
790         # env[a,b], A*[a,i,c], B[b,i,d] -> new env[c,d]
791         env = torch.einsum(
792             "ab,aic,bid->cd",
793             env,
794             A.conj(),
795             B
796         )
797
798     return env.squeeze()
799
800
801 def qtt_tensor_product(
802     a,
803     b
804 ):
805     """
806     Tensor product of two QTTs describing two INDEPENDENT coordinate
807     groups, e.g. combining QTTs defined on two independent coordinate groups
808     into one QTT over the concatenated binary modes.
809
810     Because a TT's right boundary rank and the next TT's left boundary
811     rank are both 1, simply concatenating the two core lists produces
812     exactly the product f(group1)*g(group2) in QTT form -- no rounding is needed.
813     """
814     if (
815         a.cores[-1].shape[-1] != 1
816         or b.cores[0].shape[0] != 1
817     ):
818         raise ValueError(
819             "Expected boundary TT ranks equal to one."
820         )
821
822     return tn.Tensor(
823         [
824             core.clone()
825             for core in (
826                 list(a.cores)
827                 + list(b.cores)
828             )
829         ]
830     )
831
832
833 # ===== Notebook cell 9 =====
834
835 # Cell 4 -- Unfold -> column FFT -> fold QFT

```

```

837
838 def qtt_unfold_to_physical_tt(
839     qtt,
840     bits_per_dim
841 ):
842     """
843     Step 1 of the QFT recipe: contract each coordinate's group of binary
844     QTT cores into one "physical" TT core of shape (R_left, 2**Ld, R_right).
845
846     This is literally the same recursive tensordot-and-reshape merge you'd
847     use to undo one level of the quantics bit-splitting -- we're re-forming
848     the dense axis for coordinate d, but leaving the OTHER coordinates'
849     cores untouched, so the intermediate object is still compact whenever
850     the cross-coordinate TT rank stays small.
851     """
852     bits_per_dim = list(
853         map(int, bits_per_dim)
854     )
855
856     if sum(bits_per_dim) != len(
857         qtt.cores
858     ):
859         raise ValueError(
860             "sum(bits_per_dim) must equal the number of QTT cores."
861         )
862
863     physical_cores = []
864     offset = 0
865
866     for nbits in bits_per_dim:
867
868         group = qtt.cores[
869             offset:
870             offset + nbits
871         ]
872
873         merged = group[0].clone()
874
875         for next_core in group[1:]:
876
877             merged = torch.tensordot(
878                 merged,
879                 next_core,
880                 dims=(
881                     [2],
882                     [0]
883                 )
884             )
885
886             merged = merged.reshape(
887                 merged.shape[0],
888                 -1,
889                 merged.shape[-1]
890             ).contiguous()
891
892         if merged.shape[1] != 2**nbits:
893             raise RuntimeError(
894                 "Unexpected physical size while unfolding QTT."
895             )
896
897         physical_cores.append(

```

```

898         merged
899     )
900
901     offset += nbits
902
903     return physical_cores
904
905
906 def fft_physical_tt_cores(
907     physical_cores,
908     inverse=False,
909     norm="ortho"
910 ):
911     """
912     Step 2 of the QFT recipe: FFT every "column" of every physical TT core
913     along its physical (length-N) mode.
914
915     A core has shape (rleft, N, rright). We want to Fourier-transform only
916     the N axis, independently for each of the rleft*rright choices of bond
917     indices -- so we permute N to the front, flatten the bond indices into
918     one batch axis, call torch.fft.fft along dim=0, then unflatten. Using
919     norm="ortho" keeps the transform unitary, matching torch.fft.fftn's
920     "ortho" convention used for validation in Cell 7 and in apply_H_dense().
921     """
922     transformed = []
923
924     fft_fun = (
925         torch.fft.ifft
926         if inverse
927         else torch.fft.fft
928     )
929
930     for core in physical_cores:
931
932         rleft, Ndim, rright = (
933             core.shape
934         )
935
936         complex_dtype = (
937             torch.complex128
938             if core.dtype in (
939                 torch.float64,
940                 torch.complex128
941             )
942             else torch.complex64
943         )
944
945         core_c = core.to(
946             complex_dtype
947         )
948
949         # (rL, N, rR) -> (N, rL*rR): every column is independently
950         # transformed.
951         columns = (
952             core_c
953             .permute(
954                 1,
955                 0,
956                 2
957             )
958             .reshape(

```

```

959         Ndim,
960         rleft * rright
961     )
962 )
963
964     columns_k = fft_fun(
965         columns,
966         dim=0,
967         norm=norm
968     )
969
970     core_k = (
971         columns_k
972         .reshape(
973             Ndim,
974             rleft,
975             rright
976         )
977         .permute(
978             1,
979             0,
980             2
981         )
982         .contiguous()
983     )
984
985     transformed.append(
986         core_k
987     )
988
989     return transformed
990
991
992 def _split_physical_core_exact(
993     core,
994     nbits
995 ):
996     """
997     Step 3a: split ONE physical TT core of shape (R_left, 2**nbits, R_right)
998     back into 'nbits' binary cores, via the same recursive SVD used in
999     qtt_from_dense(), but WITHOUT truncation here (full SVD rank is kept --
1000     the eventual compression happens afterward, once in
1001     physical_tt_to_qtt(), by a single qtt_round() call over the whole QTT).
1002     Keeping this step exact avoids compounding rounding error coordinate
1003     by coordinate.
1004     """
1005     rleft_external, Ndim, rright_external = (
1006         core.shape
1007     )
1008
1009     if Ndim != 2**nbits:
1010         raise ValueError(
1011             "Physical dimension is incompatible with nbits."
1012         )
1013
1014     if nbits == 1:
1015         return [
1016             core.reshape(
1017                 rleft_external,
1018                 2,
1019                 rright_external

```

```

1020         ).contiguous()
1021     ]
1022
1023     work = core.reshape(
1024         rleft_external,
1025         *([2] * nbits),
1026         rright_external
1027     ).contiguous()
1028
1029     qcores = []
1030     rleft = (
1031         rleft_external
1032     )
1033
1034     for k in range(
1035         nbits - 1
1036     ):
1037
1038         M = work.reshape(
1039             rleft * 2,
1040             -1
1041         )
1042
1043         U, S, Vh = torch.linalg.svd(
1044             M,
1045             full_matrices=False
1046         )
1047
1048         # No truncation by delta/rmax here -- keep the full-rank SVD so
1049         # this split step is numerically exact.
1050         r = int(
1051             S.numel()
1052         )
1053
1054         qcores.append(
1055             U.reshape(
1056                 rleft,
1057                 2,
1058                 r
1059             ).contiguous()
1060         )
1061
1062         remaining_bits = (
1063             nbits
1064             - k
1065             - 1
1066         )
1067
1068         work = (
1069             S[:, None]
1070             * Vh
1071         ).reshape(
1072             r,
1073             *([2] * remaining_bits),
1074             rright_external
1075         ).contiguous()
1076
1077         rleft = r
1078
1079     qcores.append(
1080         work.reshape(

```

```

1081         rleft,
1082         2,
1083         rright_external
1084     ).contiguous()
1085 )
1086
1087 return qcores
1088
1089
1090 def physical_tt_to_qtt(
1091     physical_cores,
1092     bits_per_dim,
1093     eps=QFT_EPS,
1094     rmax=QTT_RMAX
1095 ):
1096     """Step 3b: fold every coordinate's physical core back to binary QTT
1097     cores, concatenate them all, and finally round the WHOLE QTT once
1098     (this is where the FFT-induced rank growth actually gets truncated)."""
1099     bits_per_dim = list(
1100         map(int, bits_per_dim)
1101     )
1102
1103     if len(physical_cores) != len(
1104         bits_per_dim
1105     ):
1106         raise ValueError(
1107             "Need one physical TT core per physical dimension."
1108         )
1109
1110     qcores = []
1111
1112     for core, nbits in zip(
1113         physical_cores,
1114         bits_per_dim
1115     ):
1116
1117         qcores.extend(
1118             _split_physical_core_exact(
1119                 core,
1120                 nbits
1121             )
1122         )
1123
1124     return qtt_round(
1125         tn.Tensor(qcores),
1126         eps=eps,
1127         rmax=rmax
1128     )
1129
1130
1131 def qtt_qft(
1132     qtt,
1133     bits_per_dim,
1134     inverse=False,
1135     eps=QFT_EPS,
1136     rmax=QTT_RMAX,
1137     norm="ortho"
1138 ):
1139     """
1140     Multidimensional QFT/IQFT of a QTT wavefunction, tying the three steps
1141     together:

```

```

1142
1143     QTT --(unfold)--> one physical TT core per coordinate
1144         --(FFT each core's columns)--> transformed physical cores
1145         --(fold + round)--> binary QTT
1146
1147     Set inverse=True for the inverse transform (used to come back from
1148     momentum space to coordinate space after applying the kinetic
1149     propagator/operator).
1150     """
1151     physical_cores = qtt_unfold_to_physical_tt(
1152         qtt,
1153         bits_per_dim
1154     )
1155
1156     transformed = fft_physical_tt_cores(
1157         physical_cores,
1158         inverse=inverse,
1159         norm=norm
1160     )
1161
1162     return physical_tt_to_qtt(
1163         transformed,
1164         bits_per_dim,
1165         eps=eps,
1166         rmax=rmax
1167     )
1168
1169
1170
1171
1172 # ===== Notebook cell 11 =====
1173
1174 # Extra QTT helpers used by the n-dimensional block solver
1175
1176 def qtt_scale(qtt, scalar):
1177     """Multiply a QTT by a scalar without changing its TT ranks."""
1178     cores = [core.clone() for core in qtt.cores]
1179     scalar_t = torch.as_tensor(scalar, device=cores[0].device)
1180     dtype = torch.promote_types(cores[0].dtype, scalar_t.dtype)
1181     cores = [core.to(dtype) for core in cores]
1182     cores[0] = cores[0] * scalar_t.to(dtype)
1183     return tn.Tensor(cores)
1184
1185
1186 def qtt_conjugate(qtt):
1187     """Complex conjugate of a QTT."""
1188     return tn.Tensor([core.conj().clone() for core in qtt.cores])
1189
1190
1191 def qtt_linear_combination(states, coeffs, eps=QTT_EPS, rmax=QTT_RMAX):
1192     """Return sum_i coeffs[i] * states[i], rounding after additions."""
1193     if len(states) != len(coeffs):
1194         raise ValueError("states and coeffs must have the same length")
1195
1196     active = []
1197     for i, c in enumerate(coeffs):
1198         c_abs = float(torch.abs(torch.as_tensor(c)).detach().cpu())
1199         if c_abs > 1.0e-15:
1200             active.append((i, c))
1201
1202     if not active:

```

```

1203         return qtt_scale(states[0], 0.0)
1204
1205     i0, c0 = active[0]
1206     out = qtt_scale(states[i0], c0)
1207
1208     for i, c in active[1:]:
1209         out = qtt_add(
1210             out,
1211             qtt_scale(states[i], c),
1212             eps=eps,
1213             rmax=rmax
1214         )
1215
1216     return qtt_round(out, eps=eps, rmax=rmax)
1217
1218
1219 def tensor_product_many(qtts):
1220     """Tensor product of a list of independent-coordinate QTTs."""
1221     if len(qtts) == 0:
1222         raise ValueError("Need at least one QTT")
1223     out = qtts[0]
1224     for qtt in qtts[1:]:
1225         out = qtt_tensor_product(out, qtt)
1226     return out
1227
1228
1229 def sum_qtts(qtts, eps=QTT_EPS, rmax=QTT_RMAX):
1230     """Sum a list of QTTs."""
1231     if len(qtts) == 0:
1232         raise ValueError("Need at least one QTT")
1233     out = qtts[0]
1234     for qtt in qtts[1:]:
1235         out = qtt_add(out, qtt, eps=eps, rmax=rmax)
1236     return qtt_round(out, eps=eps, rmax=rmax)
1237
1238
1239 def qtt_relative_error(a, b):
1240     """Algebraic relative 2-norm ||a-b||/||a|| (grid weights cancel)."""
1241     diff = qtt_add(a, qtt_scale(b, -1.0), eps=QTT_EPS, rmax=QTT_RMAX)
1242     num = torch.sqrt(torch.clamp(qtt_inner(diff, diff).real, min=0.0))
1243     den = torch.sqrt(torch.clamp(qtt_inner(a, a).real, min=1.0e-300))
1244     return float((num / den).detach().cpu())
1245
1246
1247 # ===== Notebook cell 13 =====
1248
1249 # Physical parameters and n binary Fourier grids
1250
1251
1252 hbar = 1.0
1253
1254 # -----
1255 # 1. Choose the number of oscillators
1256 # -----
1257 # The code remains n-dimensional, but n=2 is the default because it gives
1258 # a stringent benchmark against analytical 2D product states and permits
1259 # full wavefunction heat maps.
1260 N_OSC = 3
1261 NSTATES = 20
1262
1263 # Accuracy-oriented default grid:

```

```

1264 # L=7 -> N=128 points per coordinate
1265 # The wide box makes the low Morse bound states essentially zero at the
1266 # periodic FFT boundaries, strongly reducing finite-box/periodicity error.
1267 def make_default_oscillators(n, L=7):
1268     oscillators = []
1269     for d in range(n):
1270         oscillators.append({
1271             "mass": 1.0 + 0.05 * d,
1272             "De": 20.0 - 1.5 * d,
1273             "a": 0.80 - 0.03 * d,
1274             "xe": 0.0,
1275             "L": L,
1276             "xmin": -6.0,
1277             "xmax": 18.0,
1278         })
1279     return oscillators
1280
1281 OSCILLATORS = make_default_oscillators(N_OSC, L=7)
1282
1283 # Example manual customization:
1284 # OSCILLATORS[1].update({"mass": 1.2, "De": 15.0, "a": 0.65})
1285
1286 # -----
1287 # 2. Choose uncoupled or coupled dynamics
1288 # -----
1289 # Keep "none" for the analytical benchmark. Switch to "morse_pair" to
1290 # study coupled modes, or "custom" for an arbitrary many-body potential.
1291 COUPLING_MODEL = "none" # "none", "morse_pair", or "custom"
1292 COUPLING_STRENGTH = 0.60 # used to build nearest-neighbor KAPPA
1293
1294 KAPPA = np.zeros((N_OSC, N_OSC), dtype=float)
1295 for d in range(N_OSC - 1):
1296     KAPPA[d, d + 1] = COUPLING_STRENGTH
1297     KAPPA[d + 1, d] = COUPLING_STRENGTH
1298
1299 # For a fully custom coupling set COUPLING_MODEL="custom" and replace
1300 # this function. X has shape (nsamples, N_OSC) and the return value must
1301 # have shape (nsamples,).
1302 def CUSTOM_COUPLING(X):
1303     return torch.zeros(X.shape[0], dtype=X.dtype, device=X.device)
1304
1305 if len(OSCILLATORS) != N_OSC:
1306     raise ValueError("len(OSCILLATORS) must equal N_OSC")
1307 if KAPPA.shape != (N_OSC, N_OSC):
1308     raise ValueError("KAPPA must have shape (N_OSC, N_OSC)")
1309 if not np.allclose(KAPPA, KAPPA.T):
1310     raise ValueError("KAPPA must be symmetric")
1311 if not np.allclose(np.diag(KAPPA), 0.0):
1312     raise ValueError("The diagonal of KAPPA should be zero")
1313
1314 # -----
1315 # 3. Build coordinate and momentum grids
1316 # -----
1317 BITS_PER_DIM = [int(p["L"]) for p in OSCILLATORS]
1318 GRID_SHAPE = [2**L for L in BITS_PER_DIM]
1319 TOTAL_BITS = sum(BITS_PER_DIM)
1320 TOTAL_POINTS = int(np.prod(GRID_SHAPE))
1321
1322 XGRID, PGRID, DX, T1D, V1D, MORSE_S1D = [], [], [], [], [], []
1323
1324 for d, p in enumerate(OSCILLATORS):

```

```

1325     N = GRID_SHAPE[d]
1326     dx = (p["xmax"] - p["xmin"]) / N
1327     x = p["xmin"] + dx * torch.arange(N, dtype=torch.get_default_dtype(), device=device)
1328     freq = torch.fft.fftfreq(N, d=dx, device=device)
1329     momentum = 2.0 * np.pi * hbar * freq
1330
1331     s = 1.0 - torch.exp(-p["a"] * (x - p["xe"]))
1332     V = p["De"] * s**2
1333     T = momentum**2 / (2.0 * p["mass"])
1334
1335     DX.append(dx)
1336     XGRID.append(x)
1337     PGRID.append(momentum)
1338     MORSE_S1D.append(s)
1339     V1D.append(V)
1340     T1D.append(T)
1341
1342     dvol = float(np.prod(DX))
1343
1344     parameter_table = pd.DataFrame([
1345         {
1346             "d": d,
1347             "mass": p["mass"],
1348             "De": p["De"],
1349             "a": p["a"],
1350             "xe": p["xe"],
1351             "L": p["L"],
1352             "N": GRID_SHAPE[d],
1353             "xmin": p["xmin"],
1354             "xmax": p["xmax"],
1355             "dx": DX[d],
1356         }
1357     ])
1358
1359     print(f"N_OSC      = {N_OSC}")
1360     print(f"NSTATES     = {NSTATES}")
1361     print(f"Coupling      = {COUPLING_MODEL}")
1362     print(f"BITS_PER_DIM   = {BITS_PER_DIM}")
1363     print(f"TOTAL_BITS     = {TOTAL_BITS}")
1364     print(f"Dense points   = {TOTAL_POINTS:,}")
1365     print(f"dvol          = {dvol:.6e}")
1366     display(parameter_table)
1367
1368     if COUPLING_MODEL == "morse_pair":
1369         print("\nPair-coupling matrix KAPPA:")
1370         display(pd.DataFrame(KAPPA))
1371
1372
1373
1374
1375     # ===== Notebook cell 14 =====
1376
1377     # Sanity plot: one-dimensional Morse wells
1378     #
1379     # ylim(-Dmax, Dmax) is a plotting choice only. The unshifted Morse
1380     # potential itself still has minimum 0 and dissociation limit +D.
1381     Dmax = max(p["De"] for p in OSCILLATORS)
1382
1383     plt.figure(figsize=(9, 5))
1384     for d in range(N_OSC):
1385         plt.plot(

```

```

1386         XGRID[d].detach().cpu().numpy(),
1387         V1D[d].detach().cpu().numpy(),
1388         label=f"mode {d}"
1389     )
1390 plt.ylim(0, Dmax)
1391 plt.xlabel("coordinate")
1392 plt.ylabel("Morse potential")
1393 plt.title(f"One-dimensional Morse terms for {N_OSC} modes")
1394 plt.grid(alpha=0.25)
1395 plt.legend()
1396 plt.tight_layout()
1397 plt.show()
1398
1399
1400
1401 # ===== Notebook cell 16 =====
1402
1403 # Build V(x1,...,xn) in QTT form.
1404 # Built-in uncoupled and pairwise Morse models use exact structured QTT
1405 # algebra; arbitrary custom couplings use TT-cross.
1406
1407 CROSS_EPS = 1.0e-11
1408 CROSS_RMAX = QTT_RMAX
1409 CROSS_MAX_ITER = 40
1410 CROSS_VAL_SIZE = 4000
1411 CROSS_SEED = 1234
1412
1413 np.random.seed(CROSS_SEED)
1414 torch.manual_seed(CROSS_SEED)
1415
1416 # Reusable one-dimensional QTT factors.
1417 ONES_1D_QTT = [
1418     qtt_from_dense(
1419         torch.ones(GRID_SHAPE[d], dtype=torch.get_default_dtype(), device=device),
1420         [BITS_PER_DIM[d]], eps=1.0e-14, rmax=QTT_RMAX
1421     )
1422     for d in range(N_OSC)
1423 ]
1424 V_1D_QTT = [
1425     qtt_from_dense(V1D[d], [BITS_PER_DIM[d]], eps=QTT_EPS, rmax=QTT_RMAX)
1426     for d in range(N_OSC)
1427 ]
1428 S_1D_QTT = [
1429     qtt_from_dense(MORSE_S1D[d], [BITS_PER_DIM[d]], eps=QTT_EPS, rmax=QTT_RMAX)
1430     for d in range(N_OSC)
1431 ]
1432
1433
1434 def embedded_product(factors):
1435     """Tensor product of one 1D QTT factor for every physical coordinate."""
1436     return tensor_product_many(factors)
1437
1438
1439 def add_terms_qtt(terms):
1440     if not terms:
1441         raise ValueError("Need at least one QTT term")
1442     out = terms[0]
1443     for term in terms[1:]:
1444         out = qtt_add(out, term, eps=QTT_EPS, rmax=QTT_RMAX)
1445     return qtt_round(out, eps=QTT_EPS, rmax=QTT_RMAX)
1446

```

```

1447
1448 def potential_nd(X):
1449     """Vectorized n-dimensional Morse PES evaluated at X[nsamples, N_OSC]."""
1450     if X.ndim != 2 or X.shape[1] != N_OSC:
1451         raise ValueError("X must have shape (nsamples, N_OSC)")
1452
1453     V = torch.zeros(X.shape[0], dtype=X.dtype, device=X.device)
1454     s_values = []
1455     for d, p in enumerate(OSCILLATORS):
1456         s = 1.0 - torch.exp(-p["a"] * (X[:, d] - p["xe"]))
1457         s_values.append(s)
1458         V = V + p["De"] * s**2
1459
1460     if COUPLING_MODEL == "morse_pair":
1461         kappa = torch.as_tensor(KAPPA, dtype=X.dtype, device=X.device)
1462         for i in range(N_OSC):
1463             for j in range(i + 1, N_OSC):
1464                 if KAPPA[i, j] != 0.0:
1465                     V = V + kappa[i, j] * s_values[i] * s_values[j]
1466     elif COUPLING_MODEL == "custom":
1467         V = V + CUSTOM_COUPLING(X)
1468     elif COUPLING_MODEL != "none":
1469         raise ValueError("Unknown COUPLING_MODEL")
1470
1471     return V
1472
1473
1474 if COUPLING_MODEL in ("none", "morse_pair"):
1475     terms = []
1476
1477     # Sum_d V_d(x_d)
1478     for d in range(N_OSC):
1479         factors = [V_1D_QTT[q] if q == d else ONES_1D_QTT[q] for q in range(N_OSC)]
1480         terms.append(embedded_product(factors))
1481
1482     # Sum_{i<j} kappa_ij s_i s_j
1483     if COUPLING_MODEL == "morse_pair":
1484         for i in range(N_OSC):
1485             for j in range(i + 1, N_OSC):
1486                 if KAPPA[i, j] == 0.0:
1487                     continue
1488                 factors = []
1489                 for q in range(N_OSC):
1490                     if q == i or q == j:
1491                         factors.append(S_1D_QTT[q])
1492                     else:
1493                         factors.append(ONES_1D_QTT[q])
1494                 terms.append(qtt_scale(embedded_product(factors), KAPPA[i, j]))
1495
1496     VQTT = add_terms_qtt(terms)
1497     VQTT_CROSS_INFO = {"method": "structured", "val_epss": [], "nsamples": 0}
1498     print("Built VQTT by structured QTT algebra (no TT-cross approximation).")
1499
1500 else:
1501     # TT-cross path for arbitrary custom many-body potentials.
1502     BIT_WEIGHTS = [
1503         2.0 ** torch.arange(L - 1, -1, -1, dtype=torch.get_default_dtype(), device=device)
1504         for L in BITS_PER_DIM
1505     ]
1506     BIT_OFFSETS = np.cumsum([0] + BITS_PER_DIM)
1507

```

```

1508 def binary_samples_to_coordinates(B):
1509     B = B.to(dtype=torch.get_default_dtype(), device=device)
1510     coords = []
1511     for d in range(N_OSC):
1512         i0, i1 = BIT_OFFSETS[d], BIT_OFFSETS[d + 1]
1513         idx = B[:, i0:i1] @ BIT_WEIGHTS[d]
1514         p = OSCILLATORS[d]
1515         coords.append(p["xmin"] + DX[d] * idx)
1516     return torch.stack(coords, dim=1)
1517
1518 def potential_from_binary_samples(B):
1519     return potential_nd(binary_samples_to_coordinates(B))
1520
1521 BIT_DOMAIN = [
1522     torch.tensor([0.0, 1.0], dtype=torch.get_default_dtype(), device=device)
1523     for _ in range(TOTAL_BITS)
1524 ]
1525
1526 print(f"Building custom VQTT over {TOTAL_BITS} binary modes by TT-cross...")
1527 VQTT, VQTT_CROSS_INFO = tn.cross(
1528     function=potential_from_binary_samples,
1529     domain=BIT_DOMAIN,
1530     function_arg="matrix",
1531     eps=CROSS_EPS,
1532     rmax=CROSS_RMAX,
1533     max_iter=CROSS_MAX_ITER,
1534     val_size=CROSS_VAL_SIZE,
1535     verbose=True,
1536     return_info=True
1537 )
1538 VQTT = qtt_round(VQTT, eps=QTT_EPS, rmax=QTT_RMAX)
1539
1540 Vtot_qtt = VQTT
1541 qtt_summary("VQTT", VQTT)
1542 val_epss = VQTT_CROSS_INFO.get("val_epss", [])
1543 if len(val_epss):
1544     print(f"TT-cross validation relative error = {val_epss[-1]:.3e}")
1545 if VQTT_CROSS_INFO.get("method") == "structured":
1546     print("Potential construction error is controlled only by QTT rounding.")
1547 else:
1548     print("TT-cross samples =", VQTT_CROSS_INFO.get("nsamples", "n/a"))
1549
1550
1551 # ===== Notebook cell 18 =====
1552
1553 # Full kinetic diagonal T(p1,...,pn) in QTT form
1554
1555 T_1D_QTT = [
1556     qtt_from_dense(T1D[d], [BITS_PER_DIM[d]], eps=1.0e-12, rmax=QTT_RMAX)
1557     for d in range(N_OSC)
1558 ]
1559
1560
1561 def embed_1d_qtt(qtt_1d, target_dim):
1562     factors = [
1563         qtt_1d if d == target_dim else ONES_1D_QTT[d]
1564         for d in range(N_OSC)
1565     ]
1566     return tensor_product_many(factors)
1567
1568 kinetic_terms = [

```

```

1569     embed_1d_qtt(T_1D_QTT[d], d)
1570     for d in range(N_OSC)
1571 ]
1572
1573 Ttot_qtt = sum_qtts(kinetic_terms, eps=QTT_EPS, rmax=QTT_RMAX)
1574 qtt_summary("T total", Ttot_qtt)
1575
1576
1577
1578 # ===== Notebook cell 20 =====
1579
1580 # QFT -> inverse-QFT round-trip validation on an n-dimensional product test state
1581
1582 test_factors = []
1583 for d in range(N_OSC):
1584     x = XGRID[d]
1585     p = OSCILLATORS[d]
1586     width = 1.0 + 0.15 * d
1587     f = torch.exp(-0.5 * ((x - p["xe"]) / width)**2) * torch.cos((0.7 + 0.1*d) * x)
1588     test_factors.append(
1589         qtt_from_dense(f, [BITS_PER_DIM[d]], eps=QTT_EPS, rmax=QTT_RMAX)
1590     )
1591
1592 test_qtt = tensor_product_many(test_factors)
1593 test_k_qtt = qtt_qft(test_qtt, BITS_PER_DIM, inverse=False, eps=QFT_EPS, rmax=QTT_RMAX)
1594 test_back_qtt = qtt_qft(test_k_qtt, BITS_PER_DIM, inverse=True, eps=QFT_EPS, rmax=QTT_RMAX)
1595
1596 print(f"QFT round-trip relative QTT error = {qtt_relative_error(test_qtt, test_back_qtt):.3e}")
1597 qtt_summary("test x", test_qtt)
1598 qtt_summary("test k", test_k_qtt)
1599
1600
1601
1602 # ===== Notebook cell 22 =====
1603
1604 # QTT-native Hamiltonian action
1605
1606 def apply_H_qtt(psi_qtt, eps=QTT_EPS, qft_eps=QFT_EPS):
1607     Vpsi = qtt_hadamard(Vtot_qtt, psi_qtt, eps=eps, rmax=QTT_RMAX)
1608
1609     psi_p = qtt_qft(
1610         psi_qtt, BITS_PER_DIM, inverse=False,
1611         eps=qft_eps, rmax=QTT_RMAX
1612     )
1613
1614     Tpsi_p = qtt_hadamard(Ttot_qtt, psi_p, eps=eps, rmax=QTT_RMAX)
1615
1616     Tpsi = qtt_qft(
1617         Tpsi_p, BITS_PER_DIM, inverse=True,
1618         eps=qft_eps, rmax=QTT_RMAX
1619     )
1620
1621     return qtt_add(Vpsi, Tpsi, eps=eps, rmax=QTT_RMAX)
1622
1623
1624 def qtt_energy(psi_qtt):
1625     """Rayleigh quotient. The common dvol factor cancels."""
1626     Hpsi = apply_H_qtt(psi_qtt)
1627     return (qtt_inner(psi_qtt, Hpsi) / qtt_inner(psi_qtt, psi_qtt)).real
1628
1629

```

```

1630 def qtt_normalize(psi_qtt):
1631     norm2 = dvol * qtt_inner(psi_qtt, psi_qtt).real
1632     norm = torch.sqrt(torch.clamp(norm2, min=1.0e-300))
1633     return qtt_scale(psi_qtt, 1.0 / norm)
1634
1635
1636
1637 # ===== Notebook cell 24 =====
1638
1639 # QTT-native block overlap, orthogonalization, and Rayleigh-Ritz
1640
1641 ORTHO_EIG_FLOOR = 1.0e-12
1642
1643 def qtt_overlap_matrix(states):
1644     M = len(states)
1645     S = torch.zeros((M, M), dtype=torch.complex128, device=device)
1646
1647     for i in range(M):
1648         for j in range(i, M):
1649             value = dvol * qtt_inner(states[i], states[j])
1650             S[i, j] = value
1651             S[j, i] = value.conj()
1652
1653     return 0.5 * (S + S.conj().T)
1654
1655
1656 def rotate_qtt_block(states, C, eps=QTT_EPS, rmax=QTT_RMAX):
1657     """new_state_j = sum_i old_state_i * C[i,j]."""
1658     return [
1659         qtt_linear_combination(states, C[:, j], eps=eps, rmax=rmax)
1660         for j in range(C.shape[1])
1661     ]
1662
1663
1664 def orthonormalize_qtt_block(states, passes=2):
1665     """Lowdin symmetric orthogonalization performed entirely with QTT overlaps."""
1666     states = list(states)
1667
1668     for _ in range(passes):
1669         S = qtt_overlap_matrix(states)
1670         evals, U = torch.linalg.eigh(S)
1671
1672         min_eval = float(evals.min().real.detach().cpu())
1673         if min_eval < ORTHO_EIG_FLOOR:
1674             print(f"Warning: small block-overlap eigenvalue {min_eval:.3e}; regularizing.")
1675
1676         evals_safe = torch.clamp(evals.real, min=ORTH0_EIG_FLOOR)
1677         Sinvhalf = U @ torch.diag(evals_safe.rsqrt().to(torch.complex128)) @ U.conj().T
1678         states = rotate_qtt_block(states, Sinvhalf)
1679
1680     return [qtt_normalize(state) for state in states]
1681
1682
1683 def rayleigh_ritz_qtt(states):
1684     """QTT block orthogonalization + projected-Hamiltonian diagonalization."""
1685     states = orthonormalize_qtt_block(states, passes=2)
1686     M = len(states)
1687
1688     Hstates = [apply_H_qtt(state) for state in states]
1689     Hsub = torch.zeros((M, M), dtype=torch.complex128, device=device)
1690

```

```

1691     for i in range(M):
1692         for j in range(M):
1693             Hsub[i, j] = dvol * qtt_inner(states[i], Hstates[j])
1694
1695     Hsub = 0.5 * (Hsub + Hsub.conj().T)
1696     evals, U = torch.linalg.eigh(Hsub)
1697
1698     rotated = rotate_qtt_block(states, U)
1699     rotated = [qtt_normalize(state) for state in rotated]
1700
1701     return rotated, evals.real
1702
1703
1704
1705 # ===== Notebook cell 26 =====
1706
1707 # Exact 1D Morse spectra and the k lowest product combinations
1708
1709 def morse_lambda(p):
1710     return np.sqrt(2.0 * p["mass"] * p["De"]) / (hbar * p["a"])
1711
1712 def exact_energy_1d(n, p, lam):
1713     s = lam - n - 0.5
1714     if s <= 0:
1715         raise ValueError(f"n={n} is not a bound level")
1716     return p["De"] - (hbar**2 * p["a"]**2 / (2.0 * p["mass"])) * s**2
1717
1718 def bound_energies_1d(p):
1719     lam = morse_lambda(p)
1720     values = []
1721     n = 0
1722     while lam - n - 0.5 > 0:
1723         values.append(exact_energy_1d(n, p, lam))
1724         n += 1
1725     return lam, values
1726
1727 LAMBDA = []
1728 E1D_EXACT = []
1729 for p in OSCILLATORS:
1730     lam, energies = bound_energies_1d(p)
1731     LAMBDA.append(lam)
1732     E1D_EXACT.append(energies)
1733
1734 def k_lowest_product_levels(energy_lists, k):
1735     """Return k smallest sums from sorted per-mode energy lists without full enumeration."""
1736     n = len(energy_lists)
1737     if any(len(e) == 0 for e in energy_lists):
1738         return []
1739
1740     start = tuple([0] * n)
1741     start_E = sum(energy_lists[d][0] for d in range(n))
1742     heap = [(start_E, start)]
1743     visited = {start}
1744     out = []
1745
1746     while heap and len(out) < k:
1747         E, label = heapq.heappop(heap)
1748         out.append((float(E), label))
1749
1750         for d in range(n):
1751             if label[d] + 1 >= len(energy_lists[d]):

```

```

1752         continue
1753         new_label = list(label)
1754         new_label[d] += 1
1755         new_label = tuple(new_label)
1756         if new_label in visited:
1757             continue
1758         visited.add(new_label)
1759         new_E = sum(energy_lists[q][new_label[q]] for q in range(n))
1760         heapq.heappush(heap, (new_E, new_label))
1761
1762     return out
1763
1764 target_levels = k_lowest_product_levels(E1D_EXACT, NSTATES)
1765 if len(target_levels) < NSTATES:
1766     raise RuntimeError(f"Only {len(target_levels)} bound product states are available")
1767
1768 for d in range(N_OSC):
1769     print(f"mode {d}: lambda={LAMBDA[d]:.6f}, bound levels={len(E1D_EXACT[d])}")
1770
1771 print(f"\nLowest {NSTATES} uncoupled product levels:")
1772 for k, (E, label) in enumerate(target_levels):
1773     print(f"  {k:2d}: n={label}, E={E:.12f}")
1774
1775 exact_energy_sorted = np.array([E for E, _ in target_levels], dtype=float)
1776
1777
1778 # ===== Notebook cell 28 =====
1779
1780 # Exact 1D Morse states and n-dimensional product reference QTTs
1781
1782 def exact_state_1d(n, d):
1783     p = OSCILLATORS[d]
1784     lam = LAMBDA[d]
1785     x = XGRID[d].detach().cpu().numpy()
1786     s = lam - n - 0.5
1787
1788     if s <= 0:
1789         raise ValueError(f"n={n} is not bound for mode {d}")
1790
1791     z = 2.0 * lam * np.exp(-p["a"] * (x - p["xe"]))
1792     log_prefactor = s * np.log(z) - 0.5 * z
1793     prefactor = np.exp(np.maximum(log_prefactor, -745.0))
1794     psi = prefactor * eval_genlaguerre(n, 2.0 * s, z)
1795     psi = psi / np.sqrt(DX[d] * np.vdot(psi, psi).real)
1796
1797     return torch.tensor(psi, dtype=torch.complex128, device=device)
1798
1799 _exact_1d_qtt_cache = {}
1800
1801 def exact_1d_qtt(n, d):
1802     key = (d, n)
1803     if key not in _exact_1d_qtt_cache:
1804         _exact_1d_qtt_cache[key] = qtt_from_dense(
1805             exact_state_1d(n, d),
1806             [BITS_PER_DIM[d]],
1807             eps=QTT_EPS,
1808             rmax=QTT_RMAX
1809         )
1810     return _exact_1d_qtt_cache[key]
1811
1812

```

```

1813 def exact_product_state_qtt(label):
1814     factors = [exact_1d_qtt(label[d], d) for d in range(N_OSC)]
1815     return qtt_normalize(tensor_product_many(factors))
1816
1817
1818
1819 # ===== Notebook cell 30 =====
1820
1821 # Local 1D Fourier-grid Hamiltonian eigenstates used as numerical seeds
1822
1823 E1D_FGH = []
1824 PSI1D_FGH = []
1825
1826 # We only need enough local states to cover the quantum numbers appearing
1827 # in the requested low-energy product-state block.
1828 max_needed_per_mode = [max(label[d] for _, label in target_levels) for d in range(N_OSC)]
1829
1830 for d in range(N_OSC):
1831     N = GRID_SHAPE[d]
1832
1833     # F is the unitary discrete FFT matrix. The kinetic operator in the
1834     # coordinate representation is  $F^\dagger \text{diag}(T) F$ .
1835     eye = torch.eye(N, dtype=torch.complex128, device=device)
1836     F = torch.fft.fft(eye, dim=0, norm="ortho")
1837     Tmat = F.conj().T @ (T1D[d].to(torch.complex128)[: , None] * F)
1838     H1 = Tmat + torch.diag(V1D[d].to(torch.complex128))
1839     H1 = 0.5 * (H1 + H1.conj().T)
1840
1841     evals, evecs = torch.linalg.eigh(H1)
1842
1843     # torch.linalg.eigh normalizes  $\sum_i |c_i|^2 = 1$ . Convert this to the
1844     # physical quadrature convention  $\Delta x \sum_i |\psi_i|^2 = 1$ .
1845     evecs = evecs / math.sqrt(DX[d])
1846
1847     E1D_FGH.append(evals.real)
1848     PSI1D_FGH.append(evecs)
1849
1850 # Explicitly verify local grid quality against the independent analytical
1851 # Morse spectrum before building the multidimensional block.
1852 local_rows = []
1853 for d in range(N_OSC):
1854     for n in range(max_needed_per_mode[d] + 1):
1855         E_exact = E1D_EXACT[d][n]
1856         E_fgh = float(E1D_FGH[d][n].detach().cpu())
1857         local_rows.append({
1858             "mode": d,
1859             "n": n,
1860             "E_exact": E_exact,
1861             "E_1D_FGH": E_fgh,
1862             "abs_error": abs(E_fgh - E_exact),
1863         })
1864
1865 print("1D grid-quality check (numerical FGH vs analytical Morse levels):")
1866 display(pd.DataFrame(local_rows))
1867
1868 _seed_1d_qtt_cache = {}
1869
1870 def local_fgh_seed_1d_qtt(n, d):
1871     key = (d, n)
1872     if key not in _seed_1d_qtt_cache:
1873         vec = PSI1D_FGH[d][:, n]

```

```

1874     _seed_1d_qtt_cache[key] = qtt_from_dense(
1875         vec,
1876         [BITS_PER_DIM[d]],
1877         eps=QTT_EPS,
1878         rmax=QTT_RMAX
1879     )
1880     return _seed_1d_qtt_cache[key]
1881
1882 seed_labels = [label for _, label in target_levels]
1883 states_qtt = []
1884 for label in seed_labels:
1885     factors = [local_fgh_seed_1d_qtt(label[d], d) for d in range(N_OSC)]
1886     states_qtt.append(qtt_normalize(tensor_product_many(factors)))
1887
1888 print("\nLocal-FGH product seeds:")
1889 for k, qtt in enumerate(states_qtt):
1890     qtt_summary(f"seed[{k}]", qtt)
1891
1892 print("\nInitial QTT Rayleigh-Ritz rotation...")
1893 states_qtt, E0 = rayleigh_ritz_qtt(states_qtt)
1894 print("Initial Ritz energies:")
1895 for k, E in enumerate(E0.detach().cpu().numpy()):
1896     print(f"  {k:2d}: {E:.12f}")
1897
1898
1899
1900 # ===== Notebook cell 32 =====
1901
1902 # Split-operator propagators and one QTT imaginary-time step
1903
1904 def make_qtt_propagators(dtau):
1905     # For an uncoupled potential V=sum_d V_d, exp(-dtau V/2) factorizes
1906     # exactly. This removes TT-cross error from the default benchmark.
1907     if COUPLING_MODEL == "none":
1908         potential_factors = []
1909         for d in range(N_OSC):
1910             Uv_d = torch.exp(-0.5 * dtau * V1D[d])
1911             potential_factors.append(
1912                 qtt_from_dense(Uv_d, [BITS_PER_DIM[d]], eps=QTT_EPS, rmax=QTT_RMAX)
1913             )
1914         Uv_qtt = tensor_product_many(potential_factors)
1915         info = {"method": "factorized", "val_epss": []}
1916     else:
1917         # Coupled V does not factorize, so apply the scalar map to VQTT by TT-cross.
1918         Uv_qtt, info = tn.cross(
1919             function=lambda v: torch.exp(-0.5 * dtau * v),
1920             tensors=[VQTT],
1921             eps=CROSS_EPS,
1922             rmax=CROSS_RMAX,
1923             max_iter=CROSS_MAX_ITER,
1924             val_size=CROSS_VAL_SIZE,
1925             verbose=False,
1926             return_info=True
1927         )
1928         Uv_qtt = qtt_round(Uv_qtt, eps=QTT_EPS, rmax=QTT_RMAX)
1929
1930     # Kinetic propagator always factorizes because T=sum_d T_d.
1931     kinetic_factors = []
1932     for d in range(N_OSC):
1933         Ut_d = torch.exp(-dtau * T1D[d])
1934         kinetic_factors.append(

```

```

1935         qtt_from_dense(Ut_d, [BITS_PER_DIM[d]], eps=QTT_EPS, rmax=QTT_RMAX)
1936     )
1937
1938     Ut_qtt = tensor_product_many(kinetic_factors)
1939     return Uv_qtt, Ut_qtt, info
1940
1941
1942 def imaginary_time_step_qtt(psi_qtt, Uv_qtt, Ut_qtt):
1943     # Second-order Strang splitting:
1944     #  $\exp(-dt H) = \exp(-dt V/2) \exp(-dt T) \exp(-dt V/2) + O(dt^3)$ 
1945     psi = qtt_hadamard(Uv_qtt, psi_qtt, eps=QTT_EPS, rmax=QTT_RMAX)
1946     psi = qtt_qft(psi, BITS_PER_DIM, inverse=False, eps=QTT_EPS, rmax=QTT_RMAX)
1947     psi = qtt_hadamard(Ut_qtt, psi, eps=QTT_EPS, rmax=QTT_RMAX)
1948     psi = qtt_qft(psi, BITS_PER_DIM, inverse=True, eps=QTT_EPS, rmax=QTT_RMAX)
1949     psi = qtt_hadamard(Uv_qtt, psi, eps=QTT_EPS, rmax=QTT_RMAX)
1950     return qtt_normalize(psi)
1951
1952
1953
1954 # ===== Notebook cell 34 =====
1955
1956 # n-state block imaginary-time propagation entirely in QTT form
1957 #
1958 # With local-FGH seeds the uncoupled benchmark starts extremely close to
1959 # the target eigenstates, so we use small time steps primarily as a QTT
1960 # propagation consistency test. Coupled problems need stronger filtering.
1961 if COUPLING_MODEL == "none":
1962     schedule = [
1963         (5.0e-3, 3),
1964         (1.0e-3, 4),
1965         (2.0e-4, 5),
1966     ]
1967 else:
1968     schedule = [
1969         (2.0e-2, 6),
1970         (5.0e-3, 8),
1971         (1.0e-3, 10),
1972         (2.0e-4, 10),
1973     ]
1974
1975 print(f"Starting {NSTATES}-state QTT imaginary-time propagation for {N_OSC} modes...\n")
1976 history = []
1977 global_step = 0
1978 watch = sorted(set([0, min(1, NSTATES-1), NSTATES//2, NSTATES-1]))
1979
1980 for stage, (dtau, nsteps) in enumerate(schedule, start=1):
1981     Uv_qtt, Ut_qtt, info = make_qtt_propagators(dtau)
1982     val_epss = info.get("val_epss", [])
1983     cross_err = val_epss[-1] if len(val_epss) else np.nan
1984     prop_method = info.get("method", "tt-cross")
1985
1986     print(
1987         f"Stage {stage}/{len(schedule)}: dtau={dtau:g}, steps={nsteps}, "
1988         f"Uv method={prop_method}, cross error={cross_err:.3e}"
1989     )
1990
1991     for local_step in range(1, nsteps + 1):
1992         propagated = [
1993             imaginary_time_step_qtt(state, Uv_qtt, Ut_qtt)
1994             for state in states_qtt
1995         ]

```

```

1996
1997     states_qtt, Esub = rayleigh_ritz_qtt(propagated)
1998     global_step += 1
1999
2000     E_np = Esub.detach().cpu().numpy()
2001     max_rank_now = max(max(qtt_ranks(qtt)) for qtt in states_qtt)
2002
2003     history.append({
2004         "global_step": global_step,
2005         "stage": stage,
2006         "local_step": local_step,
2007         "dtau": dtau,
2008         "max_rank": max_rank_now,
2009         **{f"E{k}": E_np[k] for k in range(NSTATES)},
2010     })
2011
2012     report = ", ".join([f"E[{k}]={E_np[k]:.10f}" for k in watch])
2013     print(
2014         f"  step {local_step:3d}/{nsteps}; "
2015         f"max rank={max_rank_now:3d}; {report}"
2016     )
2017
2018     print()
2019
2020 # Final cleanup rotation and independent Rayleigh quotients.
2021 states_qtt, E_final_rr = rayleigh_ritz_qtt(states_qtt)
2022 E_qtt_final = np.array([
2023     float(qtt_energy(state).detach().cpu())
2024     for state in states_qtt
2025 ])
2026
2027 order = np.argsort(E_qtt_final)
2028 states_qtt = [states_qtt[i] for i in order]
2029 E_qtt_final = E_qtt_final[order]
2030
2031 print("Final QTT energies:")
2032 for k, E in enumerate(E_qtt_final):
2033     print(f"  {k:2d}: {E:.12f}")
2034
2035
2036
2037 # ===== Notebook cell 36 =====
2038
2039 # Optional dense validation
2040
2041 DENSE_VALIDATION_LIMIT = 200_000
2042 DENSE_VALIDATION = TOTAL_POINTS <= DENSE_VALIDATION_LIMIT
2043
2044 E_dense_final = None
2045
2046 if DENSE_VALIDATION:
2047     print(f"Dense validation enabled ({TOTAL_POINTS:,} points <= {DENSE_VALIDATION_LIMIT:,}).")
2048
2049     # Build dense coordinate-space potential only for this validation branch.
2050     meshes = torch.meshgrid(*XGRID, indexing="ij")
2051     X_dense = torch.stack([mesh.reshape(-1) for mesh in meshes], dim=1)
2052     Vdiag_dense = potential_nd(X_dense).reshape(GRID_SHAPE)
2053
2054     # Dense kinetic diagonal by broadcasting the 1D T_d arrays.
2055     Tdiag_dense = torch.zeros(GRID_SHAPE, dtype=torch.get_default_dtype(), device=device)
2056     for d in range(N_OSC):

```

```

2057     shape = [1] * N_OSC
2058     shape[d] = GRID_SHAPE[d]
2059     Tdiag_dense = Tdiag_dense + T1D[d].reshape(shape)
2060
2061     def apply_H_dense(Psi):
2062         squeeze = False
2063         if Psi.ndim == 1:
2064             Psi = Psi[:, None]
2065             squeeze = True
2066
2067         M = Psi.shape[1]
2068         field = Psi.reshape(*GRID_SHAPE, M)
2069         fft_dims = tuple(range(N_OSC))
2070         field_p = torch.fft.fftn(field, dim=fft_dims, norm="ortho")
2071
2072         Tfield = torch.fft.ifftn(
2073             Tdiag_dense[..., None] * field_p,
2074             dim=fft_dims,
2075             norm="ortho"
2076         )
2077         Hfield = Vdiag_dense[..., None] * field + Tfield
2078         out = Hfield.reshape(TOTAL_POINTS, M)
2079         return out[:, 0] if squeeze else out
2080
2081     Psi_dense = torch.stack([
2082         qtt_to_dense(state, BITS_PER_DIM).reshape(-1)
2083         for state in states_qtt
2084     ], dim=1)
2085
2086     HPsi_dense = apply_H_dense(Psi_dense)
2087     numerator = dvol * torch.sum(Psi_dense.conj() * HPsi_dense, dim=0)
2088     denominator = dvol * torch.sum(Psi_dense.conj() * Psi_dense, dim=0)
2089     E_dense_final = (numerator / denominator).real.detach().cpu().numpy()
2090
2091     print("\nQTT vs dense-FGH expectation values:")
2092     dense_check = pd.DataFrame({
2093         "state": np.arange(NSTATES),
2094         "E_QTT": E_qtt_final,
2095         "E_dense_FGH": E_dense_final,
2096         "abs_difference": np.abs(E_qtt_final - E_dense_final),
2097     })
2098     display(dense_check)
2099 else:
2100     print(
2101         f"Dense validation skipped: {TOTAL_POINTS:,} points > "
2102         f"{DENSE_VALIDATION_LIMIT:,}. The QTT calculation itself is unaffected."
2103     )
2104
2105
2106 # ===== Notebook cell 38 =====
2107
2108 # Analyze final states using QTT overlaps and Hamiltonian residuals
2109
2110 def qtt_squared_overlap(a, b):
2111     ov = dvol * qtt_inner(a, b)
2112     na = dvol * qtt_inner(a, a).real
2113     nb = dvol * qtt_inner(b, b).real
2114     return float((torch.abs(ov)**2 / (na * nb)).real.detach().cpu())
2115
2116
2117

```

```

2118 def qtt_residual_norm(psi_qtt, energy):
2119     """Physical L2 norm ||H psi - E psi|| / ||psi||."""
2120     Hpsi = apply_H_qtt(psi_qtt)
2121     residual = qtt_add(
2122         Hpsi,
2123         qtt_scale(psi_qtt, -energy),
2124         eps=QTT_EPS,
2125         rmax=QTT_RMAX
2126     )
2127     num = dvol * qtt_inner(residual, residual).real
2128     den = dvol * qtt_inner(psi_qtt, psi_qtt).real
2129     return float(torch.sqrt(torch.clamp(num / den, min=0.0)).detach().cpu())
2130
2131
2132 residuals_qtt = [qtt_residual_norm(state, E_qtt_final[j]) for j, state in enumerate(states_qtt)]
2133
2134 if COUPLING_MODEL == "none":
2135     reference_levels = target_levels
2136     reference_qtts = [exact_product_state_qtt(label) for _, label in reference_levels]
2137
2138     overlap_matrix = np.zeros((NSTATES, NSTATES), dtype=float)
2139     for i in range(NSTATES):
2140         for j in range(NSTATES):
2141             overlap_matrix[i, j] = qtt_squared_overlap(reference_qtts[i], states_qtt[j])
2142
2143     exact_idx, num_idx = linear_sum_assignment(-overlap_matrix)
2144     assigned = {int(i): int(j) for i, j in zip(exact_idx, num_idx)}
2145
2146     rows = []
2147     for i, (E_exact, label) in enumerate(reference_levels):
2148         j = assigned[i]
2149         abs_err = abs(E_qtt_final[j] - E_exact)
2150         rows.append({
2151             "exact_rank": i,
2152             "numerical_index": j,
2153             "quantum_numbers": str(label),
2154             "E_exact": E_exact,
2155             "E_QTT": E_qtt_final[j],
2156             "abs_energy_error": abs_err,
2157             "rel_energy_error": abs_err / max(abs(E_exact), 1.0e-15),
2158             "fidelity": overlap_matrix[i, j],
2159             "residual_norm": residuals_qtt[j],
2160             "max_QTT_rank": max(qtt_ranks(states_qtt[j])),
2161             "QTT_coefficients": qtt_numcoef(states_qtt[j]),
2162         })
2163
2164     results = pd.DataFrame(rows).sort_values("exact_rank").reset_index(drop=True)
2165     print("Uncoupled exact-state comparison:")
2166     display(results)
2167
2168     print("\nSummary accuracy:")
2169     print(f"  max |E_QTT-E_exact| = {results['abs_energy_error'].max():.3e}")
2170     print(f"  min fidelity       = {results['fidelity'].min():.12f}")
2171     print(f"  max residual norm   = {results['residual_norm'].max():.3e}")
2172
2173 else:
2174     NREFERENCE = max(2 * NSTATES, NSTATES)
2175     reference_levels = k_lowest_product_levels(E1D_EXACT, NREFERENCE)
2176     reference_qtts = [exact_product_state_qtt(label) for _, label in reference_levels]
2177
2178     overlap_matrix = np.zeros((len(reference_qtts), NSTATES), dtype=float)

```

```

2179     for i in range(len(reference_qtts)):
2180         for j in range(NSTATES):
2181             overlap_matrix[i, j] = qtt_squared_overlap(reference_qtts[i], states_qtt[j])
2182
2183     rows = []
2184     for j in range(NSTATES):
2185         i_dom = int(np.argmax(overlap_matrix[:, j]))
2186         weights = overlap_matrix[:, j]
2187         captured = float(np.sum(weights))
2188         participation = (captured**2 / np.sum(weights**2)) if np.sum(weights**2) > 0 else np.nan
2189
2190         rows.append({
2191             "numerical_state": j,
2192             "E_QTT": E_qtt_final[j],
2193             "residual_norm": residuals_qtt[j],
2194             "dominant_uncoupled_label": str(reference_levels[i_dom][1]),
2195             "dominant_product_weight": overlap_matrix[i_dom, j],
2196             "captured_weight_in_reference_set": captured,
2197             "effective_product_participation": participation,
2198             "max_QTT_rank": max(qtt_ranks(states_qtt[j])),
2199             "QTT_coefficients": qtt_numcoef(states_qtt[j]),
2200         })
2201
2202     results = pd.DataFrame(rows)
2203     print("Coupled-state product-basis character:")
2204     display(results)
2205
2206     # Add dense-validation energies when available.
2207     if E_dense_final is not None:
2208         dense_map = {int(i): float(E_dense_final[i]) for i in range(NSTATES)}
2209         if COUPLING_MODEL == "none":
2210             results["E_dense_FGH"] = [dense_map[int(j)] for j in results["numerical_index"]]
2211             results["QTT_minus_dense"] = results["E_QTT"] - results["E_dense_FGH"]
2212         else:
2213             results["E_dense_FGH"] = [dense_map[int(j)] for j in results["numerical_state"]]
2214             results["QTT_minus_dense"] = results["E_QTT"] - results["E_dense_FGH"]
2215
2216     results_file = OUTPUT_DIR / "n_morse_qtt_results.csv"
2217     results.to_csv(results_file, index=False)
2218     print("Saved:", results_file)
2219
2220
2221
2222     # ===== Notebook cell 40 =====
2223
2224     # QTT-native one-coordinate marginals
2225
2226     def qtt_marginal_1d(psi_qtt, target_dim):
2227         prob_qtt = qtt_hadamard(
2228             qtt_conjugate(psi_qtt),
2229             psi_qtt,
2230             eps=QTT_EPS,
2231             rmax=QTT_RMAX
2232         )
2233
2234         cores = qtt_unfold_to_physical_tt(prob_qtt, BITS_PER_DIM)
2235
2236         # Contract all coordinates to the left of target_dim.
2237         left = torch.ones((1,), dtype=cores[0].dtype, device=device)
2238         for d in range(target_dim):
2239             summed = cores[d].sum(dim=1)

```

```

2240     left = left @ summed
2241
2242     # Contract all coordinates to the right of target_dim.
2243     right = torch.ones((1,), dtype=cores[0].dtype, device=device)
2244     for d in range(N_OSC - 1, target_dim, -1):
2245         summed = cores[d].sum(dim=1)
2246         right = summed @ right
2247
2248     marginal = torch.einsum("a,aib,b->i", left, cores[target_dim], right).real
2249
2250     # Integrate over all coordinates except target_dim.
2251     other_volume = dvol / DX[target_dim]
2252     marginal = marginal * other_volume
2253
2254     # Remove tiny negative roundoff and renormalize.
2255     marginal = torch.clamp(marginal, min=0.0)
2256     norm = DX[target_dim] * marginal.sum()
2257     marginal = marginal / torch.clamp(norm, min=1.0e-300)
2258
2259     return marginal
2260
2261 selected_states = sorted(set([0, min(1, NSTATES - 1), NSTATES - 1]))
2262
2263 for state_index in selected_states:
2264     plt.figure(figsize=(9, 5))
2265     for d in range(N_OSC):
2266         marginal = qtt_marginal_id(states_qtt[state_index], d)
2267         plt.plot(
2268             XGRID[d].detach().cpu().numpy(),
2269             marginal.detach().cpu().numpy(),
2270             label=f"mode {d}"
2271         )
2272
2273     plt.xlabel("coordinate")
2274     plt.ylabel(r"marginal probability  $P_d(x_d)$ ")
2275     plt.title(f"State {state_index}: one-coordinate marginals")
2276     plt.grid(alpha=0.25)
2277     plt.legend()
2278     plt.tight_layout()
2279     plt.show()
2280
2281
2282
2283 # ===== Notebook cell 42 =====
2284
2285 # Convergence and final-spectrum diagnostics
2286
2287 history_df = pd.DataFrame(history)
2288 history_file = OUTPUT_DIR / "n_morse_qtt_history.csv"
2289 history_df.to_csv(history_file, index=False)
2290 print("Saved:", history_file)
2291
2292 plt.figure(figsize=(10, 5))
2293 for k in range(min(NSTATES, 6)):
2294     plt.plot(history_df["global_step"], history_df[f"E{k}"], marker="o", ms=3, label=f"E{k}")
2295 plt.xlabel("QTT block iteration")
2296 plt.ylabel("Ritz energy")
2297 plt.title("Imaginary-time convergence of the lowest Ritz energies")
2298 plt.grid(alpha=0.25)
2299 plt.legend(ncol=2)
2300 plt.tight_layout()

```

```

2301 plt.show()
2302
2303 plt.figure(figsize=(10, 4))
2304 plt.plot(history_df["global_step"], history_df["max_rank"], marker="o")
2305 plt.xlabel("QTT block iteration")
2306 plt.ylabel("maximum QTT rank")
2307 plt.title("Rank growth during imaginary-time propagation")
2308 plt.grid(alpha=0.25)
2309 plt.tight_layout()
2310 plt.show()
2311
2312 plt.figure(figsize=(10, 5))
2313 plt.plot(np.arange(NSTATES), E_qtt_final, "o-", label="QTT")
2314 if COUPLING_MODEL == "none":
2315     plt.plot(np.arange(NSTATES), exact_energy_sorted, "s--", label="exact uncoupled")
2316 plt.xlabel("energy rank")
2317 plt.ylabel("energy")
2318 plt.title(f"Lowest {NSTATES} energies for {N_OSC} Morse modes")
2319 plt.grid(alpha=0.25)
2320 plt.legend()
2321 plt.tight_layout()
2322 plt.show()
2323
2324 plt.figure(figsize=(10, 6))
2325 plt.imshow(overlap_matrix, origin="lower", aspect="auto")
2326 plt.colorbar(label="squared overlap")
2327 plt.xlabel("numerical state")
2328 plt.ylabel("exact product state" if COUPLING_MODEL == "none" else "uncoupled product basis state")
2329 plt.title("QTT overlap matrix")
2330 plt.tight_layout()
2331 plt.show()
2332
2333
2334
2335 # ===== Notebook cell 44 =====
2336
2337 # Side-by-side 2D numerical/analytical heat maps
2338 #
2339 # By default show every state in the requested low-energy block. Reduce
2340 # N_HEATMAP_STATES if you want a shorter notebook output.
2341 N_HEATMAP_STATES = min(NSTATES, 12)
2342
2343 if N_OSC == 2 and COUPLING_MODEL == "none":
2344     if 'assigned' not in globals():
2345         raise RuntimeError("Run the exact-state matching cell before the heat-map cell.")
2346
2347     x1min, x1max = OSCILLATORS[0]["xmin"], OSCILLATORS[0]["xmax"]
2348     x2min, x2max = OSCILLATORS[1]["xmin"], OSCILLATORS[1]["xmax"]
2349     extent = [x2min, x2max, x1min, x1max]
2350
2351     for i in range(N_HEATMAP_STATES):
2352         E_exact, label = reference_levels[i]
2353         j = assigned[i]
2354
2355         psi_num = qtt_to_dense(states_qtt[j], BITS_PER_DIM)
2356         psi_exact = qtt_to_dense(reference_qtts[i], BITS_PER_DIM)
2357
2358         # Phase-align the numerical state to the analytical reference.
2359         overlap = dvol * qtt_inner(reference_qtts[i], states_qtt[j])
2360         if torch.abs(overlap) > 0:
2361             phase_factor = torch.conj(overlap) / torch.abs(overlap)

```

```

2362     psi_num = psi_num * phase_factor
2363
2364     psi_num_np = psi_num.detach().cpu().numpy()
2365     psi_exact_np = psi_exact.detach().cpu().numpy()
2366
2367     num_real = np.real(psi_num_np)
2368     exact_real = np.real(psi_exact_np)
2369     num_prob = np.abs(psi_num_np)**2
2370     exact_prob = np.abs(psi_exact_np)**2
2371
2372     amp_max = max(np.max(np.abs(num_real)), np.max(np.abs(exact_real)))
2373     prob_max = max(np.max(num_prob), np.max(exact_prob))
2374
2375     fidelity = overlap_matrix[i, j]
2376     E_num = E_qtt_final[j]
2377     E_err = abs(E_num - E_exact)
2378
2379     fig, axes = plt.subplots(2, 2, figsize=(11, 9), constrained_layout=True)
2380
2381     im00 = axes[0, 0].imshow(
2382         num_real, origin="lower", aspect="auto", extent=extent,
2383         cmap="RdBu_r", vmin=-amp_max, vmax=amp_max
2384     )
2385     axes[0, 0].set_title("QTT: phase-aligned Re  $\Psi$ ")
2386     axes[0, 0].set_ylabel(r" $x_1$ ")
2387     fig.colorbar(im00, ax=axes[0, 0], shrink=0.85)
2388
2389     im01 = axes[0, 1].imshow(
2390         exact_real, origin="lower", aspect="auto", extent=extent,
2391         cmap="RdBu_r", vmin=-amp_max, vmax=amp_max
2392     )
2393     axes[0, 1].set_title("Analytical: Re  $\Psi$ ")
2394     fig.colorbar(im01, ax=axes[0, 1], shrink=0.85)
2395
2396     im10 = axes[1, 0].imshow(
2397         num_prob, origin="lower", aspect="auto", extent=extent,
2398         cmap="viridis", vmin=0.0, vmax=prob_max
2399     )
2400     axes[1, 0].set_title(r"QTT:  $|\Psi|^2$ ")
2401     axes[1, 0].set_xlabel(r" $x_2$ ")
2402     axes[1, 0].set_ylabel(r" $x_1$ ")
2403     fig.colorbar(im10, ax=axes[1, 0], shrink=0.85)
2404
2405     im11 = axes[1, 1].imshow(
2406         exact_prob, origin="lower", aspect="auto", extent=extent,
2407         cmap="viridis", vmin=0.0, vmax=prob_max
2408     )
2409     axes[1, 1].set_title(r"Analytical:  $|\Psi|^2$ ")
2410     axes[1, 1].set_xlabel(r" $x_2$ ")
2411     fig.colorbar(im11, ax=axes[1, 1], shrink=0.85)
2412
2413     fig.suptitle(
2414         f"state rank {i}, n={label} | "
2415         f"E_exact={E_exact:.10f}, E_QTT={E_num:.10f}, "
2416         f"|dE|={E_err:.2e}, fidelity={fidelity:.10f}"
2417     )
2418     plt.show()
2419 else:
2420     print(
2421         "Direct QTT/analytical 2D heat maps require N_OSC == 2 and "
2422         "COUPLING_MODEL == 'none'. For coupled or higher-dimensional systems, "

```

```
2423         "use marginals and product-basis projections instead."  
2424     )
```