

Taehyun Cho

✉ talium@cml.snu.ac.kr ☎ +82-10-5519-3089 🌐 talium0713.github.io

Research Experience

- Aug 2026 – Aug 2027 **Vector Institute**, Toronto, Canada
Vector Distinguished Postdoctoral Fellow
- Co-affiliation: **Fields Institute** — Principles of Intelligence Postdoctoral Fellow.
 - Host PI: Prof. Tim G. J. Rudner, University of Toronto.
- Mar 2026 – Present **Seoul National University**, Seoul, South Korea
Postdoctoral Fellow
- Sejong Science Fellowship, National Research Foundation of Korea (NRF).
 - Adviser: Prof. Jungwoo Lee.
- Dec 2024 – May 2025 **LG AI Research**, Seoul, South Korea
Research Intern, Superintelligence Lab
- RLHF Squad Team — developing process reward models.
- Mar 2020 – Feb 2026 **Seoul National University**, Seoul, South Korea
Graduate Researcher, Cognitive Machine Learning Lab
- Adviser: Prof. Jungwoo Lee.

Education

- Mar 2022 – 2026 **Seoul National University**, Seoul, South Korea
Ph.D. in Electrical and Computer Engineering
- Mar 2020 – 2022 **Seoul National University**, Seoul, South Korea
M.S. in Electrical and Computer Engineering
- Mar 2013 – 2020 **Korea University**, Seoul, South Korea
B.S. in Mathematics

Research Interests

My research focuses on *sequential decision-making under uncertainty*, particularly in the context of *human feedback*. I have extensively studied *distributional reinforcement learning (DistRL)*, *reinforcement learning from human feedback (RLHF)*, and *regret analysis*, aiming to bridge theory and practice.

My long-term goal is to build *human-aligned, socially-aware agents* — systems that reason about people, navigate the dynamics of human interaction, and act on their preferences, values, and decisions. I draw inspiration from how humans make decisions. By seeking to understand the cognitive mechanisms underlying human choice and mathematically modeling the structure that governs interaction, I aim to develop both theoretical insights and practical algorithms for robust decision-making under uncertainty.

- **Broad:** Deep Learning, Reinforcement Learning, Stochastic Optimization.
- **Specific:** Distributional Reinforcement Learning, Regret Analysis, Reinforcement Learning from Human Feedback.

Publications

International Conferences

1. **[Spotlight, Top 2.2%]** Suhwan Kim*, **Taehyun Cho***, Geonhyeong Kim, Yujin Kim, Youngsoo Jang, Moontae Lee, Jungwoo Lee. “A Regret Minimization Framework on Preference Learning in Large Language Models.” *ICML 2026*, Seoul, Korea.
2. Jung Min Lee, Dohyeok Lee, Seokhun Ju, **Taehyun Cho**, Jin Woo Koo, Li Zhao, Sangwoo Hong, Jungwoo Lee. “MVP-LAM: Learning Action-Centric Latent Action via Cross-Viewpoint Reconstruction.” *ICML 2026*, Seoul, Korea.
3. Kyungjae Lee, Dohyeong Kim, **Taehyun Cho**, Chaeyeon Kim, Yunkyoung Ko, Seungyub Han, Seokhun Ju, Dohyeok Lee, Sungbin Lim. “Pareto Optimal Risk-Agnostic Distributional Bandits with Heavy-Tail Rewards.” *NeurIPS 2025*, San Diego, USA.
4. **[Spotlight, Top 2.6%]** **Taehyun Cho***, Seokhun Ju*, Seungyub Han, Dohyeong Kim, Kyungjae Lee, Jungwoo Lee. “Policy-based Preference Learning: Is Preference Enough for RLHF?” *ICML 2025*, Vancouver, Canada.
5. **Taehyun Cho**, Seungyub Han, Seokhun Ju, Dohyeong Kim, Kyungjae Lee, Jungwoo Lee. “Bellman Unbiasedness: Toward Provably Efficient Distributional Reinforcement Learning with General Value Function Approximation.” *ICML 2025*, Vancouver, Canada.
6. Dohyeong Kim, **Taehyun Cho**, Seungyub Han, Hojun Chung, Kyungjae Lee, Songhwai Oh. “Spectral-Risk Safe Reinforcement Learning with Convergence Guarantees.” *NeurIPS 2024*, Vancouver, Canada.
7. **Taehyun Cho**, Seungyub Han, Heesoo Lee, Kyungjae Lee, Jungwoo Lee. “Pitfall of Optimism: Distributional Reinforcement Learning by Randomizing Risk Criterion.” *NeurIPS 2023*, New Orleans, USA.
8. Dohyeok Lee, Seungyub Han, **Taehyun Cho**, Jungwoo Lee. “SPQR: Controlling Q-ensemble Independence with Spiked Random Model for Reinforcement Learning.” *NeurIPS 2023*, New Orleans, USA.
9. Seungyub Han, Yeongmo Kim, **Taehyun Cho**, Jungwoo Lee. “On the Convergence of Continual Learning with Adaptive Methods.” *UAI 2023*, Pittsburgh, USA.
10. Seungyub Han, Yeongmo Kim, **Taehyun Cho**, Jungwoo Lee. “Adaptive Methods for Nonconvex Continual Learning.” *NeurIPS 2022 Workshop*, New Orleans, USA.
11. **Taehyun Cho**, Seungyub Han, Heesoo Lee, Kyungjae Lee, Jungwoo Lee. “Perturbed Quantile Regression for Distributional Reinforcement Learning.” *NeurIPS 2022 Workshop*, New Orleans, USA.
12. Sangwoo Hong, Heecheol Yang, Youngseok Yoon, **Taehyun Cho**, Jungwoo Lee. “Chebyshev Polynomial Codes: Task Entanglement-based Coding for Distributed Matrix Multiplication.” *ICML 2021*, Vienna, Austria.

International Journals

1. Seokhun Ju, Seungyub Han, **Taehyun Cho**, Jungwoo Lee, Taeyoung Lee, Minkyung Kim, Jinho Ahn. “Learning Graph Based Individual Intrinsic Reward for Multi-Agent Reinforcement Learning.” *ICT Express*, Aug. 2025.
2. Heasung Kim, **Taehyun Cho**, Jungwoo Lee, Wonjae Shin, H. Vincent Poor. “Optimized Shallow Neural Networks for Sum-Rate Maximization in Energy Harvesting Downlink Multiuser NOMA Systems.” *IEEE Journal on Selected Areas in Communications*, vol. 39, no. 4, pp. 982–997, Apr. 2021.
3. Heasung Kim, **Taehyun Cho**, Jungwoo Lee, Wonjae Shin, H. Vincent Poor. “An Efficient Neural Network Architecture for Rate Maximization in Energy Harvesting Downlink Channels.” *IEEE ISIT 2020*, Los Angeles, USA, June 21–26, 2020.

Ph.D. Dissertation

1. **Taehyun Cho**. “A Distributional Perspective on Human-Aligned Decision Making under Uncertainty.” *Seoul National University*, Feb. 2026.
Committee: Jungwoo Lee (adviser), Songhwai Oh, Kyomin Jung, Taesup Moon, Kyungjae Lee.

Preprints & In Progress

1. *(Under submission)* Seungyub Han, **Taehyun Cho**, Dohyeong Kim, Kyungjae Lee, Jungwoo Lee. “When to Truncate Traces: Stochastic Truncation for Multi-Step Off-Policy RL.”
2. *(Under submission)* Kukyoung Jang, **Taehyun Cho**, Junrui Zhang, Ping Xu, Kyungjae Lee. “Probabilistic Smoothing with Ratio-Monotone Transforms for Global Optimization.”
3. **Taehyun Cho**, Suhwan Kim, Seungyub Han, Kyungjae Lee, Jungwoo Lee. “The Unique Off-Policy Permissibility of KL-divergence in Direct Preference Optimization.”
4. **Taehyun Cho**, Suhwan Kim, Seungyub Han, Seokhun Ju, Dohyeong Kim, Kyungjae Lee, Jungwoo Lee. “An Axiomatization of Process Score Model: Your Process-level Feedback is Not a Reward.”
5. Seungyub Han*, **Taehyun Cho***, Suhwan Kim, Seokhun Ju, Dohyeong Kim, Kyungjae Lee, Jungwoo Lee. “Off-Policy Trust Region Preference Optimization.”

Teaching Experience

Spring 2022	Teaching Assistant — <i>Dept. of Electrical and Computer Engineering, SNU</i> • Introduction to Reinforcement Learning.
Spring 2021	Teaching Assistant — <i>Dept. of Electrical and Computer Engineering, SNU</i> • Introduction to Reinforcement Learning.

Talks

July 2026	Cortiq Summit 2026: Agentic AI — <i>Coex, Seoul</i> “A Regret Minimization Framework on Preference Learning in Large Language Models.”
July 2026	RIKEN-AIP · Vector · NAIRL Workshop — <i>KAIST Seocho AICT</i> “A Regret Minimization Framework on Preference Learning in Large Language Models.”
Apr 2026	Vector Visitor Research Talk — <i>Vector Institute</i> “Rethinking Human Feedback: A Regret Minimization Perspective on Preference Learning.”
Apr 2026	SNU AI Summit — <i>Seoul National University</i> “Policy-labeled Preference Learning: Is Preference Enough for RLHF?”
May 2025	LG AI Research Seminar — <i>LG AI Research</i> “Policy Optimization with Process Score in LRMs.”
Dec 2024	LG AI Research Seminar — <i>LG AI Research</i> “Policy-labeled Preference Learning: Is Preference Enough for RLHF?”
Aug 2023	LG Tech Talk — <i>LG AI Research</i> “Pitfall of Optimism: Distributional Reinforcement Learning with Randomized Risk Criterion.”
May 2023	AiIS Spring Retreat Program — <i>Seoul National University</i> “Pitfall of Optimism: Distributional Reinforcement Learning with Randomized Risk Criterion.”

Awards and Grants

July 2026 – 2027	Research Grant at Grantmaking.ai — <i>Grantmaking.ai</i> Project: “Faithful Preference Learning: Cognitively-Aligned Post-Training” (USD 50K)
July 2026	SK Inc. AX Best Research Talk Award — <i>Cortiq Summit 2026: Agentic AI</i>

Aug 2026 – 2027	Vector Distinguished Postdoctoral Fellowship — <i>Vector Institute</i>
Aug 2026 – 2027	Principles of Intelligence Postdoctoral Fellowship — <i>Fields Institute</i>
May 2026	INMC Young Researcher Award — <i>Institute of New Media and Communications (INMC)</i>
May 2026	Gold Reviewer Award — <i>International Conference on Machine Learning (ICML)</i>
Mar 2026 – 2031	Sejong Science Fellowship — <i>National Research Foundation of Korea (NRF)</i> Project: “Distributional Regret Analysis for Human-Aligned Interactive Agentic AI under High Uncertainty” (KRW 600M)
Feb 2026	Distinguished Dissertation Award — <i>Seoul National University (SNU)</i>
Mar 2020 – 2026	Brain Korea 21 Plus Scholarship — <i>Seoul National University (SNU)</i>
Jan 2023	Certificate of Commendation — <i>Center for Applied Research in Artificial Intelligence (CARAI)</i>

Other Experience

2022 – Present	Conference & Workshop Reviewer • ICML 2022–2026 — <i>Gold Reviewer Award (Top 25%)</i> • Workshop @ ICML 2026: RLxF — Reinforcement Learning from World Feedback • NeurIPS 2022–2025 • ICLR 2023–2025 • AAAI 2023 • The Journal of Korean Institute of Communications and Information Sciences, 2026
Jun 2018 – Present	Philosophy Club — <i>Sunday Salon</i> • Discussions on Baccalauréat topics
May 2016 – 2018	Military Service — <i>Republic of Korea Air Force</i> • Served at 3rd Flight Training Wing

Skills

Programming: Python, NumPy, PyTorch.

Languages: Korean (Native), English (Fluent).