

Joint-Space Empowerment: Discovering Low-Dimensional Action Manifolds for Dexterous Manipulation

James Heald*, Vittorio Caggiano[†], Vikash Kumar[†], Maneesh Sahani*

*Gatsby Unit, University College London, London, UK

[†]MyoLab, New York, USA

Correspondence: jamesbheald@gmail.com

Abstract—Searching for effective policies is notoriously challenging in overactuated musculoskeletal systems, where individual joints are controlled by multiple muscles. Although this redundancy complicates naive policy search, it also implies that effective control can be captured by a low-dimensional action manifold. To identify such a manifold, we introduce Joint-Space Empowerment (JoSE), a novel information-theoretic objective that quantifies how much control an agent has over its mechanical degrees-of-freedom. We frame manifold discovery as an optimal precoding problem—where a state-dependent precoder maps low-dimensional latent actions to high-dimensional muscle commands—and derive an optimal precoder in closed form under control-affine Gaussian dynamics. We show that manipulation policies learned on this manifold display significantly enhanced dexterity, sample efficiency, and improved generalization. These results present optimal precoding as a general paradigm for coordinating high-dimensional actuators to control low-dimensional features, with implications for robotic manipulation.

I. INTRODUCTION

Replicating human-level dexterity is a grand challenge of robotics. One promising path to overcoming this challenge is to understand how the human brain and central nervous system resolve the redundancy inherent in the musculoskeletal system. In the human hand, 39 muscles control 23 mechanical degrees of freedom, creating an overactuated system in which many different muscle activation patterns can produce identical joint torques. It has long been hypothesized that the brain resolves this redundancy problem by controlling a small number of synergies, where each synergy recruits a group of muscles as a single functional unit, thereby constraining control to a low-dimensional manifold embedded within the high-dimensional muscle space [41, 14, 25]. Yet despite more than a century of study [35], the computational principles that underlie this low-dimensional control strategy remain poorly understood.

Understanding these principles has direct implications for robotic manipulation: reinforcement learning (RL), the dominant paradigm for learning manipulation, is particularly sensitive to action redundancy [15, 38, 46, 3, 36]. For value-based methods, treating behaviorally equivalent actions as distinct dilutes the learning signal across the action space; for policy-gradient methods, action redundancy inflates gradient variance [44]. Prior approaches define action redundancy in terms of the effects of actions on next states [37, 30, 1, 45]. However, for musculoskeletal systems, this approach assigns

equal importance to both muscle and joint states—ignoring the fact that tasks are accomplished through skeletal interactions with the environment, with muscles serving merely as the actuators (Figure 1 a).

To exploit this functional asymmetry, we develop an information-theoretic framework that defines action redundancy in terms of the effects of muscle commands on joint motion. Our framework eliminates redundancy by discovering a task-agnostic, state-dependent action manifold that gives an agent maximum control over its mechanical degrees-of-freedom.

A version of this framework achieved first place in the Manipulation Track of the NeurIPS 2024 MyoChallenge [42], where it was used to control a model of the entire human arm (27 degrees-of-freedom, 63 muscles).

II. RELATED WORK

Synergy learning methods can be split into three categories. **Demonstration-based** approaches extract synergies from observed movement, whether human behavior or the output of synthetic expert policies [26, 28, 17]. **Task-optimized** approaches derive synergies by optimizing a policy with respect to a task-specific objective function (e.g. cost-to-go or return) [40, 12, 43]. **Task-agnostic** approaches ground synergies in the biomechanics of the body, independent of task objectives [6, 33, 21].

III. JOINT-SPACE EMPOWERMENT (JOSE)

We model the musculoskeletal system as a controlled Markov process $(\mathcal{S}, \mathcal{A}, P)$, where $\mathcal{S}$ and $\mathcal{A}$ are continuous state and action spaces, and P is the state transition kernel. The state $s = (q, \dot{q}, m) \in \mathcal{S}$ consists of joint angles and velocities, q and $\dot{q}$, and muscle activations m . Each action dimension represents the neural input to a single muscle.

Our central goal is to discover a low-dimensional action manifold that maximizes how much control an agent has over its body. Rather than attempt to control the full high-dimensional state of the body, we target a low-dimensional *controlled variable* whose regulation is sufficient to govern the body’s mechanical degrees-of-freedom. In musculoskeletal systems, joint angular velocities serve this role: they are sufficient to determine body motion, while abstracting away actuator-level details such as muscle activations.

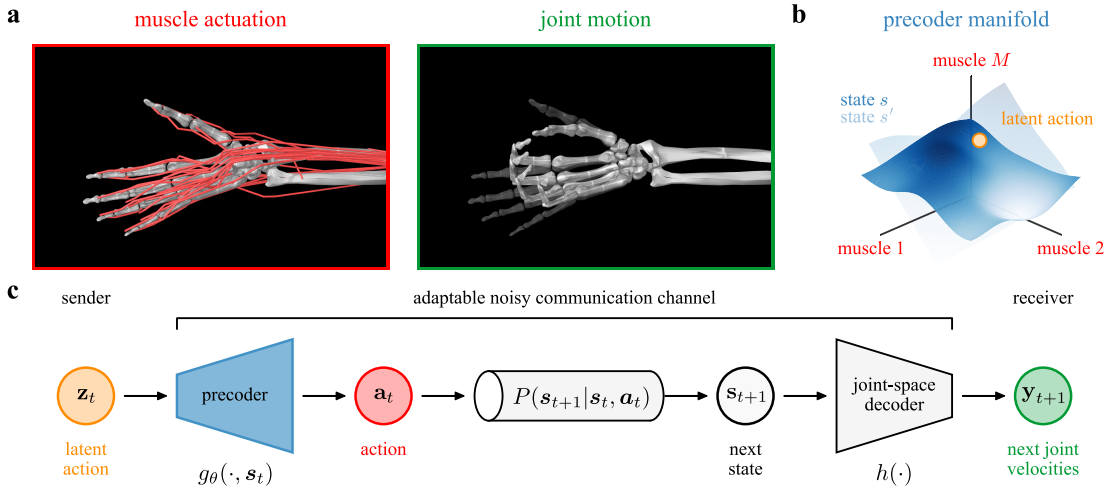


Fig. 1. **Joint-Space Empowerment (JoSE) enables dexterous motor coordination.** (a) Muscle actuation (left) produces joint motion (right) according to the dynamics of the musculoskeletal system. (b) A latent action (orange circle) maps to a point on a state-dependent low-dimensional action manifold (blue surface) parameterized by a precoder. (c) A noisy communication channel composed of a state-dependent precoder $g_\theta : \mathcal{Z} \times \mathcal{S} \rightarrow \mathcal{A}$, the state transition kernel P , and a decoder $h : \mathcal{S} \rightarrow \mathcal{Y}$. The precoder g_θ maps a low-dimensional latent action z_t to a high-dimensional action a_t , the environment transitions according to $s_{t+1} \sim P(s_{t+1}|s_t, a_t)$, and the resulting state s_{t+1} is mapped by the decoder to the joint angular velocities $y_{t+1} \in \mathcal{Y}$.

We formalize our goal as an *optimal precoding problem*. To this end, we introduce a latent action space $\mathcal{Z}$ and a state-dependent precoder $g_\theta : \mathcal{Z} \times \mathcal{S} \rightarrow \mathcal{A}$, parameterized by θ , that maps a latent action $z_t \in \mathcal{Z}$ to an action $a_t \in \mathcal{A}$. Actions produced by the precoder lie on a manifold $\mathcal{M}_\theta(s_t) \triangleq \{g_\theta(z_t, s_t) : z_t \in \mathcal{Z}\} \subseteq \mathcal{A}$ (Figure 1b). We also introduce a joint-space decoder $h : \mathcal{S} \rightarrow \mathcal{Y}$ that deterministically maps the full system state to the controlled variable $y_t \in \mathcal{Y}$, corresponding to the joint angular velocities.

We view action manifold discovery as a communication problem involving a sender, a noisy communication channel, and a receiver (Figure 1c). At time t in state s_t , the sender transmits a source signal z_t over the channel, which is received at the next time step as the value y_{t+1} . The transmitted signal z_t corresponds to a latent action, and the received signal y_{t+1} corresponds to the next joint velocities. The noisy communication channel $p_\theta(y_{t+1}|z_t, s_t)$ is defined by the composition of the precoder, transition kernel and decoder.

Given a source distribution $p(z_t|s_t)$ and a communication channel $p_\theta(y_{t+1}|z_t, s_t)$, the average information the received signal y_{t+1} contains about the transmitted signal z_t in state s_t is given by the mutual information:

$$\mathcal{I}_\theta(z_t; y_{t+1}|s_t) = \left\langle \log \frac{p_\theta(y_{t+1}, z_t|s_t)}{p(z_t|s_t)p_\theta(y_{t+1}|s_t)} \right\rangle_{p_\theta(y_{t+1}, z_t|s_t)},$$

where $\langle \cdot \rangle_p$ denotes expectation with respect to distribution p . The greater the information transmitted across the channel, the more control or influence an agent has over the joint velocities. The mutual information can be expressed as:

$$\mathcal{I}_\theta(z_t; y_{t+1}|s_t) = \mathcal{H}_\theta[y_{t+1}|s_t] - \mathcal{H}_\theta[y_{t+1}|z_t, s_t],$$

where $\mathcal{H}_\theta[y_{t+1}|s_t]$ is the marginal entropy of the next joint velocities, and $\mathcal{H}_\theta[y_{t+1}|z_t, s_t]$ is the conditional entropy of the next joint velocities given the latent action, which reflects how

noisy the channel is. Mutual information is high when a wide diversity of next joint velocities can be reached (high marginal entropy) in a reliable manner (low conditional entropy). We refer to the maximum mutual information over all source distributions as Joint-Space Empowerment (JoSE).

Definition III.1 (Joint-Space Empowerment). *Given a communication channel $p_\theta(y_{t+1}|z_t, s_t)$, the Joint-Space Empowerment (JoSE) in a state $s_t \in \mathcal{S}$ is defined as the maximum mutual information between latent actions and next joint velocities over all source distributions $p(z_t|s_t)$ (i.e., the channel capacity):*

$$\mathcal{E}_\theta(s_t) = \max_{p(z_t|s_t)} \mathcal{I}_\theta(z_t; y_{t+1}|s_t).$$

Since the communication channel $p_\theta(y_{t+1}|z_t, s_t)$ depends on the precoder g_θ , whose parameters θ are learnable, the channel is itself adaptable. Our goal is to learn the optimal precoder that maximizes JoSE:

$$\theta^* = \arg \max_{\theta} \mathcal{E}_\theta(s_t).$$

IV. MODEL-BASED MANIFOLD DISCOVERY

To develop a tractable and data-efficient solution to mutual information estimation, we adopt a model-based approach, exploiting the insight that the effects of muscle commands on joint motions can be well approximated by a state-dependent control-affine mapping (see Appendix D). Specifically, we model the mapping from states and muscle commands to next joint velocities as control-affine Gaussian:

$$p(y_{t+1}|s_t, a_t) = \mathcal{N}(f(s_t) + G(s_t)a_t, Q(s_t)). \quad (1)$$

Crucially, while the dependence on action is affine, the dynamics parameters $\phi \triangleq \{f, G, Q\}$ depend nonlinearly on the state s_t . We parameterize this relationship using a neural network ψ that maps the state s_t to the dynamics parameters

ϕ . We consider two variants for the input matrix $\mathbf{G}$: a full-rank parameterization and a low-rank factorization that reflects biomechanical coupling between joints (see Appendix E-B1).

For a system of the form Equation (1), an optimal precoder maximizing JoSE is linear and can be expressed in closed form as a state-dependent matrix $\boldsymbol{\theta}^*(s_t) \in \mathbb{R}^{\dim(\mathcal{A}) \times \dim(\mathcal{Z})}$ such that $g_{\boldsymbol{\theta}^*}(z_t, s_t) = \boldsymbol{\theta}^*(s_t)z_t$, where $\dim(\mathcal{Z}) = \text{rank}(\mathbf{G})$ (Theorem B.1, Corollary B.2, and Appendix B). This matrix is a function of the dynamics parameters ϕ and inherits their state dependence; we denote this functional relationship by $\boldsymbol{\theta}^* = \text{PRECODERGENERATOR}(\phi)$. The columns of the precoder matrix span a $\dim(\mathcal{Z})$ -dimensional subspace $\mathcal{M}_{\boldsymbol{\theta}^*}(s_t) := \{\boldsymbol{\theta}^*(s_t)z_t : z_t \in \mathcal{Z}\} \subset \mathcal{A}$, with $\dim(\mathcal{Z}) = \dim(\mathcal{Y})$ in the full-rank case and $\dim(\mathcal{Z}) < \dim(\mathcal{Y})$ in the low-rank case. The subspace defines an equivalence class of empowerment-optimal precoders (Corollary B.3).

V. COMPOSITIONAL MANIFOLD POLICIES

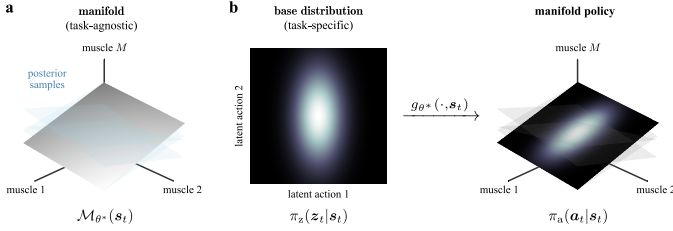


Fig. 2. **Composing the manifold policy:** A policy $\pi_a(a_t|s_t)$ is formed by composing two learned, state-dependent components: (1) a task-agnostic manifold $\mathcal{M}_{\boldsymbol{\theta}^*}(s_t)$ and (2) a task-specific base distribution $\pi_z(z_t|s_t)$. The two components are learned by optimizing different objectives: the manifold is learned by model-based JoSE maximization, and the base distribution is learned via model-free RL. The manifold and the base distribution can be learned simultaneously or sequentially. (a) The action manifold $\mathcal{M}_{\boldsymbol{\theta}^*}(s_t)$; posterior samples visualize manifold uncertainty. (b) The base distribution $\pi_z(z_t|s_t)$ over latent actions (left) is pushed forward through the precoder $g_{\boldsymbol{\theta}^*}(\cdot, s_t)$ onto the manifold, yielding the action density $\pi_a(a_t|s_t)$ (right).

We introduce a class of compositional policies that build task-specific skills on top of task-agnostic action manifolds (Figure 2). An action generated by the policy $a_t \sim \pi_a(a_t|s_t)$ is expressed as a transformation $g_{\boldsymbol{\theta}}$ of a latent action sampled from the base distribution:

$$a_t = g_{\boldsymbol{\theta}}(z_t, s_t) \quad \text{where} \quad z_t \sim \pi_z(z_t|s_t).$$

Many RL algorithms incorporate the policy entropy $\mathcal{H}(\pi_a(\cdot|s_t)) = -\langle \log \pi_a(a_t|s_t) \rangle_{\pi_a}$ into the objective. For actions on the manifold $\mathcal{M}_{\boldsymbol{\theta}^*}(s_t)$, we compute the action density $\pi_a(a_t|s_t)$ via the generalized change-of-variables formula for injective maps (Appendix E-A).

VI. JOSEPI

We train the compositional policy $\pi_a(a_t|s_t)$ using SAC [20]. During the policy improvement step, we only update the base distribution $\pi_z(z_t|s_t)$, with gradients propagated through the fixed precoder. To update the precoder, we train the dynamics network ψ via maximum likelihood estimation on transitions $(s_t^{(i)}, a_t^{(i)}, y_{t+1}^{(i)})$ sampled from the replay buffer

Algorithm 1 JoSEPI

Input: initial base distribution π_z , initial (or pretrained) dynamics network ψ , empty buffer $\mathcal{B}$

for each iteration do

for each environment step do

$z_t \sim \pi_z(z_t|s_t)$ *sample latent action*

$\hat{\phi} \sim q_{\psi}(\phi|s_t)$ *sample dynamics*

$\boldsymbol{\theta} \leftarrow \text{PRECODERGENERATOR}(\hat{\phi})$ *generate precoder*

$a_t \leftarrow \boldsymbol{\theta}z_t$ *generate action*

$s_{t+1} \sim P(s_{t+1}|s_t, a_t)$ *step environment*

$\mathcal{B} \leftarrow \mathcal{B} \cup \{s_t, a_t, r_t, s_{t+1}\}$ *store transition*

end for

for each gradient step do

train ψ via MLE *update dynamics model*

end for

for each gradient step do

train π_z with RL *update base distribution*

end for

end for

$\mathcal{B}$. We apply Monte Carlo (MC) dropout to the hidden units [18], inducing an implicit variational posterior $q_{\psi}(\phi|s_t)$ over the state-dependent dynamics parameters; this in turn induces a posterior over the optimal precoder (Section E-B3). For the low-rank approximation to the input matrix $\mathbf{G}$, we use rank 15, yielding a 15-dimensional state-dependent manifold $\mathcal{M}_{\boldsymbol{\theta}^*}(s_t)$ in a 39-dimensional action space $\mathcal{A}$. We only report results for low-rank factorization variant in the main text.

The full procedure, JoSEPI, is summarized in Algorithm 1; hyperparameters and implementation details in Appendix E.

VII. EXPERIMENTS

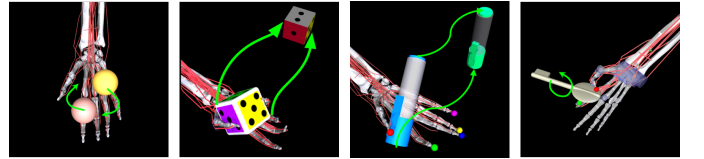


Fig. 3. **In-hand object manipulation with the MyoHand:** We evaluate JoSEPI on a suite of MuJoCo-based manipulation tasks from MyoSuite: BaodingBalls (rotate two balls simultaneously), DieReorient (match target die orientations), PenTwirl (reorient a pen), and KeyTurn (rotate a key).

We evaluate JoSEPI on a suite of contact-rich in-hand manipulation tasks using the MyoHand [11] (23 joints, 39 muscles; four tasks are shown in Figure 3; see Appendix F for full details). Performance is measured as *fraction solved*, the fraction of episode steps in a solved state.

A. Baselines

We compare JoSEPI against SAC [20], Lattice [12] (SAC and rPPO variants), SAR [5], and SD-SAR (Appendix F-C2). All SAC-based methods share identical hyperparameters (Appendix E). Results are averaged over 5 random seeds.

VIII. RESULTS

A. Simultaneous Manifold and Task Learning

We first evaluate JoSEPi in the most demanding setting, where both the state-dependent action manifold $\mathcal{M}_\theta(s_t)$ and the task-specific base distribution $\pi_z(z_t|s_t)$ are learned simultaneously. As shown in Figure 4, JoSEPi consistently outperforms the baselines across all tasks, achieving faster learning, higher asymptotic performance, and improved training stability (see Supplementary Videos S1-S4 for trained JoSEPi behavior). This demonstrates that JoSE-optimized action manifolds provide a strong inductive bias for coordination and exploration in dexterous manipulation tasks.

A post-hoc analysis confirms that our control-affine approximation maintains high prediction accuracy even during local discontinuities caused by contact events (Section G-A). Furthermore, JoSEPi is robust to the dimensionality of the action manifold, determined by the input matrix rank (Section H-B).

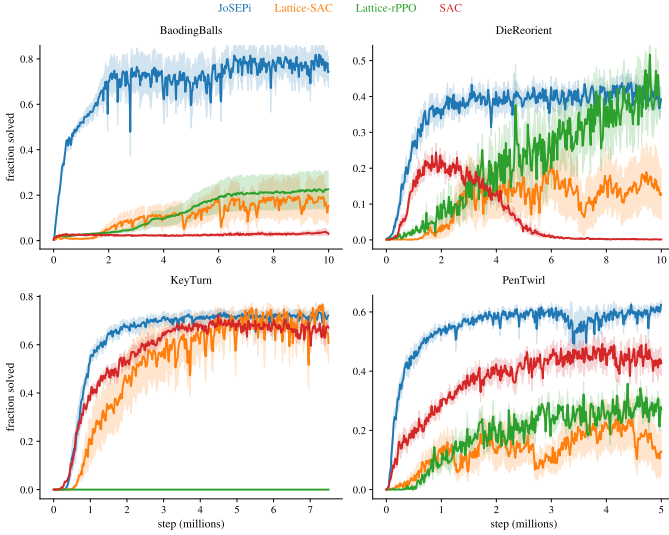


Fig. 4. **Simultaneous manifold and task learning:** Performance of JoSEPi and baselines on BaodingBalls, DieReorient, KeyTurn, and PenTwirl. Data show solved fraction over training (mean $\pm$ s.e.m. across 5 seeds).

B. Dynamics-Based Manifold Discovery

A key feature of JoSEPi is that the action manifold is derived from a dynamics model. Therefore, the learning signal for manifold discovery is dense, as every state transition provides a data point for updating the dynamics model. Model-free task-optimized approaches, in contrast, learn synergies from rewards—which can be sparse. As shown in Figure 5, JoSEPi remains robust under sparse rewards, substantially outperforming task-optimized baselines in both learning speed and final task performance.

C. Decoupled Manifold and Task Learning

In JoSEPi, the policy factorizes into a manifold and a base distribution (Figure 2). We exploit this compositional structure by first discovering the manifold in a play phase, and then

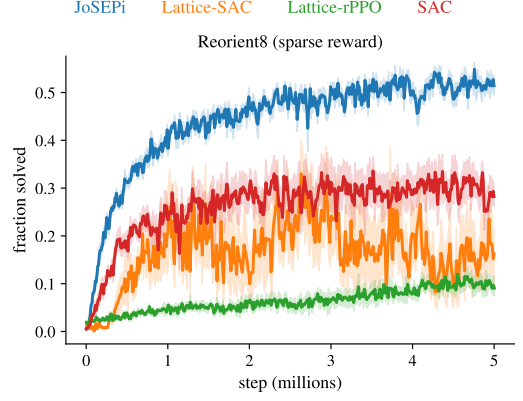


Fig. 5. **JoSEPi discovers manifolds from dynamics, not rewards:** Performance of JoSEPi and baselines on a sparse-reward variant of Reorient8. Data show solved fraction over training (mean $\pm$ s.e.m. across 5 seeds).

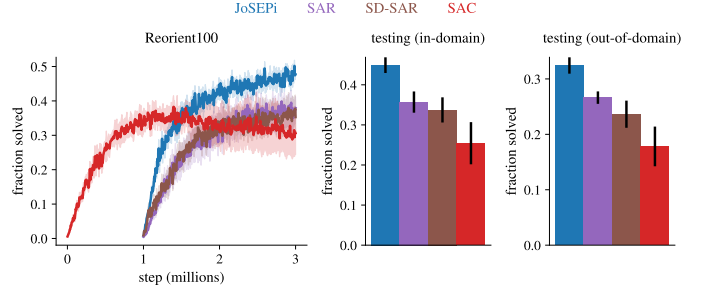


Fig. 6. **Decoupled manifold and task learning:** Performance of JoSEPi and baselines on Reorient100 (left; JoSEPi, SAR and SD-SAR offset to account for play phase) and zero-shot generalization to unseen objects (middle, right). Data show solved fraction (mean $\pm$ s.e.m. across 5 seeds).

using the manifold to accelerate learning on a downstream task (see Appendix F for protocol details).

As shown in Figure 6 (left), JoSEPi substantially outperforms SAR and SD-SAR on the downstream task. The JoSE objective yields a manifold that provides a reusable foundation for building new policies as it is task-agnostic. In contrast, the demonstration-based baselines extract a manifold that reflects specific play-phase behaviors, limiting transfer to new tasks.

After training on the downstream task, we evaluate zero-shot generalization of the trained policy to 1000 unseen objects with in-domain (ReorientID) and out-of-domain (ReorientOOD) properties. JoSEPi achieves the best performance in both conditions (Figure 6, right), demonstrating that JoSE-optimized manifolds provide an effective scaffold for learning policies that generalize more robustly to novel objects.

IX. DISCUSSION

Here we develop an information-theoretic framework for learning action manifolds that maximize control over low-dimensional state features. Although instantiated in musculoskeletal systems, the core approach is relevant to any system where targeted control is sufficient for dexterous behavior.

REFERENCES

- [1] Alekh Agarwal, Sham Kakade, Akshay Krishnamurthy, and Wen Sun. FLAMBE: Structural complexity and representation learning of low rank MDPs. In *Advances in Neural Information Processing Systems*, volume 33, pages 20095–20107. Curran Associates, Inc., 2020.
- [2] Pierre Baldi and Kurt Hornik. Neural networks and principal component analysis: Learning from examples without local minima. *Neural Networks*, 2(1):53–58, January 1989. ISSN 08936080. doi: 10.1016/0893-6080(89)90014-2.
- [3] Nir Baram, Shie Mannor, and Ron Meir. Action redundancy in reinforcement learning. In *Proceedings of the 37th Conference on Uncertainty in Artificial Intelligence*, volume 161, pages 376–385. PMLR, 2021.
- [4] Thomas Bendokat, Ralf Zimmermann, and P.-A. Absil. A Grassmann manifold handbook: Basic geometry and computational aspects. *Advances in Computational Mathematics*, 50(1):6, 2024. doi: 10.1007/s10444-023-10090-8.
- [5] Cameron H Berg, Vittorio Caggiano, and Vikash Kumar. SAR: Generalization of physiological dexterity via synergistic action representation. In *Robotics: Science and Systems*, July 2023. doi: 10.15607/RSS.2023.XIX.007.
- [6] Max Berniker, Anthony M. Jarc, Emilio Bizzi, and Matthew C. Tresch. Simplified and effective motor control based on muscle synergies to exploit musculoskeletal dynamics. *Proceedings of the National Academy of Sciences*, 106:7601 – 7606, 2009.
- [7] Homanga Bharadhwaj, Mohammad Babaeizadeh, Dumitru Erhan, and Sergey Levine. INFOrmation prioritization through EmPOWERment in visual model-based RL. In *International Conference on Learning Representations*, 2022.
- [8] H. Bourlard and Y. Kamp. Auto-association by multilayer perceptrons and singular value decomposition. *Biological Cybernetics*, 59(4):291–294, September 1988. ISSN 1432-0770. doi: 10.1007/BF00332918.
- [9] Stephen Boyd and Lieven Vandenbergh. *Convex optimization*. Cambridge University Press, 2004. ISBN 9780521833783.
- [10] Vittorio Caggiano, Guillaume Durandau, Huwawei Wang, Alberto Chiappa, Alexander Mathis, Pablo Tano, Nisheet Patel, Alexandre Pouget, Pierre Schumacher, Georg Martius, Daniel Haeufle, Yiran Geng, Boshi An, Yifan Zhong, Jiaming Ji, Yuanpei Chen, Hao Dong, Yaodong Yang, Rahul Siripurapu, Luis Eduardo Ferro Diez, Michael Kopp, Vihang Patil, Sepp Hochreiter, Yuval Tassa, Josh Merel, Randy Schultheis, Seungmoon Song, Massimo Sartori, and Vikash Kumar. MyoChallenge 2022: Learning contact-rich manipulation using a musculoskeletal hand. In *Proceedings of the NeurIPS 2022 competition track*, volume 220, pages 233–250. PMLR, 2022.
- [11] Vittorio Caggiano, Huawei Wang, Guillaume Durandau, Massimo Sartori, and Vikash Kumar. MyoSuite: A contact-rich simulation suite for musculoskeletal motor control. In *Proceedings of the 4th Annual Learning for Dynamics and Control Conference*, volume 168, pages 492–507. PMLR, 2022.
- [12] Alberto Silvio Chiappa, Alessandro Marin Vargas, Ann Zixiang Huang, and Alexander Mathis. Latent exploration for reinforcement learning. In *Advances in Neural Information Processing Systems*, volume 36. Curran Associates, Inc., 2023.
- [13] Thomas M. Cover and Joy A. Thomas. *Elements of information theory*. Wiley, 2nd edition, September 2005. ISBN 978-0-471-24195-9 978-0-471-74882-3. doi: 10.1002/047174882X.
- [14] Andrea d’Avella, Philippe Saltiel, and Emilio Bizzi. Combinations of muscle synergies in the construction of a natural motor behavior. *Nature Neuroscience*, 6(3): 300–308, March 2003.
- [15] Gabriel Dulac-Arnold, Richard Evans, Peter van den Driessche, Peter Glenz, Jakub M. Czarnecki, Johannes Rauh, Tom Schaul, and Iain Leahy. Deep reinforcement learning in large discrete action spaces. *arXiv preprint arXiv:1512.07679*, 2015.
- [16] Benjamin Eysenbach, Abhishek Gupta, Julian Ibarz, and Sergey Levine. Diversity is all you need: Learning skills without a reward function. In *International Conference on Learning Representations*, 2019.
- [17] Yusen Feng, Xiyan Xu, and Libin Liu. MuscleVAE: Model-based controllers of muscle-actuated characters. In *SIGGRAPH Asia 2023 conference papers*. Association for Computing Machinery, 2023. ISBN 979-8-4007-0315-7. doi: 10.1145/3610548.3618137.
- [18] Yarin Gal and Zoubin Ghahramani. Dropout as a Bayesian approximation: Representing model uncertainty in deep learning. In *Proceedings of the 33rd International Conference on Machine Learning*, volume 48, pages 1050–1059. PMLR, 2016.
- [19] Karol Gregor, Danilo Jimenez Rezende, and Daan Wierstra. Variational intrinsic control. In *ICLR Workshop*, 2017.
- [20] Tuomas Haarnoja, Aurick Zhou, Kristian Hartikainen, George Tucker, Sehoon Ha, Jie Tan, Vikash Kumar, Henry Zhu, Abhishek Gupta, Pieter Abbeel, and Sergey Levine. Soft Actor-Critic algorithms and applications, 2018. arXiv:1812.05905.
- [21] Kaibo He, Chenhui Zuo, Chengtian Ma, and Yanan Sui. DynSyn: Dynamical synergistic representation for efficient learning and control in overactuated embodied systems. In *Proceedings of the 41st International Conference on Machine Learning*, volume 235, pages 18115–18132, Vienna, Austria, 2024. PMLR.
- [22] Tobias Jung, Daniel Polani, and Peter Stone. Empowerment for continuous agent-environment systems. *Adaptive Behavior*, 19(1):16–39, 2011.
- [23] Diederik P. Kingma and Jimmy Ba. Adam: A method for stochastic optimization. In *International Conference on Learning Representations*, 2015.
- [24] A S Klyubin, D Polani, and C L Nehaniv. Empowerment: A universal agent-centric measure of control. In

- IEEE Congress on Evolutionary Computation*, Edinburgh, Scotland, UK, 2005.
- [25] Mark L. Latash. *Synergy*. Oxford University Press, 2008.
 - [26] Seunghwan Lee, Moonseok Park, Kyoungmin Lee, and Jehee Lee. Scalable muscle-actuated human simulation and control. *ACM Trans. Graph.*, 38(4), July 2019. ISSN 0730-0301. doi: 10.1145/3306346.3322972.
 - [27] Shakir Mohamed and Danilo Jimenez Rezende. Variational information maximisation for intrinsically motivated reinforcement learning. In *Advances in Neural Information Processing Systems*, volume 28, pages 2116–2124. Curran Associates, Inc., 2015.
 - [28] Jungnam Park, Moon Seok Park, Jehee Lee, and Jungdam Won. Bidirectional GaitNet: A bidirectional prediction model of human gait and anatomical conditions. In *ACM SIGGRAPH 2023 conference proceedings*. Association for Computing Machinery, 2023. ISBN 979-8-4007-0159-7. doi: 10.1145/3588432.3591492.
 - [29] Antonin Raffin, Ashley Hill, Adam Gleave, Anssi Kanervisto, Maximilian Ernestus, and Noah Dormann. Stable-baselines3: Reliable reinforcement learning implementations. *Journal of Machine Learning Research*, 22(268): 1–8, 2021.
 - [30] Sahand Rezaei-Shoshtari, Rosie Zhao, Prakash Panagaden, David Meger, and Doina Precup. Continuous MDP homomorphisms and homomorphic policy gradient. In *Advances in Neural Information Processing Systems*, volume 35. Curran Associates, Inc., 2022.
 - [31] Christoph Salge, Cornelius Glackin, and Daniel Polani. Changing the environment based on empowerment as intrinsic motivation. *Entropy. An International and Interdisciplinary Journal of Entropy and Information Studies*, 16(5):2789–2819, May 2014. ISSN 1099-4300. doi: 10.3390/e16052789.
 - [32] Alain Sarlette and Rodolphe Sepulchre. Consensus optimization on manifolds. *SIAM Journal on Control and Optimization*, 48(1):56–76, 2009. doi: 10.1137/060673400.
 - [33] Pierre Schumacher, Daniel F.B. Haeufle, Dieter Büchler, Syn Schmitt, and Georg Martius. DEP-RL: Embodied exploration for reinforcement learning in overactuated and musculoskeletal systems. In *International Conference on Learning Representations*, May 2023.
 - [34] Archit Sharma, Shixiang Gu, Sergey Levine, Vikash Kumar, and Karol Hausman. Dynamics-aware unsupervised discovery of skills. In *International Conference on Learning Representations*, 2020.
 - [35] C. S. Sherrington. Flexion-reflex of the limb, crossed extension-reflex, and reflex stepping and standing. *The Journal of Physiology*, 40(1-2):28–121, April 1910. ISSN 0022-3751. doi: 10.1113/jphysiol.1910.sp001362.
 - [36] Roland Stolz, Hanna Krasowski, Jakob Thumm, Michael Eichelbeck, Philipp Gassert, and Matthias Althoff. Excluding the irrelevant: Focusing reinforcement learning through continuous action masking. In *Advances in Neural Information Processing Systems*, volume 37, pages 95067–95094. Curran Associates, Inc., 2024. doi: 10.52202/079017-3013.
 - [37] Jonathan Taylor, Doina Precup, and Prakash Panagaden. Bounding performance loss in approximate MDP homomorphisms. In *Advances in Neural Information Processing Systems*, volume 21, pages 1649–1656. Curran Associates, Inc., 2008.
 - [38] Gal Tennenholtz, Shie Mannor, and Ron Meir. The natural language of actions. In *Proceedings of the 36th International Conference on Machine Learning*, volume 97, pages 10802–10811. PMLR, 2019.
 - [39] William R Thompson. On the likelihood that one unknown probability exceeds another in view of the evidence of two samples. *Biometrika*, 25(3/4):285–294, 1933.
 - [40] Emanuel Todorov and Michael I. Jordan. Optimal feedback control as a theory of motor coordination. *Nature Neuroscience*, 5:1226–1235, 2002.
 - [41] M C Tresch, P Saltiel, and E Bizzi. The construction of movement by the spinal cord. *Nature Neuroscience*, 2(2): 162–167, February 1999.
 - [42] Cheryl Wang, Chun Kwang Tan, Balint K. Hodossy, Shirui Lyu, Pierre Schumacher, James Heald, Kai Biegun, Samo Hromadka, Maneesh Sahani, Gunwoo Park, Beomsoo Shin, JongHyun Park, Seungbum Koo, Chenhui Zuo, Chengtian Ma, Yanan Sui, Nicklas Hansen, Stone Tao, Yuan Gao, Hao Su, Seungmoon Song, Letizia Gionfrida, Massimo Sartori, Guillaume Durandau, Vikash Kumar, and Vittorio Caggiano. MyoChallenge 2024: A new benchmark for physiological dexterity and agility in bionic humans. In *NeurIPS 2025 Datasets and Benchmarks Track*, 2025.
 - [43] Yunyue Wei, Chenhui Zuo, and Yanan Sui. Scalable exploration for high-dimensional continuous control via value-guided flow. In *International Conference on Learning Representations*, 2026.
 - [44] Cong Wu, Aravind Rajeswaran, Igor Mordatch, and Vikash Kumar. Variance reduction for policy gradient with action-dependent factorized baselines. In *International Conference on Learning Representations*, 2018.
 - [45] Lin Yang and Mengdi Wang. Reinforcement learning in feature space: Matrix bandit, kernels, and regret bound. In *Proceedings of the 37th International Conference on Machine Learning*, volume 119, pages 10746–10756. PMLR, November 2020.
 - [46] Tom Zahavy, Matan Haroush, Nadav Merlis, Daniel J. Mankowitz, and Shie Mannor. Learn what not to learn: Action elimination with deep reinforcement learning. In *Advances in Neural Information Processing Systems*, volume 31, pages 3566–3577. Curran Associates, Inc., 2018.

APPENDIX A EMPOWERMENT

The amount of control an agent can have over the world can be quantified as the capacity of the communication channel between its actions and future states, known as *empowerment* [24] (Definition A.1). Empowerment has been used for intrinsic motivation [22, 31, 27], as an objective for learning latent state representations [7], and as a criterion for unsupervised skill discovery [19, 16, 34].

Definition A.1 (Empowerment). *Given a communication channel $p(s_{t+1}|s_t, a_t)$, the empowerment $\mathcal{E}(s_t)$ in a state $s_t \in \mathcal{S}$ is defined as the maximum mutual information between actions and next states over all source distributions $p(a_t|s_t)$ (i.e., the channel capacity):*

$$\mathcal{E}(s_t) = \max_{p(a_t|s_t)} \mathcal{I}(a_t; s_{t+1}|s_t).$$

APPENDIX B THE OPTIMAL PRECODING PROBLEM

A. The Setting

Consider a communication problem involving a sender, a receiver and a state-dependent communication channel. The sender transmits a signal drawn from a Gaussian source distribution $z_t \sim \mathcal{N}(\mathbf{0}, \Sigma(s_t))$. This signal is spatially directed by a learnable, state-dependent precoder matrix $\theta(s_t)$. The precoded signal $a_t = \theta(s_t)z_t$ is sent to the receiver via a control-affine Gaussian channel:

$$p(y_{t+1}|a_t, s_t) = \mathcal{N}(f(s_t) + G(s_t)a_t, Q(s_t)).$$

Here, $G \in \mathbb{R}^{\dim(\mathcal{Y}) \times \dim(\mathcal{A})}$ and $\theta \in \mathbb{R}^{\dim(\mathcal{A}) \times \dim(\mathcal{Z})}$.

The composition of the precoder and the channel defines an *effective communication channel* from the source signal to the received signal:

$$p(y_{t+1}|z_t, s_t) = \mathcal{N}(f(s_t) + G_{\text{eff}}(s_t)z_t, Q(s_t)),$$

where $G_{\text{eff}}(s_t) := G(s_t)\theta(s_t)$.

The optimal precoding problem is to find a state-dependent precoder $\theta(s_t)$ that maximizes the capacity of this effective communication channel (i.e., joint-space empowerment), subject to a transmit power constraint.

For the remainder of the section, we suppress the dependence of the parameters f , G , θ , Q and Σ on the state s_t to simplify notation.

B. The Problem

Maximizing the capacity of the communication channel with respect to the precoder is equivalent to maximizing the mutual information with respect to both the precoder and the covariance of the source distribution:

$$\begin{aligned} & \underset{\theta, \Sigma}{\text{maximize}} && \mathcal{I}_\theta(z_t; y_{t+1}|s_t) \\ & \text{subject to} && \Sigma \succeq 0, \text{Tr}(\theta\Sigma\theta^\top) = P, \end{aligned} \quad (\text{P1a})$$

where $\Sigma \succeq 0$ denotes that Σ is positive semidefinite (PSD), and $P > 0$ is a scalar transmit power budget. The constraint on the total transmit power $\text{Tr}(\theta\Sigma\theta^\top)$ (the trace of the covariance

of the transmitted signal) ensures that the mutual information is bounded; the capacity of an affine-Gaussian channel is infinite otherwise.

Because the channel noise is independent of the source signal, maximizing the mutual information $\mathcal{I}_\theta(z_t; y_{t+1}|s_t) = \mathcal{H}_\theta[y_{t+1}|s_t] - \mathcal{H}_\theta[y_{t+1}|z_t, s_t]$ is equivalent to maximizing the marginal entropy $\mathcal{H}_\theta[y_{t+1}|s_t]$. As the marginal distribution for the setting considered here is $p_\theta(y_{t+1}|s_t) = \mathcal{N}(f, G\theta\Sigma\theta^\top G^\top + Q)$, the optimization problem in P1a can be equivalently written as

$$\begin{aligned} & \underset{\theta, \Sigma}{\text{maximize}} && \det(G\theta\Sigma\theta^\top G^\top + Q) \\ & \text{subject to} && \Sigma \succeq 0, \text{Tr}(\theta\Sigma\theta^\top) = P. \end{aligned} \quad (\text{P1b})$$

C. The Solution

Let $Q = \Psi\Lambda\Psi^\top$ denote the eigendecomposition of the noise covariance matrix Q , and define the whitened channel matrix $\tilde{G} = \Lambda^{-\frac{1}{2}}\Psi^\top G$. Since $\det(G\theta\Sigma\theta^\top G^\top + Q) = \det(Q) \det(\tilde{G}\theta\Sigma\theta^\top \tilde{G}^\top + I)$ and $\det(Q)$ does not depend on θ or Σ , the argmax of P1b coincides with that of

$$\begin{aligned} & \underset{\theta, \Sigma}{\text{maximize}} && \det(\tilde{G}\theta\Sigma\theta^\top \tilde{G}^\top + I) \\ & \text{subject to} && \Sigma \succeq 0, \text{Tr}(\theta\Sigma\theta^\top) = P, \end{aligned} \quad (\text{P1c})$$

where $\tilde{G}\theta\Sigma\theta^\top \tilde{G}^\top + I$ denotes the marginal covariance expressed in the whitened basis.

Next, let $\tilde{G} = U_r S_r V_r^\top$ denote the compact singular value decomposition (SVD) of the rank- r matrix $\tilde{G}$, retaining only nonzero singular values and their corresponding singular vectors; thus $U_r \in \mathbb{R}^{\dim(\mathcal{Y}) \times r}$, $S_r \in \mathbb{R}^{r \times r}$, and $V_r \in \mathbb{R}^{\dim(\mathcal{A}) \times r}$.

Theorem B.1 (Optimal Precoder). *A solution to the optimization problem in P1a is given by*

$$\theta^* = V_r,$$

the right singular vectors of $\tilde{G}$ associated with nonzero singular values.

Proof: The objective depends on θ and Σ only through the covariance of the transmitted signal $\Phi \triangleq \theta\Sigma\theta^\top$, so problem P1c is equivalent to

$$\begin{aligned} & \underset{\Phi}{\text{maximize}} && \det(\tilde{G}\Phi\tilde{G}^\top + I) \\ & \text{subject to} && \Phi \succeq 0, \text{Tr}(\Phi) = P. \end{aligned} \quad (\text{P1d})$$

Any component of Φ that lies in the nullspace of $\tilde{G}$ makes zero contribution to $\tilde{G}\Phi\tilde{G}^\top$. Since $\det(\tilde{G}\Phi\tilde{G}^\top + I)$ is monotone non-decreasing on the PSD cone [9], redirecting power from the nullspace into $\text{range}(V_r)$ — its orthogonal complement under the compact SVD $\tilde{G} = U_r S_r V_r^\top$ — cannot decrease the objective. An optimal solution can therefore be parameterized as $\Phi = V_r \Gamma V_r^\top$ for some $\Gamma \succeq 0$, where the power constraint reduces to $\text{Tr}(\Gamma) = P$ by the cyclic property of the trace.

Substituting $\tilde{G} = U_r S_r V_r^\top$, $\Phi = V_r \Gamma V_r^\top$ and applying $\det(I + AB) = \det(I + BA)$ gives

$$\det(\tilde{G}\Phi\tilde{G}^\top + I) = \det(S_r \Gamma S_r^\top + I),$$

and hence P1d reduces to

$$\begin{aligned} & \underset{\mathbf{\Gamma}}{\text{maximize}} \quad \det(\mathbf{S}_r \mathbf{\Gamma} \mathbf{S}_r^\top + \mathbf{I}) \\ & \text{subject to} \quad \mathbf{\Gamma} \succeq 0, \text{Tr}(\mathbf{\Gamma}) = P. \end{aligned} \quad (\text{P1e})$$

By the Hadamard inequality, for any positive semidefinite matrix $\mathbf{K}$,

$$\det(\mathbf{K}) \leq \prod_i [\mathbf{K}]_{ii},$$

with equality iff $\mathbf{K}$ is diagonal. Applied to $\mathbf{K} = \mathbf{S}_r \mathbf{\Gamma} \mathbf{S}_r^\top + \mathbf{I}$, the objective is bounded by $\prod_i (\sigma_i^2[\mathbf{\Gamma}]_{ii} + 1)$, where σ_i are the singular values in $\mathbf{S}_r$. This bound depends only on the diagonal of $\mathbf{\Gamma}$, which suggests discarding the off-diagonal entries. To this end, let $\mathbf{\Gamma}_d$ denote the matrix obtained by setting the off-diagonal entries of $\mathbf{\Gamma}$ to zero. Since both $\mathbf{\Gamma}_d$ and $\mathbf{S}_r$ are diagonal, $\mathbf{K} = \mathbf{S}_r \mathbf{\Gamma}_d \mathbf{S}_r^\top + \mathbf{I}$ is diagonal, and hence by the equality condition above, $\mathbf{\Gamma}_d$ attains the bound:

$$\det(\mathbf{S}_r \mathbf{\Gamma} \mathbf{S}_r^\top + \mathbf{I}) \leq \prod_i (\sigma_i^2[\mathbf{\Gamma}]_{ii} + 1) = \det(\mathbf{S}_r \mathbf{\Gamma}_d \mathbf{S}_r^\top + \mathbf{I}).$$

Therefore, replacing any feasible $\mathbf{\Gamma}$ by $\mathbf{\Gamma}_d$ does not decrease the objective, so an optimal solution $\mathbf{\Gamma}^*$ to P1e can be taken to be diagonal. Recalling that $\mathbf{\Phi} = \mathbf{V}_r \mathbf{\Gamma} \mathbf{V}_r^\top$, an optimal transmit covariance for P1d is $\mathbf{\Phi}^* = \mathbf{V}_r \mathbf{\Gamma}^* \mathbf{V}_r^\top$, which is realized by the precoder $\mathbf{\theta}^* = \mathbf{V}_r$ together with a diagonal source covariance $\mathbf{\Sigma}^* = \mathbf{\Gamma}^*$. This proves that $\mathbf{V}_r$, the right singular vectors of $\tilde{\mathbf{G}}$ associated with nonzero singular values, is an optimal precoder. ■

Corollary B.2 (Optimal Number of Synergies). *The closed-form solution $\mathbf{\theta}^* = \mathbf{V}_r \in \mathbb{R}^{\dim(\mathcal{A}) \times r}$ enforces $\dim(\mathcal{Z}) = r = \text{rank}(\tilde{\mathbf{G}})$. Each column of $\mathbf{V}_r$ corresponds to a single synergy, and the r synergies span an r -dimensional subspace of the action space $\mathcal{A}$. Any synergy orthogonal to this subspace lies in the nullspace of $\tilde{\mathbf{G}}$, contributes nothing to the objective, and receives zero power at optimality. For overactuated musculoskeletal systems, where muscles outnumber joints ($\dim(\mathcal{A}) > \dim(\mathcal{Y})$), the full-rank case gives $r = \dim(\mathcal{Y})$, that is, the number of synergies equals the number of joints.*

Corollary B.3 (Optimal Subspace). *The empowerment-optimal subspace is $\text{range}(\mathbf{V}_r) \in \text{Gr}(k, \dim(\mathcal{A}))$ (Section C), the Grassmannian of k -dimensional subspaces of the action space—the intrinsic parameter space for the precoder. Any precoder whose columns span the same subspace as $\mathbf{V}_r$ is empowerment-optimal. Formally, for any invertible matrix $\mathbf{T}$, the precoder $\mathbf{V}_r \mathbf{T}$ together with source covariance $\mathbf{T}^{-1} \mathbf{\Sigma}^* \mathbf{T}^{-\top}$ induces the same transmit covariance $\mathbf{\Phi}^*$, and hence achieves the same objective value. Theorem B.1 therefore identifies $\mathbf{V}_r$ as one canonical representative of this optimal subspace.*

Remark B.4 (Optimal Power Allocation). The proof of Theorem B.1 does not require the explicit values of the diagonal entries of $\mathbf{\Gamma}^*$. For completeness, they are given by the standard waterfilling solution [13]:

$$\max_{\mathbf{\Gamma}_{ii} \geq 0, \sum_i \mathbf{\Gamma}_{ii} = P} \prod_i (\sigma_i^2 \mathbf{\Gamma}_{ii} + 1),$$

which allocates more power to directions with larger singular values σ_i of the whitened channel matrix $\tilde{\mathbf{G}}$.

D. Choice of Subspace Basis

While the right singular vectors of $\tilde{\mathbf{G}}$ provide a principled and theoretically optimal characterization of the action subspace, they are not always the most convenient basis for use in learning algorithms. In particular, singular vectors are defined only up to sign, and their ordering depends on the corresponding singular values. As a result, small changes in $\tilde{\mathbf{G}}$ —for example due to variation in the state $\mathbf{s}_t$ or updates to the dynamics model parameters ψ —can induce discontinuous changes in the basis, such as sign flips or reordering of vectors. Moreover, repeated SVD computation during training can be computationally expensive for large matrices.

To address these practical considerations, we adopt a more stable and well-behaved basis for the same subspace. Specifically, we use the rows of $\tilde{\mathbf{G}}$, normalized to have unit norm. These vectors span the same subspace as the right singular vectors $\mathbf{V}_r$, while varying continuously with $\tilde{\mathbf{G}}$ and avoiding the ambiguities inherent to the SVD representation. By Corollary B.3, any basis spanning the same subspace as $\mathbf{V}_r$ is empowerment-optimal; the row-normalized basis therefore preserves optimality while varying continuously with $\tilde{\mathbf{G}}$.

APPENDIX C THE GRASSMANNIAN

Here we introduce the necessary background on the Grassmannian, which arises in our theoretical analysis (Section B) and provides the geometric framework for the subspace distance and dispersion measures used in our empirical analyses (Section G-C).

The set of all k -dimensional subspaces of $\mathbb{R}^M$ forms a smooth manifold known as the Grassmannian, denoted $\text{Gr}(k, M)$ [4]. Points on the manifold can be represented in several forms, including as an orthonormal basis $\mathbf{V} \in \mathbb{R}^{M \times k}$ or as a unique projection matrix $\mathbf{P} = \mathbf{V} \mathbf{V}^\top \in \mathbb{R}^{M \times M}$. The choice of representation gives rise to different extrinsic distances between points on $\text{Gr}(k, M)$, measured via the Frobenius norm: the Procrustes distance (based on the orthonormal basis representation) and the chordal distance (based on the projection matrix representation). The geodesic distance, the length of the shortest path along the manifold itself, is a representation-independent, intrinsic alternative. In our analyses, we use the chordal distance, as it allows the variance of a set of subspaces to be computed analytically, unlike the geodesic and Procrustes distances, which require iterative optimization.

a) *Chordal distance.*: The chordal distance d_C between any two subspaces $\mathcal{U}, \mathcal{V} \in \text{Gr}(k, M)$ is defined via the Frobenius norm of the difference of their projection matrices:

$$d_C(\mathcal{U}, \mathcal{V}) = \frac{1}{\sqrt{2}} \|\mathbf{P}_{\mathcal{U}} - \mathbf{P}_{\mathcal{V}}\|_F. \quad (2)$$

This distance corresponds to the ℓ_2 -norm of the sines of the principal angles $\boldsymbol{\varphi} = (\varphi_1, \dots, \varphi_k)$, satisfying $d_C^2 = \sum_{i=1}^k \sin^2 \varphi_i$. The principal angles, ordered $0 \leq \varphi_1 \leq \dots \leq \varphi_k \leq \pi/2$,

measure how far apart two subspaces are, from their most aligned direction to their most misaligned; the i -th principal angle φ_i is the smallest angle between any unit vector in $\mathcal{U}$ and any unit vector in $\mathcal{V}$, subject to both vectors being orthogonal to the directions that define $\varphi_1, \dots, \varphi_{i-1}$.¹

b) Normalized chordal distance.: The maximum chordal distance on $\text{Gr}(k, M)$ is $d_{C, \max} = \sqrt{\min(k, M-k)}$. This maximum is determined by the intersection constraints of subspaces:

- When $k \leq M/2$, subspaces can be maximally separated (i.e., all k principal angles equal $\pi/2$), yielding $d_{C, \max} = \sqrt{k}$.
- When $k > M/2$, any two subspaces must share an intersection of dimension at least $2k - M$, which forces $2k - M$ principal angles to zero; only the remaining $M - k$ angles can reach $\pi/2$, yielding $d_{C, \max} = \sqrt{M - k}$.

Scaling (2) by this value yields a dimension-invariant distance $\bar{d}_C \in [0, 1]$:

$$\bar{d}_C(\mathcal{U}, \mathcal{V}) = \frac{d_C(\mathcal{U}, \mathcal{V})}{d_{C, \max}}, \quad (3)$$

which enables comparison of subspace distances across different dimensionalities.

c) Fréchet variance.: To quantify the dispersion of a collection of subspaces $\{\mathcal{U}_1, \dots, \mathcal{U}_N\} \subset \text{Gr}(k, M)$, we use the Fréchet variance:

$$\text{Var}_{\mathcal{F}} = \frac{1}{N} \sum_{i=1}^N d_C^2(\mathcal{U}_i, \bar{\mathcal{U}}), \quad (4)$$

where $\bar{\mathcal{U}} = \arg \min_{\mathcal{U} \in \text{Gr}(k, M)} \sum_{i=1}^N d_C^2(\mathcal{U}_i, \mathcal{U})$ is the Fréchet mean. The Fréchet variance is the natural analogue of Euclidean variance for data on curved manifolds.

A key advantage of the chordal distance is that it admits an analytic solution for the Fréchet mean and variance, avoiding the need for iterative manifold optimization.² Specifically, while the ambient arithmetic mean $\bar{\mathbf{P}} = \frac{1}{N} \sum_{i=1}^N \mathbf{P}_i$ does not itself lie on the Grassmannian, the Fréchet mean $\mathbf{P}_{\bar{\mathcal{U}}}$, which does, can be obtained directly from $\bar{\mathbf{P}}$. By Proposition 2 of Sarlette and Sepulchre [32], projecting $\bar{\mathbf{P}}$ back onto $\text{Gr}(k, M)$ corresponds to constructing $\mathbf{P}_{\bar{\mathcal{U}}}$ from its k principal eigenvectors. The exact Fréchet variance is then evaluated via the explicit formula:

$$\text{Var}_{\mathcal{F}} = \frac{1}{2N} \sum_{i=1}^N \|\mathbf{P}_i - \mathbf{P}_{\bar{\mathcal{U}}}\|_F^2. \quad (5)$$

This expression provides a principled, exact measure of alignment consistency for a set of subspaces.

¹The intrinsic geodesic is typically obtained by equipping $\text{Gr}(k, M)$ with the canonical Riemannian metric, defined via the Frobenius inner product $\langle \Delta_1, \Delta_2 \rangle = \text{tr}(\Delta_1^T \Delta_2)$ for $\Delta_1, \Delta_2 \in \mathbb{R}^{M \times M}$ tangent to $\text{Gr}(k, M)$. This metric is equivalent to the standard Euclidean inner product of the vectorized matrices, $\langle \text{vec}(\Delta_1), \text{vec}(\Delta_2) \rangle$. The resulting geodesic distance is the ℓ_2 -norm of the principal angles, $d(\mathcal{U}, \mathcal{V}) = \|\varphi\|_2$, with maximum $d_{\max} = \sqrt{\min(k, M-k)} \pi/2$. Since $\sin \varphi \leq \varphi$ for $\varphi \in [0, \pi/2]$, the chordal distance satisfies $d_C \leq d$, with equality in the small-angle limit.

²Under the geodesic metric, the Fréchet mean lacks an analytic solution and must be computed via iterative optimization (e.g., gradient descent on the manifold).

APPENDIX D MUSCULOSKELETAL DYNAMICS

The dynamics of the musculoskeletal system are modeled as a third-order system combining individual first-order muscle activation and contractile dynamics with second-order body dynamics:

$$\dot{m}_i = (a_i - m_i)/\tau(a_i, m_i),$$

$$f_{m,i} = m_i \cdot f_a(l_i) \cdot f_v(\dot{l}_i) + f_p(l_i),$$

$$f_{\text{mtu},i} = f_{m,i} \cos \alpha_i,$$

$$\mathbf{M}(\mathbf{q})\ddot{\mathbf{q}} + \mathbf{C}(\mathbf{q}, \dot{\mathbf{q}})\dot{\mathbf{q}} + \mathbf{g}(\mathbf{q}) = \mathbf{J}^T \mathbf{f}_{\text{mtu}},$$

where m_i and $\dot{m}_i$ are the activation and its time derivative for muscle i ; a_i is the corresponding neural control input; $\tau(a_i, m_i)$ is the activation/deactivation time constant; l_i , $\dot{l}_i$, and α_i are the muscle length, velocity, and pennation angle, respectively; f_a , f_v , and f_p are the active force-length, force-velocity, and passive force-length functions; $f_{m,i}$ and $f_{\text{mtu},i}$ are the muscle and muscle-tendon unit forces; and $\mathbf{f}_{\text{mtu}} = (f_{\text{mtu},1}, \dots, f_{\text{mtu},M})^T$ is the vector of all muscle-tendon forces. The global skeletal dynamics are governed by the mass matrix $\mathbf{M}$, Coriolis forces $\mathbf{C}$, gravity $\mathbf{g}$, and the Jacobian transpose $\mathbf{J}^T$, which maps the aggregate muscle-tendon forces to generalized joint coordinates.

The joint accelerations are an affine function of the activations. Because the MuJoCo simulator updates joint velocities via (semi-implicit) Euler integration ($\dot{\mathbf{q}}_{t+1} = \dot{\mathbf{q}}_t + \ddot{\mathbf{q}}_t \Delta t$), and integration is a linear operation, the discrete-time joint velocities are correspondingly affine in the activations. Consequently, if $\tau(\mathbf{a}, \mathbf{m})$ is independent of the neural control signals (i.e., $\tau(\mathbf{a}, \mathbf{m}) \equiv \tau(\mathbf{m})$), the next-step joint velocities are also an affine function of the controls, making the control-affine modelling assumption in Equation (1) exact. Often, a different time constant is used for muscle activation versus deactivation:

$$\tau(\mathbf{a}, \mathbf{m}) = \begin{cases} \tau_{\text{act}}(\mathbf{m}) & \mathbf{a} - \mathbf{m} > 0 \\ \tau_{\text{deact}}(\mathbf{m}) & \mathbf{a} - \mathbf{m} \leq 0, \end{cases}$$

in which case the dynamics are piecewise affine in the control signal, and our control-affine assumption remains a robust approximation for the effects of neural control signals on the joint velocities.

APPENDIX E IMPLEMENTATION

For Lattice-SAC and Lattice-rPPO, we use the authors' publicly available implementation on GitHub. All other methods (JoSEPi, SAR, SAC) are based on the Stable-Baselines3 JAX (SBX) [29] implementation of SAC (hyperparameters in Table I). We parallelize agent-environment interactions across 20 CPU cores to accelerate training.

A. Action Density on the Manifold

For JoSEPi, we compute the density of an action $\mathbf{a}_t$ on the manifold $\mathcal{M}_\theta(\mathbf{s}_t)$ using the change-of-variables formula:

$$\pi_{\mathbf{a}}(\mathbf{a}_t|\mathbf{s}_t) = \pi_{\mathbf{z}}(g_\theta^\dagger(\mathbf{a}_t, \mathbf{s}_t)|\mathbf{s}_t) \times \left(\det \left(\mathbf{J}_{g_\theta}^\top(g_\theta^\dagger(\mathbf{a}_t, \mathbf{s}_t), \mathbf{s}_t) \mathbf{J}_{g_\theta}(g_\theta^\dagger(\mathbf{a}_t, \mathbf{s}_t), \mathbf{s}_t) \right) \right)^{-\frac{1}{2}},$$

where $\mathbf{J}_{g_\theta}(\mathbf{z}_t, \mathbf{s}_t) = \frac{\partial g_\theta(\mathbf{z}_t, \mathbf{s}_t)}{\partial \mathbf{z}_t}$ is the Jacobian of g_θ with respect to $\mathbf{z}_t$, and $g_\theta^\dagger(\cdot, \mathbf{s}_t) : \mathcal{M}_\theta(\mathbf{s}_t) \rightarrow \mathcal{Z}$ is the left inverse of $g_\theta(\cdot, \mathbf{s}_t)$, such that $g_\theta^\dagger(g_\theta(\mathbf{z}_t, \mathbf{s}_t), \mathbf{s}_t) = \mathbf{z}_t$ for all $\mathbf{z}_t \in \mathcal{Z}$. The transformation $g_\theta(\cdot, \mathbf{s}_t)$ is assumed to be differentiable and injective, ensuring the existence of a left inverse and a well-defined density on the action manifold.

B. Dynamics Model

We consider two variants of the control-affine dynamics model in Equation (1): a full-rank parameterization of $\mathbf{G}$ and a low-rank approximation. Let $n = \dim(\mathcal{Y})$ and $m = \dim(\mathcal{A})$ denote the number of joints and muscles, respectively.

1) *Low-Rank Dynamics Approximation*: Due to biomechanical coupling between joints—arising from shared musculature and tendon networks—muscle commands cannot independently control all joint velocities. Consequently, the input matrix $\mathbf{G}$ is generally not full rank.

In the low-rank variant, we parameterize the input matrix as a product of factors:

$$\mathbf{G} \approx \mathbf{L}\mathbf{R},$$

where $\mathbf{L} \in \mathbb{R}^{n \times k}$ and $\mathbf{R} \in \mathbb{R}^{k \times m}$, with rank $k < n < m$.

Unless stated otherwise, we use the low-rank model with $k = 15$, which yields slightly improved performance compared to the full-rank parameterization. In Section H-B, we present the results of a sensitivity analysis in which we sweep over k .

2) *Dynamics network*: We parameterize the model using a neural network ψ that maps the state $\mathbf{s}_t$ to the dynamics parameters. In both cases, the network outputs $\mathbf{f} \in \mathbb{R}^n$ and $\tilde{\mathbf{q}} \in \mathbb{R}^n$. The covariance matrix is modeled as diagonal and constructed as $\mathbf{Q} = \text{diag}(\text{softplus}(\tilde{\mathbf{q}})^2)$. The parameterization of the input matrix $\mathbf{G} \in \mathbb{R}^{n \times m}$ depends on the variant:

- **Full-rank**: The network outputs $\mathbf{G} \in \mathbb{R}^{n \times m}$ directly. The mapping is $\psi : \mathcal{S} \rightarrow \mathbb{R}^n \times \mathbb{R}^{n \times m} \times \mathbb{R}^n$.
- **Low-rank**: The network outputs the factors $\mathbf{L}$ and $\mathbf{R}$. The mapping is $\psi : \mathcal{S} \rightarrow \mathbb{R}^n \times \mathbb{R}^{n \times k} \times \mathbb{R}^{k \times m} \times \mathbb{R}^n$. The input matrix is then reconstructed as $\mathbf{G} = \mathbf{L}\mathbf{R}$.

In both cases, since $\mathbf{Q}$ is diagonal, $\tilde{\mathbf{G}} = \mathbf{Q}^{-1/2}\mathbf{G}$ amounts to row-wise rescaling, so the row spaces of $\mathbf{G}$ and $\tilde{\mathbf{G}}$ coincide. In the low-rank case, since $\mathbf{L}$ has full column rank (as holds in practice because $k < n$), the row space of $\mathbf{G} = \mathbf{L}\mathbf{R}$ further coincides with the row space of $\mathbf{R}$. This common row space $\text{row}(\mathbf{G}(\mathbf{s}_t)) \in \text{Gr}(k, m)$ directly characterizes the optimal action subspace (Section B-D).

Training details are described below.

3) *Epistemic Manifold Uncertainty*: The closed-form construction in Theorem B.1 assumes that the dynamics parameters are known. However, in practice, these parameters must be learned from data. This creates a fundamental conflict: learning a globally accurate dynamics model requires exploring the full action space $\mathcal{A}$, yet the compositional policy $\pi_{\mathbf{a}}(\mathbf{a}_t|\mathbf{s}_t)$ confines actions to a low-dimensional manifold.

To resolve this, we represent uncertainty about the dynamics parameters in the form of a distribution $p(\phi|\mathbf{s}_t)$. This induces an implicit distribution over the optimal precoder:

$$p(\boldsymbol{\theta}^*|\mathbf{s}_t) = \langle \delta(\boldsymbol{\theta}^* - \text{PRECODERGENERATOR}(\phi)) \rangle_{p(\phi|\mathbf{s}_t)}, \quad (6)$$

where $\delta(\cdot)$ denotes the Dirac delta function. A sample from the optimal precoder distribution can be obtained by drawing a sample $\hat{\phi} \sim p(\phi|\mathbf{s}_t)$ and then evaluating $\text{PRECODERGENERATOR}(\hat{\phi})$.

Early in learning, high uncertainty in $p(\phi|\mathbf{s}_t)$ produces a broad distribution over the optimal precoder $p(\boldsymbol{\theta}^*|\mathbf{s}_t)$, enabling the agent to explore effectively across the entire action space $\mathcal{A}$. As dynamics uncertainty diminishes during learning, the induced distribution progressively concentrates around a consistent and stable manifold—a natural, uncertainty-driven transition from wide-ranging exploration to precise, synergistic control (Section G-C). This procedure is closely related to Thompson sampling [39].

We confirm empirically that the MC dropout ensemble provides meaningful estimates of epistemic uncertainty (Section G-B) and this uncertainty is crucial for training stability and asymptotic task performance (Section H-A).

4) *Training procedure*: Training is performed via maximum likelihood on transitions $(\mathbf{s}^{(i)}, \mathbf{a}^{(i)}, \mathbf{y}'^{(i)})$ sampled from the replay buffer $\mathcal{B}$, where each transition consists of a state $\mathbf{s}$, an action $\mathbf{a}$, and the next joint velocities $\mathbf{y}' = h(\mathbf{s}')$. Optimization is carried out using the ADAM optimizer [23], with transitions sampled uniformly with replacement. We use MC dropout [18] as a lightweight approximation to Bayesian inference in deep networks. Dropout induces an implicit variational posterior $q_\psi(\phi|\mathbf{s}_t)$ over the state-dependent dynamics parameters $\phi \triangleq \{\mathbf{f}, \mathbf{G}, \tilde{\mathbf{q}}\}$. Sampling different dropout masks provides a computationally efficient estimate of epistemic uncertainty, inducing a corresponding distribution over the optimal precoder (Equation (6)). Network hyperparameters are listed in Table II.

5) *Simultaneous Manifold and Task Learning*: When simultaneously learning the dynamics and SAC networks (policy and Q-functions) (Sections VIII-A and VIII-B), all networks are trained using the same shared replay buffer. The dynamics network is trained with the same training frequency, number of gradient steps, and batch size as the SAC networks (Table I). At each training iteration, we first update the dynamics network and then update the SAC networks, ensuring that the policy and Q-functions are optimized with respect to the most up-to-date action manifold.

6) *Learning Dynamics from Play*: In the play setting (Section VIII-C), we train the dynamics network using offline data collected during a dedicated play phase. We train the

model for 1M iterations, where each iteration consists of 20 gradient steps on batches of 256 transitions sampled uniformly from the play-phase buffer.

TABLE I
SAC HYPERPARAMETERS. THE SAME PARAMETERS WERE USED IN ALL EXPERIMENTS AND BY ALL MODELS.

parameter	
policy learning rate	3.0×10^{-4}
Q-function learning rate	3.0×10^{-4}
policy hidden units	[400, 300]
Q-function hidden units	[400, 300]
buffer size	300,000
batch size	256
soft update coefficient	0.02
discount factor	0.98
learning starts (steps)	3,000
train frequency (steps)	1
gradient steps	20
target entropy	$-\dim(\text{support}(\pi))$
normalize obs	true

TABLE II
DYNAMICS NETWORK HYPERPARAMETERS. *FOR SIMULTANEOUS MANIFOLD AND TASK LEARNING, THE LEARNING RATE IS LINEARLY DECAYED TO 0 OVER THE COURSE OF TRAINING, WHEREAS WHEN TRAINING FROM THE STATIC PLAY DATASET A CONSTANT LEARNING RATE IS USED.

parameter	value
learning rate*	1.0×10^{-5}
hidden units	[400, 300]
dropout rate	[0.5, 0.5]

APPENDIX F EXPERIMENTS

A. Environments

We evaluate on four contact-rich in-hand manipulation tasks from MyoSuite [11]: In **BaodingBalls**, the MyoHand must simultaneously rotate two free-floating balls over the palm, requiring precise coordination across multiple fingers. In **DieReorient**, the MyoHand must reconfigure a die to match target orientations, demanding delicate control to manipulate the object without dropping it. In **KeyTurn**, the MyoHand must rotate a key, primarily using the index finger and thumb, with intermittent contact throughout the task. In **PenTwirl**, the MyoHand must rotate a pen to a target orientation while maintaining stability through intermittent multi-finger contacts.

In the **Reorient** family of environments, the MyoHand rotates objects drawn from four geometry classes (ellipsoid, box, capsule, sphere), with either two (**Reorient8**) or twenty-five (**Reorient100**) distinct geometric configurations, toward predetermined target orientations. After training on Reorient100, zero-shot generalization is evaluated on 1000 novel objects with in-domain (**ReorientID**) or out-of-domain (**ReorientOOD**) geometric configurations.

Performance is evaluated using a *fraction solved* metric, defined as the ratio between the number of time steps

considered solved within an episode and the maximum episode length. This is generally less than 1.0 even for the optimal policy, as it takes time to reach solved states.

With the exception of BaodingBalls and DieReorient, we use the default MyoSuite environment settings for all tasks, including reward function parameters and episode lengths. BaodingBalls and DieReorient were the two tasks featured in the NeurIPS MyoChallenge 2022 competition [10], where participants were required to design custom reward functions. For these tasks, we adopt the environment configurations and reward parameters used in the Lattice paper [12]. The authors of that work were members of the team that achieved first place in the BaodingBalls challenge, making these settings a strong and well-motivated reference for evaluation.

a) *Sparse Rewards.*: In addition to these standard environments, we conduct a targeted ablation to examine the effect of reward sparsity on manifold discovery and task learning. For this experiment, we modify the dense reward function of Reorient8 to a sparse success signal: the agent receives a reward of 1.0 in solved states and 0.0 otherwise. We retain the environment’s default success definition, under which a state is considered solved when the cosine similarity between the current and goal object orientation vectors exceeds a fixed threshold of 0.95.

B. Baselines

We compare JoSEPi with current model-free RL methods in overactuated systems: SAC [20], Lattice [12], Synergistic Action Representation (SAR) [5], and a state-dependent variant of SAR (SD-SAR, Section F-C2). Like JoSEPi, SAR and SD-SAR are integrated with SAC. Lattice was originally integrated with both SAC and PPO; in the PPO variant, recurrent LSTM actor and critic networks address partial observability (e.g. unobserved object properties such as friction, mass, and shape). We implement both and denote them Lattice-SAC and Lattice-rPPO.

For fair comparison, all SAC-based methods (JoSEPi, Lattice-SAC, SAR, SD-SAR, SAC) use the same standard hyperparameters and network architecture (Table I). For Lattice-rPPO, we use the hyperparameters published in Chiappa et al. [12]. Following Haarnoja et al. [20], actions are bounded via a tanh nonlinearity, with the change-of-variables formula applied to compute log-probabilities for entropy estimation. All results are averaged across 5 random seeds.

C. Play-Based Manifold Discovery Protocol

For the experiments in Section VIII-C, we follow the protocol introduced in [5]. A SAC agent is first trained for 1M environment steps on the Reorient8 task, producing a play-phase buffer consisting of state–action–next-state transitions. SAR and JoSEPi use the same play phase to acquire an action manifold, but differ in how it is identified and later exploited during downstream learning.

1) *SAR*: In SAR, muscle activation time series generated by the play-phase policy are collected and analyzed offline.

Principal component analysis (PCA) and independent component analysis (ICA) are jointly applied to the time series data (ICAPCA), a technique commonly used in motor neuroscience to identify muscle synergies. The first 20 synergies, capturing over 80% of the variance, are used to define a low-dimensional action subspace.

The SAR policy outputs two action vectors: a low-dimensional synergy-based action and a directly parameterized high-dimensional action. These are combined via a convex combination, with weights 0.66 for the synergy-based action and 0.34 for the direct action. The code used to generate the synergies and policy of SAR is taken from the MyoSuite GitHub repository.

2) *State-Dependent SAR (SD-SAR)*: We also consider a variant of SAR (SD-SAR) that uses a state-dependent generalization of PCA to derive a synergy subspace. To motivate our approach, we first note that PCA can be framed as a representation learning problem defined over a data distribution. Specifically, consider a linear autoencoder with encoder matrix $E \in \mathbb{R}^{k \times d}$ and decoder matrix $D \in \mathbb{R}^{d \times k}$ constrained to be orthonormal ($D^\top D = I$). For a zero-mean data distribution $p(\mathbf{x})$, training this autoencoder to minimize the expected Euclidean reconstruction error $\langle \|\mathbf{x} - DE\mathbf{x}\|^2 \rangle_{p(\mathbf{x})}$ recovers an orthonormal basis for the subspace spanned by the top k principal components of the distribution, given by the columns of the optimal decoder D [8, 2]. Moreover, for any fixed orthonormal decoder, the optimal encoder matrix is given in closed form by $E = D^\top$. Substituting this optimal encoder back into the expected loss simplifies the system to the standard population PCA projection objective:

$$\min_D \langle \|\mathbf{x} - DD^\top \mathbf{x}\|^2 \rangle_{p(\mathbf{x})} \quad \text{s.t.} \quad D^\top D = I.$$

We generalize this population framework to a state-dependent setting by replacing the state-independent decoder D with a state-conditioned decoder $D(\mathbf{s}) \in \mathbb{R}^{d \times k}$. We parameterize the mapping $\xi: \mathcal{S} \rightarrow \mathbb{R}^{d \times k}$ from the state $\mathbf{s}$ to the decoder matrix via a nonlinear neural network. Specifically, the network outputs a raw matrix $\tilde{D} \in \mathbb{R}^{d \times k}$, and the orthonormal decoder $D(\mathbf{s})$ is constructed by orthonormalizing the columns of $\tilde{D}$ using the Modified Gram–Schmidt procedure.

Because $D(\mathbf{s})$ is orthonormal by construction for any state, the optimal state-dependent encoder matrix remains available in closed form as $E(\mathbf{s}) = D(\mathbf{s})^\top$. To train the network using standard stochastic optimization, we minimize the generalized reconstruction objective defined as an expectation over pairs of states and data vectors sampled from the joint distribution:

$$\mathcal{L}(\xi) = \langle \|\mathbf{x} - D(\mathbf{s})D(\mathbf{s})^\top \mathbf{x}\|^2 \rangle_{p(\mathbf{s}, \mathbf{x})}.$$

We train the decoder network using offline data consisting of concurrent states and muscle activations (the target data $\mathbf{x}$) collected during a dedicated play phase (Section VIII-C). We train the model for 50K iterations (enough for convergence), where each iteration consists of 20 gradient steps on batches of 256 state-activation pairs sampled uniformly from the play-phase buffer. Network hyperparameters are listed in Table III.

TABLE III
SD-SAR DECODER NETWORK HYPERPARAMETERS.

parameter	
learning rate	3.0×10^{-4}
hidden units	[400, 300]

3) *JoSEPi*: In JoSEPi, the play-phase experience is used to train a dynamics model that serves as the basis for action manifold discovery. The dynamics network is trained on transitions sampled from the play-phase replay buffer (see Appendix E-B for details). The dynamics model is then held fixed, inducing a static action manifold that is used to learn the Reorient100 task downstream.

APPENDIX G EMPIRICAL EVALUATION OF THE DYNAMICS MODEL

To evaluate both the accuracy of our control-affine dynamics approximation and the calibration of our epistemic uncertainty estimates, we conduct a post-hoc statistical analysis of the dynamics model. The dataset used for this analysis is created by unrolling the policy trained in each environment from Section VIII-A for 100 episodes and aggregating all transitions. During rollouts, actions are generated deterministically: the latent action is taken as the mean of the Gaussian base distribution, and the action subspace is made a fixed function of the state by deactivating dropout. At each transition $(\mathbf{s}, \mathbf{a}, \mathbf{y}')$, we draw 50 samples $\hat{\phi}^{(s)} \sim q_\psi(\phi|\mathbf{s})$ from the implicit variational posterior of the dynamics parameters. These samples were used to construct a predictive distribution over the next joint velocities:

$$p(\mathbf{y}'|\mathbf{s}, \mathbf{a}) = \frac{1}{50} \sum_{s=1}^{50} p(\mathbf{y}'|\mathbf{s}, \mathbf{a}, \hat{\phi}^{(s)}).$$

For the i -th transition in the dataset, $(\mathbf{s}^{(i)}, \mathbf{a}^{(i)}, \mathbf{y}'^{(i)})$, we defined:

- 1) the **prediction error** ($e^{(i)}$) as the mean squared error (MSE) between the expected next joint velocities $\hat{\mathbf{y}}'^{(i)} = \langle \mathbf{y}' \rangle_{p(\mathbf{y}'|\mathbf{s}^{(i)}, \mathbf{a}^{(i)})}$ and the true next joint velocities $\mathbf{y}'^{(i)}$, averaged across the J joint velocity dimensions:

$$e^{(i)} = \frac{1}{J} \left\| \hat{\mathbf{y}}'^{(i)} - \mathbf{y}'^{(i)} \right\|_2^2;$$

- 2) the **predictive uncertainty** ($u^{(i)}$) as the trace of the predictive covariance matrix divided by J :

$$u^{(i)} = \frac{1}{J} \sum_{d=1}^J \text{Var}_{p(\mathbf{y}'|\mathbf{s}^{(i)}, \mathbf{a}^{(i)})} [\mathbf{y}'_d].$$

This represents the average epistemic uncertainty of the model across the joint velocity dimensions;

- 3) the **contact switch label** ($c^{(i)}$) as a binary indicator where $c^{(i)} = 1$ if at least one fingertip switches between contact and non-contact states during the transition, and $c^{(i)} = 0$ otherwise.

Computing these metrics across all transitions yielded a final evaluation pool of 273,940 tuples, $(e^{(i)}, u^{(i)}, c^{(i)})$, which serves as the foundation for the following analyses.

A. Accuracy of the Control-Affine Approximation

Contact-rich manipulation tasks are characterized by intermittent contact events that introduce local discontinuities into the system dynamics. To assess how well the control-affine approximation in Equation (1) handles these events, we stratify the evaluation pool based on the contact switch label c .

As shown in Table IV, the mean prediction error (mean $\pm$ s.e.m.) is $e = 0.152 \pm 0.001$ at transitions with a contact switch compared to 0.149 ± 0.001 without—a small difference in means relative to the within-group variability (s.t.d. = 0.465 and 0.353, respectively). A linear mixed-effects model with a fixed effect of the contact switch label and random intercepts for each experimental run (environment–seed combination) confirms that while the large sample size (273,520 transitions) yields a highly significant effect ($\beta = 0.035$, s.e. = 0.002, $z = 17.01$, $p < 1 \times 10^{-16}$), the increase in prediction error is modest. Taken together, these results suggest that the model predictions remain largely stable across task phases, with limited degradation during contact switching events, verifying that the control-affine assumption serves as a reliable basis for optimizing JoSE. We hypothesize that the contact discontinuities are gracefully accommodated by the control-affine formulation as it allows for arbitrary nonlinearities with respect to the state. We also speculate that by incorporating richer contact-related observations, such as tactile or force feedback, prediction errors during contact switches could be reduced further.

TABLE IV
DYNAMICS MODEL PREDICTION ERROR GROUPED BY CONTACT SWITCH LABEL. DATA POOLED ACROSS ALL ENVIRONMENTS, TRAINING RUNS, EPISODES, AND TIME STEPS (273,520 TRANSITIONS), STRATIFIED BY WHETHER A CONTACT SWITCH OCCURRED ($c = 1$) OR NOT ($c = 0$).

switch	mean	s.e.m.	count	s.t.d.
yes	0.152	0.001	148,021	0.465
no	0.149	0.001	125,499	0.353

B. MC Dropout Uncertainty Calibration

Using the same evaluation dataset, we examine whether the epistemic uncertainty represented by the MC dropout ensemble is meaningful in practice. Specifically, we calculate the Pearson correlation coefficient between the predictive uncertainty u and the prediction error e across all tuples.

As shown in Table V, a consistent positive correlation ($r \in [0.50, 0.66]$) is observed across all environments, indicating that the dropout ensemble provides a meaningful signal about regions of high model uncertainty. This correlation is remarkably high given the lightweight nature of the training procedure (single-mask dropout is used at training time) and the high-dimensional nature of the dynamics network input (the state spaces of the MyoHand and each rigid-body object are 85- and 13-dimensional, respectively).

TABLE V
CORRELATION BETWEEN MC DROPOUT PREDICTIVE UNCERTAINTY AND PREDICTION ERROR. DATA SHOW PEARSON CORRELATION COEFFICIENT (MEAN $\pm$ S.E.M. ACROSS 5 TRAINING RUNS).

Environment	Pearson r
DieReorient	0.658 ± 0.032
BaodingBalls	0.501 ± 0.088
KeyTurn	0.549 ± 0.015
PenTwirl	0.507 ± 0.033

It is important to note that the predictive uncertainty specifically isolates epistemic uncertainty rather than total prediction error. Systematic errors stemming from model mismatch are irreducible and do not inflate MC dropout variance: under persistent model bias, all dropout samples converge to the same biased prediction, exhibiting zero epistemic variance. The correlation observed therefore reflects reducible uncertainty from limited data—precisely the type of metric required to drive intrinsically motivated exploration.

C. Action Subspace Uncertainty

As described in Section E-B3, the distribution over the optimal action subspace should concentrate as dynamics uncertainty diminishes during learning. Here we verify this directly by examining how subspace variability evolves throughout training. For each MC dropout sample $s = 1, \dots, 50$, the dynamics network produces a realization $G^{(s)}(s_t)$ whose row space $\text{row}(G^{(s)}(s_t)) \in \text{Gr}(k, M)$ is the optimal action subspace for that realization (Sections B-D and E-B2). To quantify how tightly an ensemble of 50 action subspaces concentrate on $\text{Gr}(k, M)$, we compute the Fréchet variance (subspace uncertainty) under the projection-matrix embedding (5).

We repeat the data collection procedure described above at 8 equally spaced points throughout training (as well as at model initialization). In each state encountered, we compute both the Fréchet variance and the predictive uncertainty (Figure 7).

Both quantities decrease rapidly early in training, indicating that the dropout ensemble quickly converges to consistent predictions and action subspaces (Figure 7). This aids uncertainty-guided exploration: high uncertainty early in training induces a broad distribution over the optimal action subspace, giving the agent access to the full action space. As uncertainty reduces, the agent converges to a consistent and reliable pattern of muscle coordination aligned with its learned dynamics. Notably, KeyTurn diverges from the other tasks later in training: subspace uncertainty increases slightly after its initial drop rather than remaining low, suggesting the dynamics model does not converge to a consistent action subspace. This may explain why KeyTurn is also the only task that shows a marginal improvement in learning efficiency when dropout is removed (Figure 8).

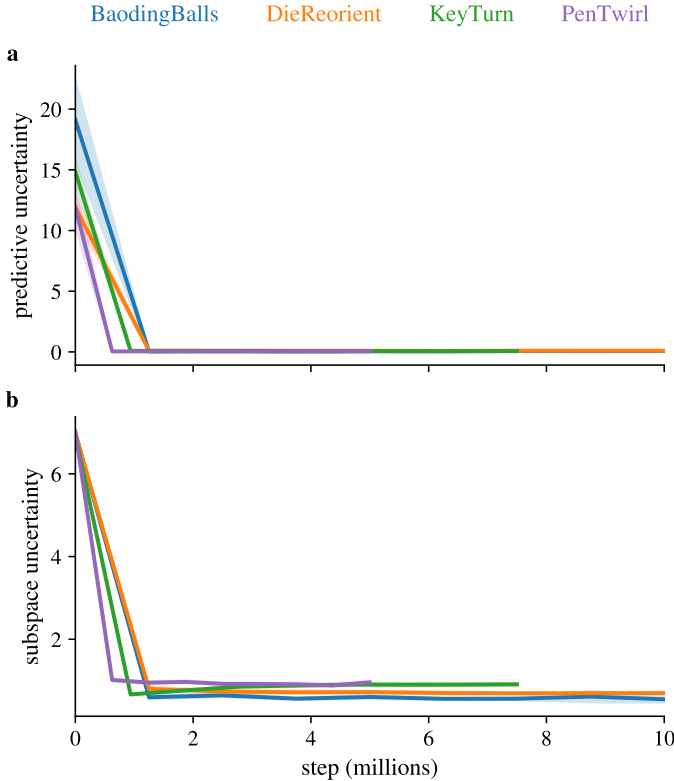


Fig. 7. **Epistemic uncertainty evolution during training.** (a) Predictive uncertainty (variance of dynamics predictions across 50 MC dropout samples) averaged over all states as a function of training step (mean $\pm$ s.e.m. across 5 training runs). Step 0 corresponds to a randomly initialized policy and dynamics network. (b) Subspace uncertainty (Fréchet variance of the action subspace across 50 MC dropout samples) averaged over all states (mean $\pm$ s.e.m. across 5 training runs).

APPENDIX H ABLATIONS AND SENSITIVITY ANALYSES

A. Effect of Dynamics Uncertainty

To isolate the impact of representing dynamics uncertainty, we set the MC dropout rate to 0.0 (thereby learning a single dynamics model rather than an ensemble) and repeat the benchmark evaluation from Section VIII-A. We find that modeling epistemic uncertainty via MC dropout substantially enhances both asymptotic performance and training stability (Figure 8).

B. Effect of Input Matrix Factorization Rank

The factorization rank k of the input matrix $\mathbf{G}$ (Section E-B1) directly determines the dimensionality of both the latent space and the learned action manifold. To characterize the sensitivity of JoSEPi to the manifold dimensionality, we sweep over $k \in \{9, 12, 15, 18, 23\}$ and repeat the benchmark evaluation from Section VIII-A. As shown in Table VI, performance remains stable across a wide range of dimensionalities ($k \geq 12$). The first meaningful performance drop occurs at $k = 9$; this degradation is most pronounced on BaodingBalls, a particularly challenging task that requires in-hand manipulation

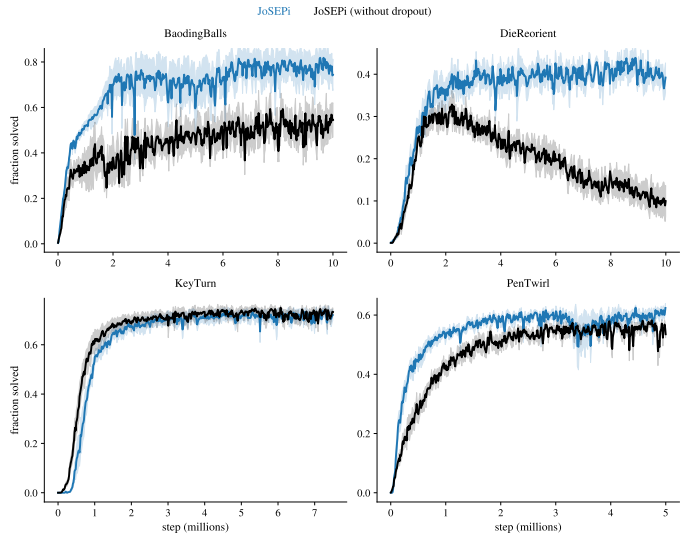


Fig. 8. **MC dropout ablation.** Performance of JoSEPi with different per-layer MC dropout rates in the dynamics model ($p = 0.5$, blue; $p = 0.0$, black) on BaodingBalls, DieReorient, KeyTurn, and PenTwirl. Data show solved fraction over training (mean $\pm$ s.e.m. across 5 random seeds).

of multiple objects. Although we currently treat k as a user-specified hyperparameter, future extensions could automate this selection, for example, by inspecting the singular values of the full-rank dynamics matrix $\mathbf{G}$ to determine its effective rank.

TABLE VI
EFFECT OF INPUT MATRIX FACTORIZATION RANK ON PERFORMANCE.
DATA SHOW SOLVED FRACTION AT 5M STEPS OF TRAINING (MEAN $\pm$ S.E.M. ACROSS 5 RANDOM SEEDS) AS A FUNCTION OF INPUT MATRIX RANK k .

k	BaodingBalls	DieReorient	KeyTurn	PenTwirl
23	0.572 ± 0.052	0.355 ± 0.023	0.688 ± 0.023	0.609 ± 0.016
18	0.687 ± 0.067	0.344 ± 0.020	0.750 ± 0.013	0.600 ± 0.011
15	0.611 ± 0.102	0.380 ± 0.031	0.717 ± 0.021	0.624 ± 0.015
12	0.700 ± 0.056	0.357 ± 0.015	0.711 ± 0.014	0.619 ± 0.009
9	0.378 ± 0.075	0.284 ± 0.030	0.641 ± 0.029	0.556 ± 0.012