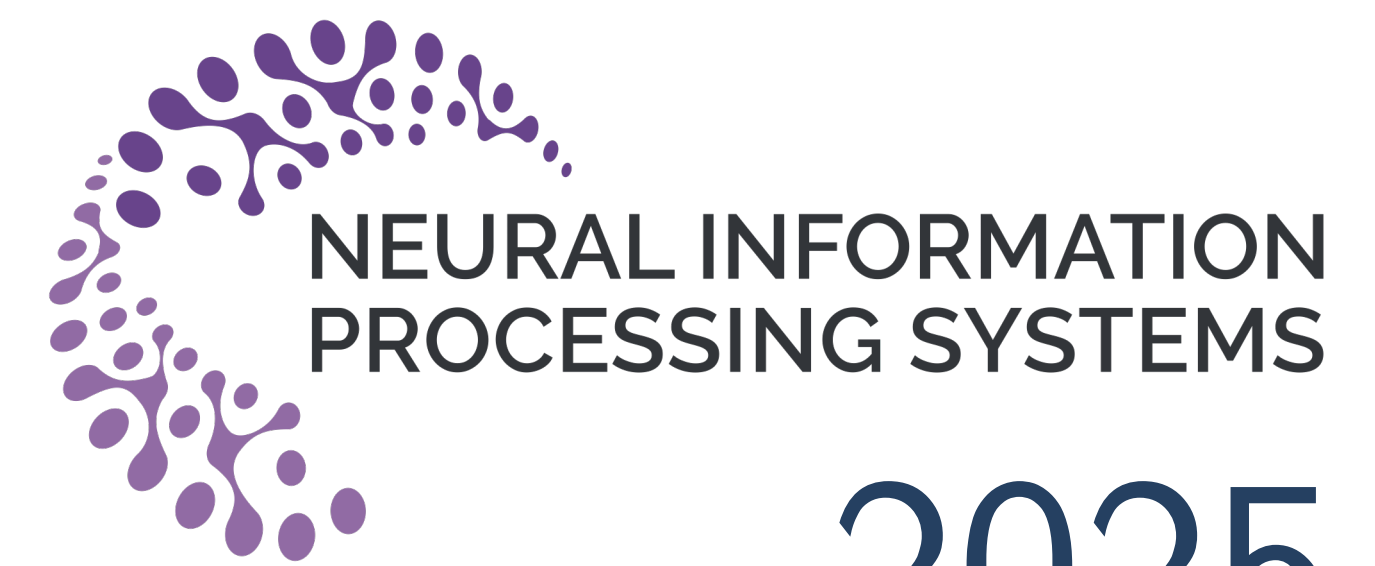


Revisiting **Orbital Minimization Method** for Neural Operator Decomposition

Jongha (Jon) Ryu
Samuel Zhou
Gregory W. Wornell

MIT EECS



2025

Problem: Operator Spectral Decomposition

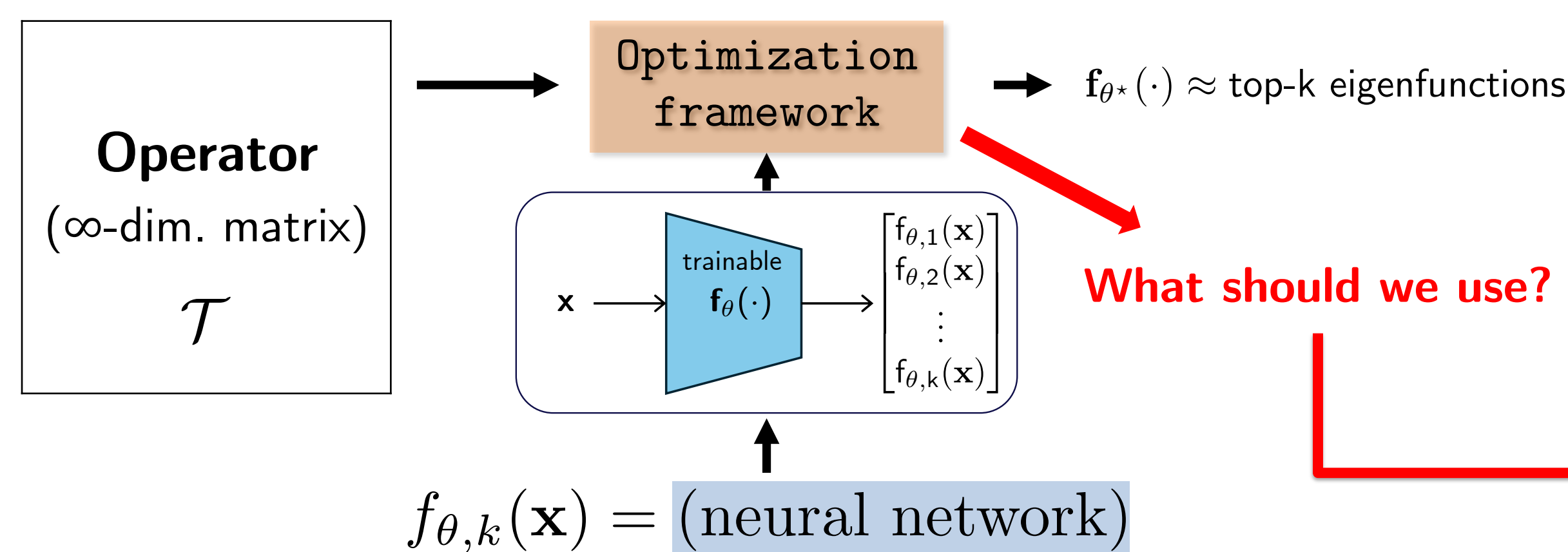
Given a target positive semidefinite operator
find the top- K eigenvalues/functions

$$\mathcal{T} = \sum_{i=1}^{\infty} \lambda_i |\psi_i\rangle \langle \psi_i| \xrightarrow{?} \mathcal{T} |\psi_i\rangle = \lambda_i |\psi_i\rangle \quad i = 1, \dots, K$$

Applications (operator)

- Quantum chemistry (Hamiltonian)
- Representation learning (conditional expectation operator)
- Graph learning (graph Laplacians)
- ...

Parametric Spectral Decomposition

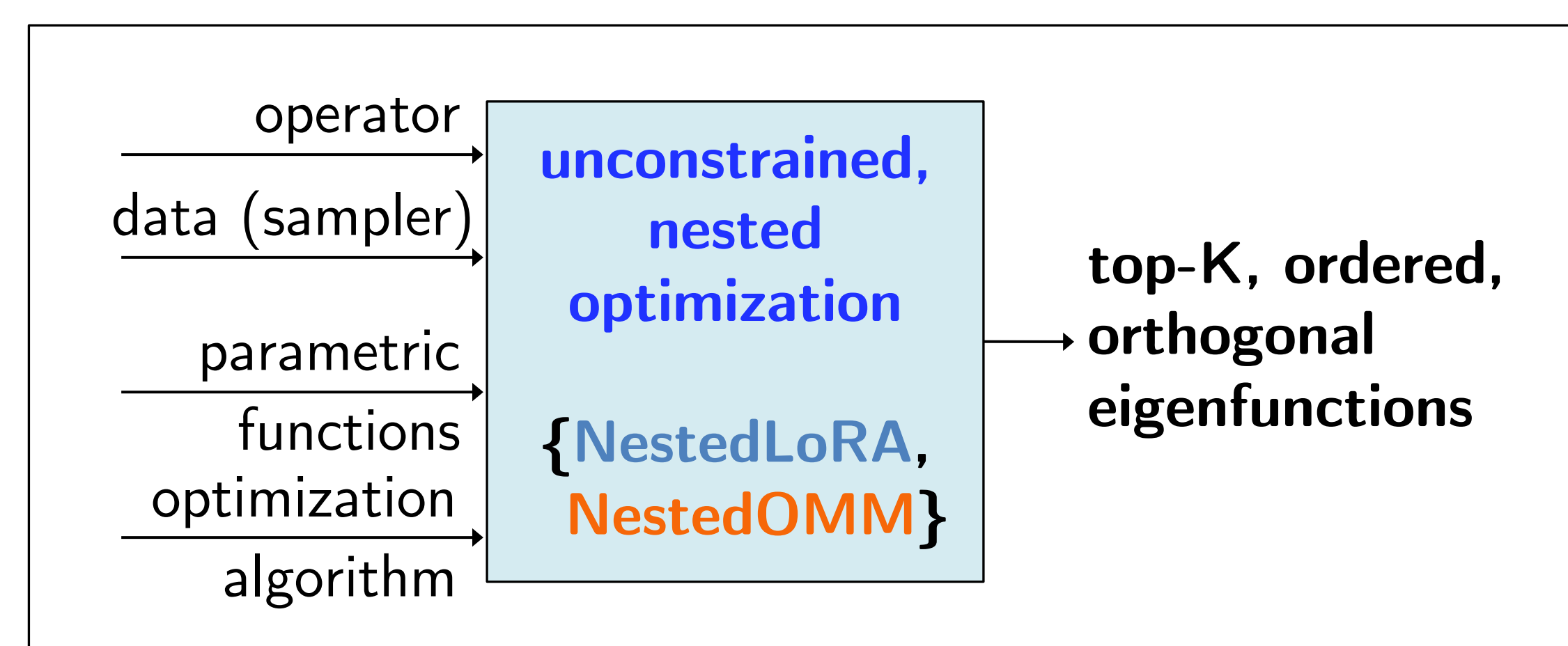


Classical approaches use very limited form of trial functions

- Finite element methods $f_{\theta,k}(\mathbf{x}) = \sum_{i=1}^N v_{k,i} \delta(\mathbf{x} - \mathbf{x}_i)$
- Galerkin type methods $f_{\theta,k}(\mathbf{x}) = \mathbf{a}_k^T \zeta_{1:M}(\mathbf{x})$

Using neural networks can be **much more scalable** with enough inductive bias (proven to be powerful in *quantum chemistry*)

Our Approach: Nested Optimization with Unconstrained Variational Principle



- Compact operators: **NestedLoRA** (Ryu et al., ICML 2024)
- PSD operators: **NestedOMM** (this paper!)

What's **Orbital Minimization Method (OMM)**?

$$A = \sum_{i=1}^d \lambda_i \mathbf{w}_i \mathbf{w}_i^T, \quad \lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_d \geq 0$$

$$\max_{V \in \mathbb{R}^{d \times k} : V^T V = I_k} \text{tr}(V^T A V) \quad \text{standard Rayleigh quotient maximization}$$

$$\max_{V \in \mathbb{R}^{d \times k}} \text{tr}((V^T V)^{-1} V^T A V) \quad \text{unconstrained version}$$

$$\min_{V \in \mathbb{R}^{d \times k}} \mathcal{L}_{\text{omm}}(V), \quad \text{where } \mathcal{L}_{\text{omm}}(V) \triangleq -\text{tr}((2I_k - V^T V) V^T A V)$$

All characterize the top-k eigensubspace as a global optima
(OMM was proposed by Mauri et al., (1993) in density functional theory)

Q1. Where does it come from?

(Original: ad-hoc Neumann series expansion)
We give a **simple yet intuitive** derivation:

$$\mathcal{L}_{\text{omm}}(V) = \text{tr}((I_d - VV^T)^2 A) - \text{tr}(A)$$

$$\mathcal{L}_{\text{omm}}^{(p)}(V) \triangleq \text{tr}((I_d - VV^T)^{2p} A) - \text{tr}(A)$$

Q2. How can we show the consistency?

We provide an **elementary linear-algebraic proof**

Key observation: OMM only depends on the PSD matrix VV^T

NestedOMM: OMM with Nesting

We can apply the nesting technique (Ryu et al., ICML 2024) to learn ordered eigenvectors $\rightarrow$ **NestedOMM**

- **Joint nesting:** update $V_{1:k}$ by $\sum_{i=1}^k w_i \mathcal{L}_{\text{omm}}(V_{1:i})$
- **Sequential nesting:** update $\mathbf{v}_i$ using gradient of $\mathcal{L}_{\text{omm}}(V_{1:i})$, for each i

Special Case: Streaming PCA

$$\text{PCA: } A = \mathbb{E}_{p(\mathbf{x})}[\mathbf{x}\mathbf{x}^T] \quad / \quad \text{streaming PCA: } A_t = \frac{1}{B} \sum_{b=1}^B \mathbf{x}_b^{(t)} (\mathbf{x}_b^{(t)})^T$$

- **OMM with sequential nesting** (+ gradient descent)

$$\mathbf{v}_i^{(t+1)} \leftarrow \mathbf{v}_i^{(t)} + \eta_t \left\{ (I - V_{1:i}^{(t)} (V_{1:i}^{(t)})^T) A_t \mathbf{v}_i^{(t)} + A_t (I - V_{1:i}^{(t)} (V_{1:i}^{(t)})^T) \mathbf{v}_i^{(t)} \right\}$$

- cf. Generalized Hebbian algorithm (a.k.a. Sanger's algorithm)

$$\mathbf{v}_i^{(t+1)} \leftarrow \mathbf{v}_i^{(t)} + \eta_t (I - V_{1:i}^{(t)} (V_{1:i}^{(t)})^T) A_t \mathbf{v}_i^{(t)}$$

Comparison with **Low-Rank Approximation (LoRA)**

$$\mathcal{L}_{\text{omm}}(V) = \text{tr}((I_d - VV^T)^2 A) - \text{tr}(A)$$

$$V^*(V^*)^T = W_{1:k} W_{1:k}^T$$

Only applicable for EVD of PSD matrices
Can outperform LoRA when A is sparse

$$\mathcal{L}_{\text{lora}}(V) = \|A - VV^T\|_F^2 - \|A\|_F^2$$

$$V^*(V^*)^T = W_{1:k} \Lambda_{1:k} W_{1:k}^T$$

Applicable for SVD of any matrices

From Matrix to Operator

$$\mathcal{T}: L_\rho^2(\mathcal{X}) \rightarrow L_\rho^2(\mathcal{X}) \quad L_\rho^2(\mathcal{X}) \triangleq \left\{ f: \mathcal{X} \rightarrow \mathbb{R} \mid \int_{\mathcal{X}} f(\mathbf{x})^2 \rho(d\mathbf{x}) < \infty \right\}$$

$$\langle f_0, f_1 \rangle_\rho \triangleq \int_{\mathcal{X}} f_0(\mathbf{x}) f_1(\mathbf{x}) \rho(d\mathbf{x})$$

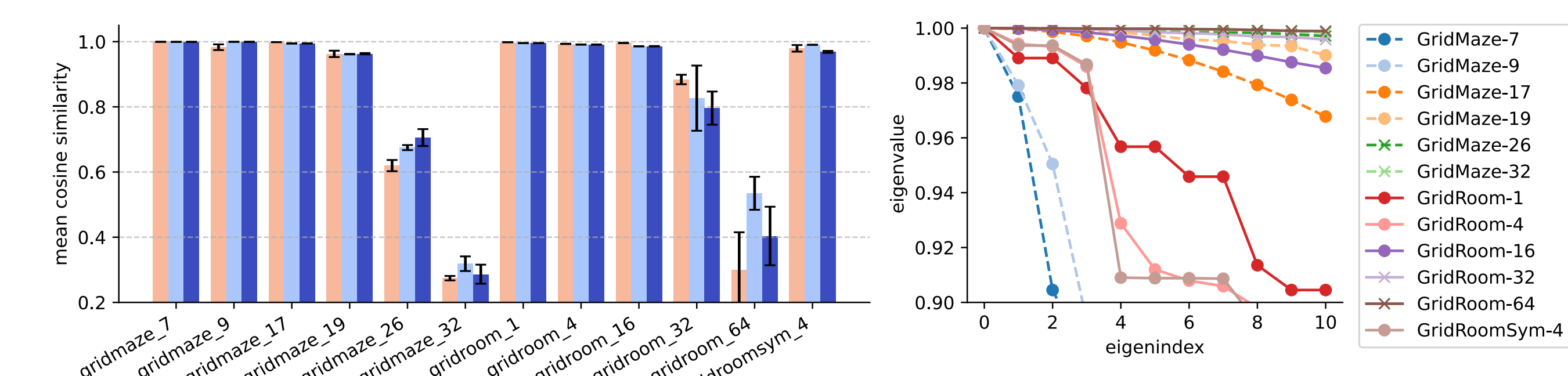
REPLACE

$$\text{overlap matrix } V^T V \quad \text{projected matrix } V^T A V \quad \text{w/} \quad M_\rho[\mathbf{f}] \triangleq \int \mathbf{f}(\mathbf{x}) \mathbf{f}(\mathbf{x})^T \rho(d\mathbf{x})$$

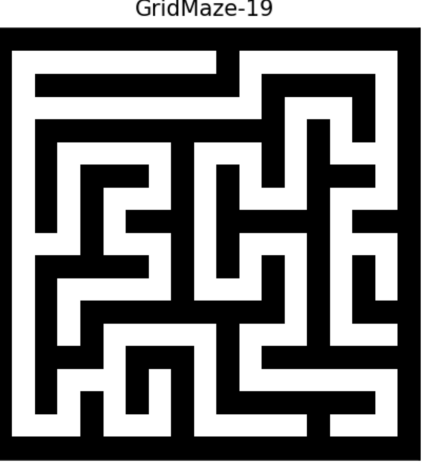
$$M_\rho[\mathbf{f}, \mathcal{T}\mathbf{f}] \triangleq \int \mathbf{f}(\mathbf{x}) (\mathcal{T}\mathbf{f})(\mathbf{x})^T \rho(d\mathbf{x})$$

Experiments

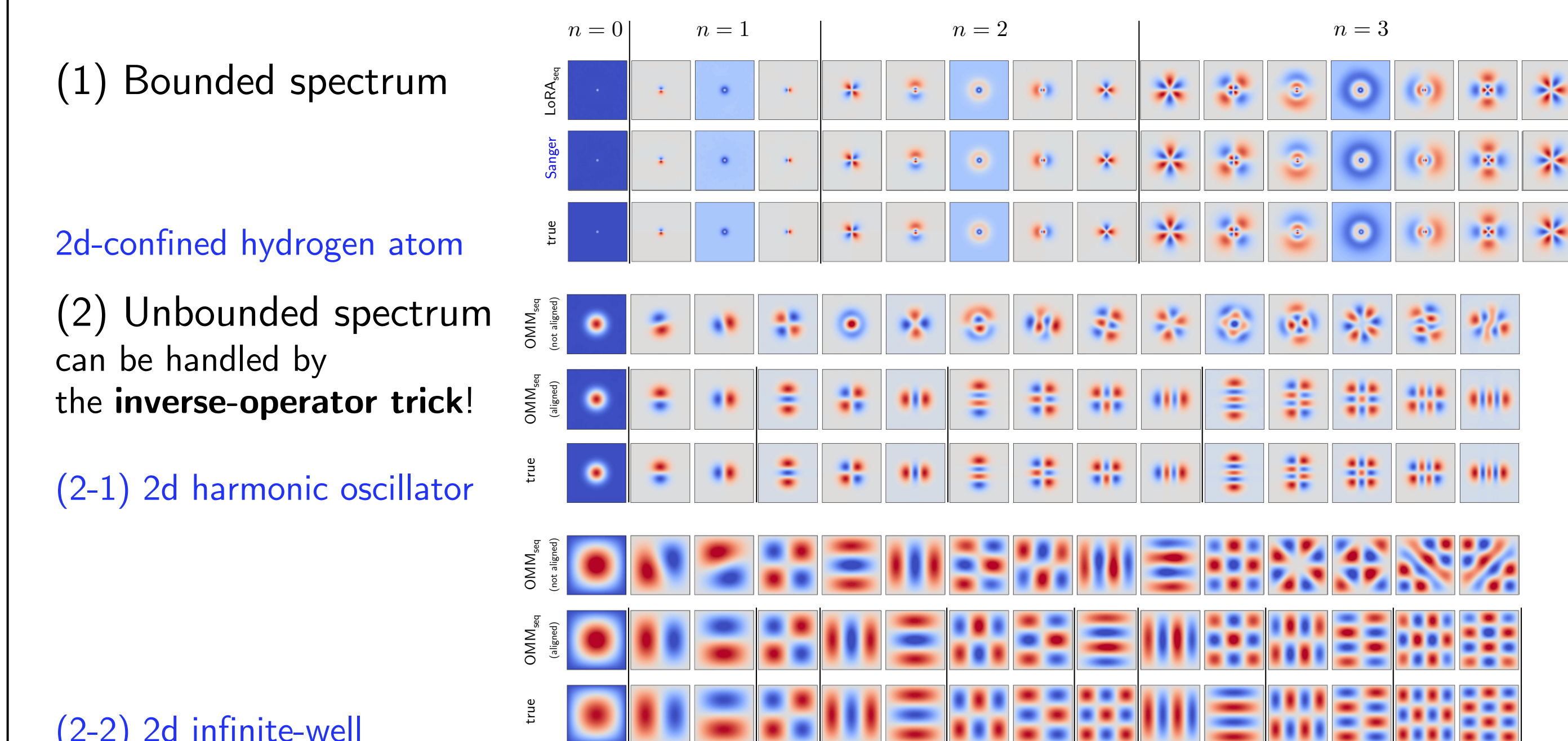
(1) Representation learning for reinforcement learning



Compared to ALLO, **OMM does not need any hyperparameter tuning and robust to small spectral gap!**



(2) Solving time-independent Schrödinger equation



(3) Linear-algebraic self-supervised representation learning

Method	With projector		DirectCLR (top-64 dim.)	
	Top-1	Top-5	Top-1	Top-5
OMM ($p = 1$)	60.02	87.13	59.83	86.65
OMM _{jnt} ($p = 1$)	61.30	85.22	59.92	85.13
OMM _{seq} ($p = 1$)	59.91	87.02	52.96	81.22
OMM ($p = 2$)	63.92	89.08	61.27	87.07
OMM ($p = 1$) + OMM ($p = 2$)	64.77	89.18	63.99	88.88
SimCLR [3]	66.50	89.28	N/A	N/A

Higher-order OMM helps learn better representation
Though not SOTA, motivates further investigation!