

1. Motivation

Many branches of science deal with complicated functions of many variables (e.g. classical or quantum field theories...).

These functions often have structure of many different length scales.

(i) Long scales: demand large domain of definition.

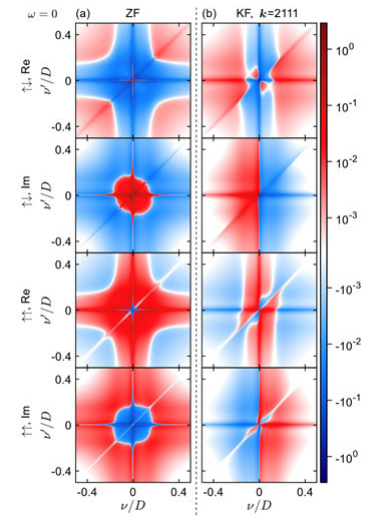
(ii) Short length scales demand dense grid.

(i) and (ii) together imply: numerics requires exploding costs in memory and computation time = 'curse of dimensionality'

Can explosion of costs be avoided without loss of control of accuracy?

Yes, if functions are 'compressible'!

4-point vertex functions of single-impurity Anderson model



Famous example: quantum many-body wavefunctions.

They can always be expressed as tensors:

'degree of tensor' = number of legs = ℓ

'local index' = $\sigma_\ell = 1, \dots, d$

'local dimension' = d

Total number of elements of F_σ :

$$d^\ell \Rightarrow F_\sigma \in \mathbb{C}^{d^\ell}$$

exponentially large in $\ell \Rightarrow$ 'curse of dimensionality'

$$|\Psi\rangle = \sum_{\sigma} F_{\sigma_1 \dots \sigma_\ell} |\sigma_1\rangle \otimes |\sigma_2\rangle \otimes \dots \otimes |\sigma_\ell\rangle$$

$$F_\sigma = \overbrace{\sigma_1 \sigma_2 \dots \sigma_\ell}^{\sigma = (\sigma_1, \dots, \sigma_\ell)}$$

Any tensor can be 'decomposed' or 'unfolded' into a 'matrix product state' (MPS) or 'tensor train' (TT):

$$F_\sigma \approx \tilde{F}_\sigma = \prod_{\ell=1}^L M_\ell^{\sigma_\ell} = [M_1]_{1a_1}^{\sigma_1} [M_2]_{a_1 a_2}^{\sigma_2} \dots [M_L]_{a_{L-1} 1}^{\sigma_L}, \quad (1)$$

$$F \approx \begin{array}{c} \overbrace{\sigma_1 \sigma_2 \dots \sigma_\ell \dots \sigma_L}^F \\ \approx \begin{array}{c} 1 \quad M_1 \quad \chi_1 \quad M_2 \quad \chi_2 \quad \dots \quad M_\ell \quad \chi_\ell \quad \dots \quad M_L \quad 1 \\ a_0 \quad \sigma_1 \quad a_1 \quad \sigma_2 \quad a_2 \quad \dots \quad \sigma_\ell \quad a_\ell \quad \dots \quad \sigma_L \quad a_L \end{array} \end{array} \quad (2)$$

$$\begin{array}{c} M_\ell \quad \chi_\ell \quad M_{\ell+1} \\ a_{\ell-1} \quad \sigma_\ell \quad a_\ell \quad \sigma_{\ell+1} \quad a_{\ell+1} \end{array} = [M_\ell]_{a_{\ell-1} a_\ell}^{\sigma_\ell} [M_{\ell+1}]_{a_\ell a_{\ell+1}}^{\sigma_{\ell+1}}$$

implicit Einstein summation over repeated indices: $\sum_{a_\ell=1}^{\chi_\ell}$

'bond index' = $a_\ell = 1, \dots, \chi_\ell$

'bond dimension' = χ_ℓ

(3)

Standard tool for unfolding tensors via repeated singular value decomposition (SVD):

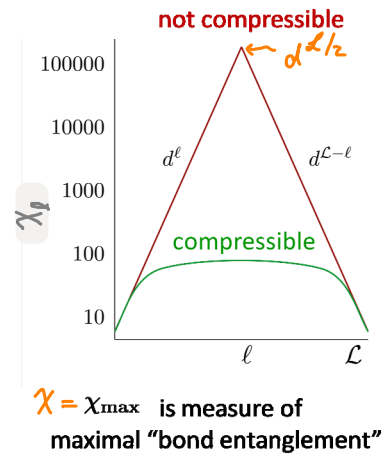
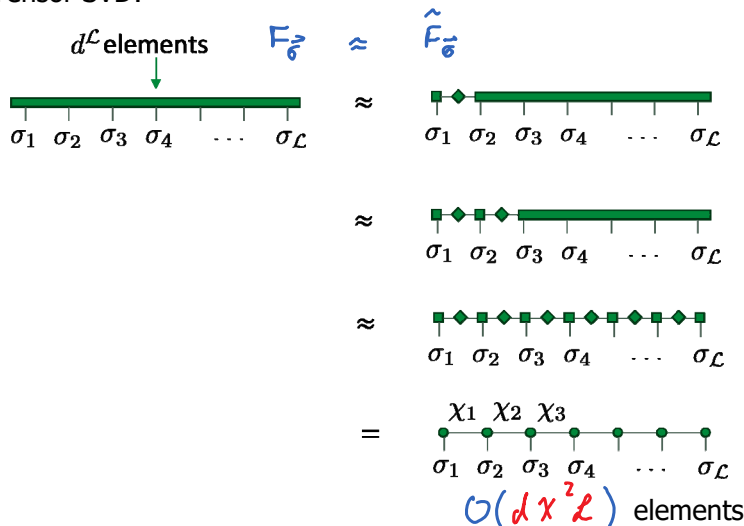
Matrix SVD: $M = U S V^\dagger$

$$\begin{pmatrix} \text{grid of dots} \end{pmatrix} \approx \begin{pmatrix} \text{grid of dots} \end{pmatrix} \begin{pmatrix} \text{grid of dots} \end{pmatrix} \begin{pmatrix} \text{grid of dots} \end{pmatrix}$$

or $\square \approx \square \diamond \square$

$\chi \leftarrow$ rank of approximation

Tensor SVD:



Tensor SVD is

'optimal': minimizes $\|F_\sigma - \tilde{F}_\sigma\|$ in Frobenius norm.

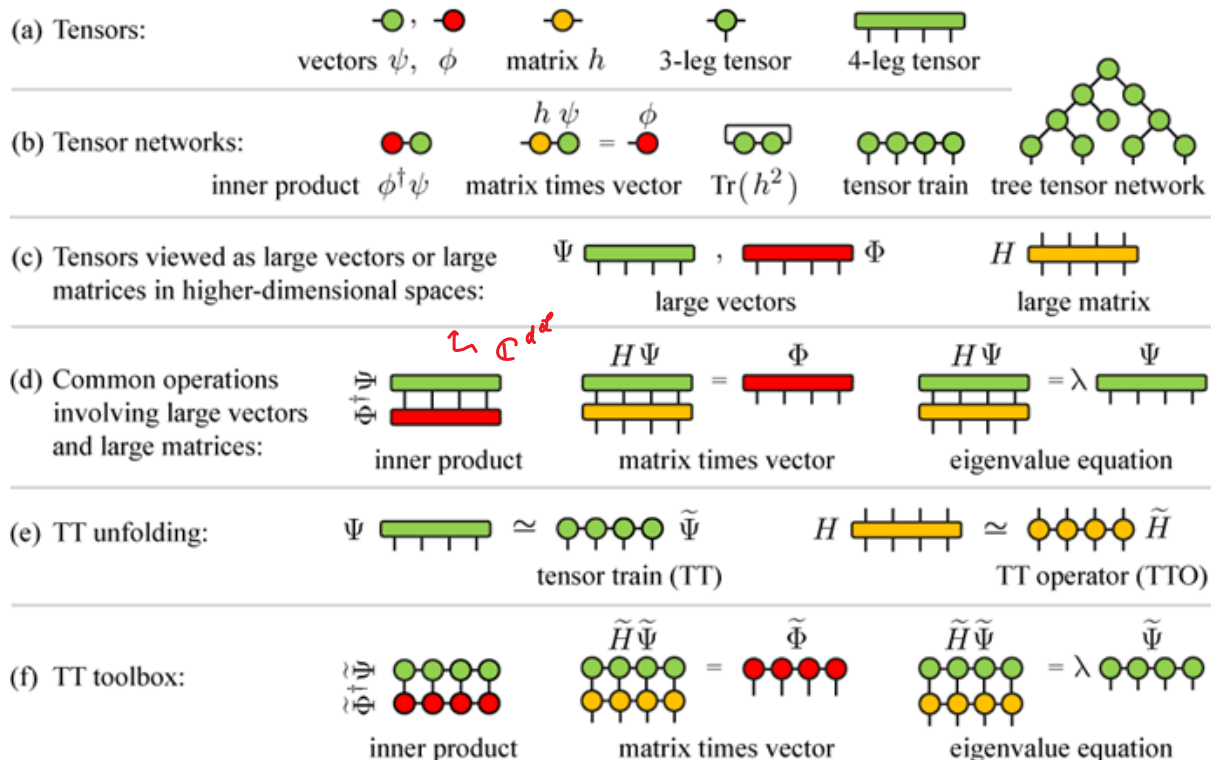
'rank-revealing': reveals correct rank (bond dimension) of every bond.

'expensive': needs all d^L elements of F_σ (else decomposition could not be optimal)!

F_σ is strongly compressible if it can be approximated by TT $\tilde{F}_\sigma$ with low bond dimension: $\chi_{\max} \ll d^{L/2}$

Then the number of elements, $O(d \chi^2 L) \ll d^L$, does not increase exponentially with L .

Miraculous fact: ground-state wave-functions of 1-dimensional Hamiltonians with short-ranged interactions are compressible! This led to development of DMRG and a powerful 'standard toolbox' for manipulating TTs (and more general tensor networks, e.g. tree tensor networks, ...)



Recent (since ~2010) insight: compressible tensors arise not only in quantum many-body theory.

Any discretized function can be expressed as a tensor. Often, that tensor is compressible!

If so, standard function manipulations (addition, multiplication, integration, convolution, Fourier transform) can be performed using tensor network toolbox, with exponential savings in numerical cost (memory & time)!

Bottleneck: standard strategy for unfolding tensor to a TT uses SVD, which requires knowledge of all d^ℓ components of tensor. So, SVD-based unfolding is exponentially costly.

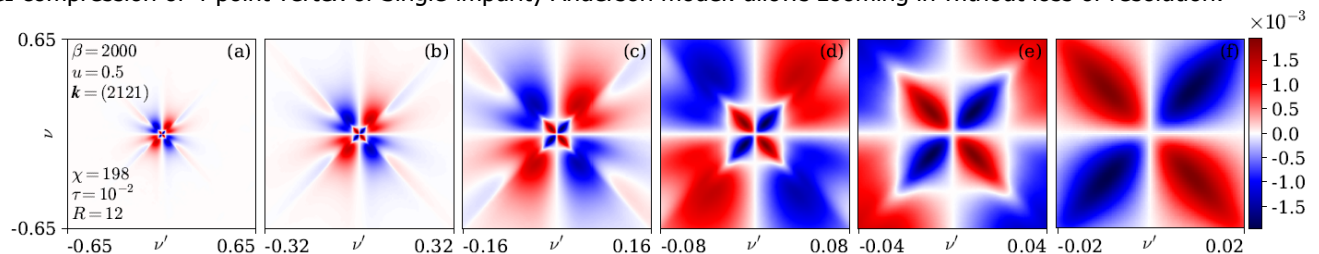
Recent progress: novel unfolding strategy based on tensor cross interpolation (TCI).

If tensor is compressible, TCI finds TT representations at runtime costs $\mathcal{O}(d\chi^3\ell)$
(If tensor is not compressible, TCI signals this by converging very slowly or not at all.)

TCI 'learns' compressed representation by interrogating only tiny subset of all elements of tensor.

This is akin to active machine learning, using a structured model (namely the TT) to represent the data.

TCI compression of 4-point vertex of Single-impurity Anderson model: allows zooming in without loss of resolution:



Historical remarks:

TCI pioneering papers (published in applied maths journals):

[Oseledets2010, Oseledets2011, Savostyanov2011, Savostyanov2014, Dolgov2020]

Early applications in physics (and beyond):

[Sozykin2022] Find minima of functions

[Fernandez2023] efficient (sign-problem-free) alternative to Monte Carlo sampling for calculating high-dimensional integrals arising in Feynman diagrams for quantum many-body systems

[Ritter2023] Compute topological invariants

[Jolly2023] Compute overlaps between atomic orbitals, solve Schrödinger equation of H_2^+ ion

[Sakurai2024] In mathematical finance, to speed up Fourier-transform-based option pricing

Open-source toolbox:

<https://tensor4all.org/>



Consider integral over function of d variables: $I = \int_U d\vec{x} f(\vec{x})$ (1)

$$f : U = D^d \subset \mathbb{R}^d \rightarrow \mathbb{R}, \quad \vec{x} = (x_1, \dots, x_d) \mapsto f(\vec{x}) \quad (2)$$

(for simplicity, same integration domain for each variable)

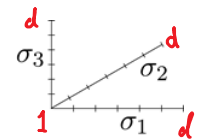
Standard approach: Monte-Carlo sampling of integrand.

Well-known issue: sign problem -- very slow convergence if integrand has many sign changes.

TT approach: discretize function, represent it as tensor, unfold it to TT: $f(\vec{x}) = F_{\vec{\sigma}} \approx \tilde{F}_{\vec{\sigma}} = \prod_{\ell} M_{\ell}$
This yields d independent integrals that can be computed separately. (3)

Discretization: (assume same grid for all variables)

Represent each variable on grid of d distinct points: $x_{\ell} \rightarrow \{p_1, \dots, p_{\sigma_{\ell}}, \dots, p_d\}$
enumerated by an index $\sigma_{\ell} = 1, \dots, d$ (4)



'Natural tensor representation' of function: $F_{\vec{\sigma}} = f(p_{\sigma_1}, \dots, p_{\sigma_d}) =$ (5)

Unfold this tensor into a TT [using SVD (expensive) or TCI (much cheaper)]:

will be explained in detail later!

$$f(x_1, \dots, x_d) \approx \tilde{F}_{\vec{\sigma}} = \begin{array}{c} M_1 \quad M_d \quad M_d \\ \downarrow \quad \downarrow \quad \downarrow \\ x_1 \quad x_d \quad x_d \end{array} = M_1(x_1) \dots M_d(x_d) \dots M_d(x_d) \quad (6)$$

$$\text{shorthand for } x_{\ell}(\sigma_{\ell}) \rightsquigarrow x_{\ell} = M_1^{\sigma_1} \dots M_d^{\sigma_d} \dots M_d^{\sigma_d} \quad (7)$$

When $\tilde{F}$ is low rank ($\chi_{max} =$ small), f is almost separable' (it would be separable if $\chi = 1$)

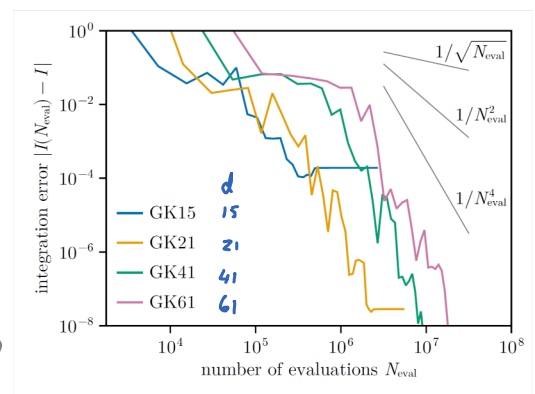
Integral separates into d one-dimensional integrals, which in practice are performed as discrete sums, followed by sequence of matrix-vector multiplications:

$$I = \int_U d\vec{x} f(\vec{x}) \approx \int_D d x_1 M_1(x_1) \dots \int_D d x_d M_d(x_d) \dots \int_D d x_d M_d(x_d) \quad (8)$$

$$= \left(\sum_{\sigma_1} M_1^{\sigma_1} \right) \cdot \left(\sum_{\sigma_d} M_d^{\sigma_d} \right) \cdot \left(\sum_{\sigma_d} M_d^{\sigma_d} \right) \quad (9)$$

Example: 10-dimensional integral of rapidly oscillating function:

$$I = 10^3 \int_{[-1,1]^{10}} d^{10}x \cos\left(10 \sum_{\ell=1}^{10} x_{\ell}^2\right) \exp\left[-10^{-3} \left(\sum_{\ell=1}^{10} x_{\ell}\right)^4\right] \quad (10)$$



Convergence with number of function evaluations: TCI: $\sim N_{eval}^{-4}$ (fast!) Monte Carlo: $\sim N_{eval}^{-1/2}$ (slow!)

TCI unfolding gives huge speed-up for multi-dimensional integration (without sign problem)!

Consider $m \times n$ matrix $A = (\vec{a}_1, \vec{a}_2, \dots, \vec{a}_n) = \begin{pmatrix} \text{grid} \end{pmatrix}$ (1)

with column vectors $\vec{a}_j \in \mathbb{C}^m$

'range of A' = $\text{Span} \{ \vec{a}_1, \vec{a}_2, \dots, \vec{a}_n \}$ (2)

'rank of A' = dimension of range of A = number of linearly independent columns (3)

= number of linearly independent columns

If A has rank = χ , then every column can be represented

as linear combination of subset of χ columns, $\vec{a}_j = \sum_{i=1}^{\chi} \vec{b}_i c_{ij}$ (4)

Let $\vec{B}$ be submatrix of $\vec{A}$ containing those columns: $\vec{B} = (\vec{b}_1, \dots, \vec{b}_{\chi})$ (5)

Matrix elements of A: $A_{kj} = [\vec{a}_j]_k \stackrel{(4)}{=} \sum_{i=1}^{\chi} [\vec{b}_i]_k c_{ij} = \sum_{i=1}^{\chi} B_{ki} c_{ij}$ (6)

So, A can be factorized: $A = \begin{matrix} m \times n \\ B \end{matrix} \cdot \begin{matrix} m \times \chi \\ C \end{matrix} = \begin{matrix} m \\ \chi \end{matrix} \cdot \begin{matrix} \chi \times n \end{matrix}$ (7)

This is compressed representation:

Memory cost: $m \cdot n$ elements for A, $(m + n)\chi$ for B and C.

SVD reveals rank: $A = U S V^T : \begin{pmatrix} \text{grid} \end{pmatrix} = \begin{pmatrix} \text{grid} \end{pmatrix} \begin{pmatrix} * & \cdot \end{pmatrix} \begin{pmatrix} \text{grid} \end{pmatrix}$ (8)

but computing it requires entire matrix A

Alternative low-rank approximation techniques exist that require only subset of rows and columns of matrix:

(i) matrix cross interpolation (CI), and (ii) partial rank-revealing LU decomposition (prrLU)

this section next section

Matrix cross interpolation (CI)

Let A be $m \times n$ matrix of rank χ with elements A_{ij} . $\begin{pmatrix} \text{grid} \end{pmatrix}$ (9)

$i \in \mathbb{I} = \{1, \dots, m\}$ list of all row indices, $j \in \mathbb{J} = \{1, \dots, n\}$ list of column indices (10)

$\mathcal{I} = \{i_1, \dots, i_{\hat{\chi}}\} \subset \mathbb{I}$ sublist of row indices, $\mathcal{J} = \{j_1, \dots, j_{\hat{\chi}}\} \subset \mathbb{J}$ sublist of column indices (11)

$A[\mathcal{I}, \mathcal{J}]$ = 'slice' of A containing all intersections of $\mathcal{I}$ -rows and $\mathcal{J}$ -columns: $\begin{pmatrix} \text{grid} \end{pmatrix}$ (12)

$A[\mathcal{I}, \mathcal{J}]_{\alpha\beta} = A_{i_{\alpha}, j_{\beta}} \quad \forall \alpha, \beta \in \{1, \dots, \hat{\chi}\}$ In particular: $A[\mathbb{I}, \mathbb{J}] = A$ (13)

Assume $\hat{\chi} \leq \chi$, with $\mathcal{I}, \mathcal{J}$ chosen such that $A[\mathcal{I}, \mathcal{J}]$ is non-singular: $\det(A[\mathcal{I}, \mathcal{J}]) \neq 0$ (14)

Define: $P = A[\mathcal{I}, \mathcal{J}] = \begin{pmatrix} \circ & \circ & \circ \\ \circ & \circ & \circ \\ \circ & \circ & \circ \end{pmatrix}$,
 'pivot matrix', its elements are called 'pivots'
 $\det P \neq 0 \Rightarrow P^{-1}$ exists (15)

$C = A[\mathcal{I}, \mathcal{J}] = \begin{pmatrix} \circ & \circ & \circ & \circ & \circ \\ \circ & \circ & \circ & \circ & \circ \\ \circ & \circ & \circ & \circ & \circ \end{pmatrix}$,
 gathers all columns containing pivots, i.e. all 'pivot columns' (15)

$R = A[\mathcal{I}, \mathcal{J}] = \begin{pmatrix} \circ & \circ & \circ & \circ & \circ \\ \circ & \circ & \circ & \circ & \circ \\ \circ & \circ & \circ & \circ & \circ \end{pmatrix}$,
 gathers all rows containing pivots, i.e. all 'pivot rows' (15)

CI formula:

gives rank- $\tilde{\chi}$ approximation $\tilde{A}$ of A

external indices i, j are fixed,

internal bonds represent sums

over pivot lists: $\sum_{i' \in \mathcal{I}} \sum_{j' \in \mathcal{J}}$

$$A \approx C \cdot P^{-1} \cdot R = \tilde{A} \quad (16)$$

$$A[\mathcal{I}, \mathcal{J}] \approx A[\mathcal{I}, \mathcal{J}] \cdot P^{-1} \cdot A[\mathcal{I}, \mathcal{J}] \quad (17)$$

$$A_{ij} = \sum_{i' \in \mathcal{I}} \sum_{j' \in \mathcal{J}} A_{i'j'} \approx \sum_{i' \in \mathcal{I}} \sum_{j' \in \mathcal{J}} A_{i'j'} \cdot P^{-1}_{i'j'} \cdot A_{i'j'} \quad (18)$$

$$\begin{pmatrix} \circ & \circ & \circ & \circ & \circ \\ \circ & \circ & \circ & \circ & \circ \\ \circ & \circ & \circ & \circ & \circ \end{pmatrix} \approx \begin{pmatrix} \circ & \circ & \circ & \circ & \circ \\ \circ & \circ & \circ & \circ & \circ \\ \circ & \circ & \circ & \circ & \circ \end{pmatrix}^{-1} \begin{pmatrix} \circ & \circ & \circ & \circ & \circ \\ \circ & \circ & \circ & \circ & \circ \\ \circ & \circ & \circ & \circ & \circ \end{pmatrix} \quad (19)$$

Remark: C and R contain original columns and rows of A itself, not linear combinations of those. By contrast, for SVD, U and V^+ contain linear combinations of rows and columns of A .

Properties of CI formula:

(i) For $\hat{\chi} \leq \chi$, $\tilde{A}$ yields an 'interpolation', i.e. it exactly reproduces all $\mathcal{I}$ -rows and all $\mathcal{J}$ -columns of A ,
 i.e. 'all rows and columns from which $\tilde{A}$ was built'. (21)

Proof of (i): consider only pivot rows ($\mathcal{I}$ -rows) or pivot columns ($\mathcal{J}$ -columns):

$$\tilde{A}[\mathcal{I}, \mathcal{J}] = A[\mathcal{I}, \mathcal{J}] \cdot P^{-1} \cdot A[\mathcal{I}, \mathcal{J}] = A[\mathcal{I}, \mathcal{J}] \quad (22)$$

$$\tilde{A}[\mathcal{I}, \mathcal{J}] = A[\mathcal{I}, \mathcal{J}] \cdot P^{-1} \cdot A[\mathcal{I}, \mathcal{J}] = A[\mathcal{I}, \mathcal{J}] \quad (23)$$

In general, $\tilde{A}$ has no information about the 'other' (non-pivot) rows and columns of A . If these are unrelated to the pivot rows and columns (e.g. because A is a random matrix), CI formula will get them completely wrong. However, if the non-pivot rows and columns are somewhat / very / fully related to the pivot rows and columns, i.e. if some / most / all of them can be expressed as linear combination of pivot rows and columns, the CI formula will reproduce them somewhat well / very well / exactly.

(ii) For $\hat{\chi} = \chi$, $\tilde{A}$ exactly reproduces the entire matrix, $\tilde{A} = A$ (20)

'Plausibility argument': for $\hat{\chi} = \chi$, all columns of A can be expressed as linear combinations of a subset of $\hat{\chi}$ columns. Choose these to be pivot columns. Since CI formula reproduces all pivot columns exactly, it is plausible that it reproduces all columns of A exactly.

Caveat: does the CI for the non-pivot columns yield the correct linear combinations of pivot columns? Answer is yes. Proof of (ii) follows further below.

(iii) Error of CI formula is determined by $[A/P]$, the Schur complement of P in A (21)

Definition: for block matrix: $A = \begin{pmatrix} A_{11} & A_{12} \\ A_{21} & A_{22} \end{pmatrix} = \begin{array}{|c|c|} \hline \text{ } & \text{ } \\ \hline \end{array}$ with A_{11} assumed square and invertible,

the 'Schur complement of A_{11} in A is: $[A/A_{11}] \equiv A_{22} - A_{21}(A_{11})^{-1}A_{12}$ (22)

block dimensions match: $\square - \square = \square$ note index structure: 1 meets 1

Schur determinant identity: $\det A = \det A_{11} \det[A/A_{11}]$ proof: see (TCI.5.3-5) (23)

Proof of (iii): consider arbitrary matrix A and pivot matrix P .

Permute rows and columns of A such that all pivots lie in first $\hat{\chi}$ rows and columns, labeled

$\mathcal{I}_1 = \mathcal{J}_1 = \{1, \dots, \hat{\chi}\}$. Denote labels of remaining rows and columns by $\mathcal{I}_2 = \mathcal{I} \setminus \mathcal{I}_1$, $\mathcal{J}_2 = \mathcal{J} \setminus \mathcal{J}_1$

Permuted matrix:

$$A(\mathcal{I}, \mathcal{J}) = \begin{pmatrix} A(\mathcal{I}_1, \mathcal{J}_1) & A(\mathcal{I}_1, \mathcal{J}_2) \\ A(\mathcal{I}_2, \mathcal{J}_1) & A(\mathcal{I}_2, \mathcal{J}_2) \end{pmatrix} = \begin{pmatrix} A_{11} & A_{12} \\ A_{21} & A_{22} \end{pmatrix}, \quad P = A_{11} = A(\mathcal{I}_1, \mathcal{J}_1) \quad (24)$$

$$\text{CI formula: } \tilde{A} = \begin{pmatrix} A_{11} \\ A_{21} \end{pmatrix} (A_{11})^{-1} \begin{pmatrix} A_{11} & A_{12} \end{pmatrix} = \begin{pmatrix} A_{11} & A_{12} \\ A_{21} & A_{21}(A_{11})^{-1}A_{12} \end{pmatrix} \quad (25)$$

$$\text{CI error: } A - \tilde{A} = \begin{pmatrix} 0 & 0 \\ 0 & [A/A_{11}] \end{pmatrix} = [A/P] \quad (26)$$

Error depends on inverse of pivot matrix. So, strategy for reducing error is to choose pivots such that $|\det P|$ is maximal. This is known as 'maximum volume principle' (more rigorous argument: see below).

Finding 'best pivots', i.e. ones that satisfy maximum volume principle, is in general exponentially difficult; but good heuristics exist that get close to optimum in practice (see below).

Proof of (ii): (rigorous version of plausibility argument given above)

For $\hat{\chi} = \chi$, pivot matrix $P = A[\mathcal{I}, \mathcal{J}]$ has dimension $\chi \times \chi$

Pick arbitrary 'new' element (x_0, y_0) of A that does not sit on a pivot row or pivot column.

Consider $(\chi+1) \times (\chi+1)$

submatrix of A , built from pivot matrix and new element:

$$A' = \begin{pmatrix} A'_{11} & A'_{12} \\ A'_{21} & A'_{22} \end{pmatrix} = \begin{pmatrix} A(\mathcal{I}, \mathcal{J}) & A(\mathcal{I}, y_0) \\ A(x_0, \mathcal{J}) & A(x_0, y_0) \end{pmatrix} = \begin{array}{|c|c|} \hline P_{old} & \text{ } \\ \hline \text{ } & \text{ } \\ \hline \end{array} \quad (27)$$

$$\text{Recall Schur determinant identity: } \det A' = \det A'_{11} \det[A/A'_{11}] \quad (28)$$

Applied to A' :

$$\det A' = \det[A(\mathcal{I}, \mathcal{J})] \det[A(x_0, y_0) - A(x_0, \mathcal{J})A^{-1}(\mathcal{I}, \mathcal{J})A(\mathcal{I}, y_0)]$$

Applied to A' :

$$\det A' = \underbrace{\det[A(\mathcal{I}, \mathcal{J})]}_{\neq 0} \det[A(x_0, y_0) - A(x_0, \mathcal{J})A^{-1}(\mathcal{I}, \mathcal{J})A(\mathcal{I}, y_0)] \quad (24)$$

A' has $X+1$ columns, hence they are linearly dependent

pivot matrix is non-singular

Therefore:

$$A(x_0, y_0) = A(x_0, \mathcal{J})A^{-1}(\mathcal{I}, \mathcal{J})A(\mathcal{I}, y_0) \stackrel{(17)}{=} \hat{A}(x_0, y_0) \quad (30)$$

Thus, CI formula exactly reproduces any 'new' element that does not sit on pivot row or pivot column.