

David Abel

Staff Research Scientist, Google DeepMind
Honorary Fellow, University of Edinburgh

📍 *Edinburgh, UK*
🏠 david-abel.github.io
✉ dmabel@deepmind.com
✉ david.abel@ed.ac.uk

EDUCATION

- 2020 **Ph.D in Computer Science**, *Brown University*, Providence, RI.
Advisor: Michael L. Littman.
Committee: George Konidaris, Stefanie Tellex, Peter Stone, Will Dabney.
Thesis: A Theory of Abstraction in Reinforcement Learning.
- 2019 **M.A. in Philosophy**, *Brown University*, Providence, RI.
Advisor: Joshua Schechter.
Thesis: Concepts in Bounded Rationality: Perspectives from Reinforcement Learning.
- 2015 **M.S. in Computer Science**, *Brown University*, Providence, RI.
Advisor: Stefanie Tellex.
Thesis: Learning to Plan in Complex Stochastic Domains.
- 2013 **B.A. in Computer Science & Philosophy**, *Carleton College*, Northfield, MN.
Advisors: David Liben-Nowell (CS), Anna Moltchanova (Philosophy).
Thesis: Toward the Defense of Hypercomputation.

PROFESSIONAL EXPERIENCE

- 2025–Present **Staff Research Scientist**, *Google DeepMind*.
- 2024–Present **Honorary Fellow**, *University of Edinburgh*, School of Informatics.
- 2022–2025 **Senior Research Scientist**, *Google DeepMind*.
- 2020–2022 **Research Scientist**, *Google DeepMind*.
- Summer 2019 **Research Intern**, *Google DeepMind*, Hosted by Dr. Will Dabney.
- Summer 2016 **Visiting Researcher**, *University of Oxford*, Hosted by Dr. Owain Evans.
- Summer 2015 **Research Intern**, *Microsoft Research NYC*, Hosted by Dr. Fernando Diaz.

SELECTED AWARDS

Principal's Medal, University of Edinburgh, 2025.
Outstanding Paper Award, "On the Expressivity of Markov Reward", NeurIPS 2021.
Presidential Award for Excellence in Teaching, Brown University.
Runner-up, AAAI/ACM SIGAI Doctoral Dissertation Award, 2020.
9 × Top Reviewer Award, ICML / NeurIPS / ICLR / AISTATS.
Long-listed, Supervisor of the Year, University of Edinburgh.

TEACHING AND ADVISING

- Student Ben Sanati, *Ph.D Student*, University of Edinburgh, 2024–Present.
- Supervision Samuel Garcin, *Ph.D Student*, University of Edinburgh, 2025–Present.
Massimiliano Tamborski, *Ph.D Student*, University of Cambridge, 2025–Present.

Massimiliano Tamborski, *Masters Student*, University of Edinburgh, 2024–2025.

Post-doctoral Supervision Alper Demir, University of Edinburgh, 2024–2025.

External Examiner Jacob Beck, *Oxford University*, Ph.D Student of Shimon Whiteson.
Riccardo Zamboni, *Politecnico di Milano*, Ph.D student of Marcello Restelli.
Theodore Moskovitz, *UCL*, Ph.D Student of Maneesh Sahani.

Instructor Reinforcement Learning, *200 students*, University of Edinburgh.
A First Byte of Computer Science, *110 students*, Brown University.
Artificial Intelligence and Society, *20 students*, Summer at Brown.

Guest Lecturer Beyond the Markov Property, Cambridge.
Three Dogmas of Reinforcement Learning, Alberta.
Continual Reinforcement Learning, Stanford.
On Research (2022), MIT.
On Research (2021a, 2021b), Harvard University.
On Research (2020), Harvard University.
Reinforcement Learning, Brown University.
Data Science, Brown University.
A First Byte of Computer Science, Brown University.
Data Science Summer Program, Microsoft Research NYC Summer School.

Teaching Assistant A First Byte of Computer Science, Brown University.
Artificial Intelligence, Brown University.
Data Structures, Carleton College.
Logic, Carleton College.
Intro to Computer Science $\times 3$, Carleton College.

PUBLICATIONS

Journal Papers

The Syntax and Semantics of Goals

David Abel, Mark K. Ho.

Topics in Cognitive Science, 2026.

Optimizing Return Distributions with Distributional Dynamic Programming.

Bernardo Ávila Pires, Mark Rowland, Diana Borsa, Zhaohan Daniel Guo, Khimya Khetarpal, André Barreto, David Abel, Rémi Munos, Will Dabney.
JMLR 2025.

People Construct Simplified Mental Representations to Plan.

Mark K. Ho, David Abel, Carlos G. Correa, Michael L. Littman, Jonathan D. Cohen, Thomas L. Griffiths.

Nature 2022.

The Value of Abstraction.

Mark K. Ho, [David Abel](#), Thomas L. Griffiths, Michael L. Littman.
Current Opinions in Behavioral Sciences 2019.

Conference Papers

Forgetting is Everywhere.

Ben Sanati, Thomas L. Lee, Trevor McInroe, Aidan Scannell, Nikolay Malkin, [David Abel](#), Amos Storkey.
CoLLAs 2026.

Fairness over Equality: Correcting Social Incentives in Asymmetric Sequential Social Dilemmas.

Alper Demir, Hüseyin Aydın, Kale-ab Abebe Tessera, [David Abel](#), Stefano V. Albrecht.
AAMAS 2026.

Probing Dec-POMDP Reasoning in Cooperative MARL.

Kale-ab Tessera, Leonard Hinckeldey, Riccardo Zamboni, [David Abel](#), Amos Storkey.
AAMAS 2026.

Discovering Coordinated Joint Options via Inter-Agent Relative Dynamics.

Raul D Steleac, Mohan Sridharan, [David Abel](#).
ICLR 2026.

Plasticity as the Mirror of Empowerment.

[David Abel](#), Michael Bowling, André Barreto, Will Dabney, Shi Dong, Steven Hansen, Anna Harutyunyan, Khimya Khetarpal, Clare Lyle, Razvan Pascanu, Georgios Piliouras, Doina Precup, Jonathan Richens, Mark Rowland, Tom Schaul, Satinder Singh.
NeurIPS 2025.

Enhancing Tactile-based Reinforcement Learning for Robotic Control.

Elle Miller, Trevor McInroe, [David Abel](#), Oisín Mac Aodha, Sethu Vijayakumar.
NeurIPS 2025.

Skill-Driven Neurosymbolic State Abstractions.

Alper Ahmetoglu, Steven James, Cameron Allen, Sam Lobel, [David Abel](#), George Konidaris.
NeurIPS 2025.

General Agents Contain World Models.

Jonathan Richens, [David Abel](#), Alexis Bellot, Tom Everitt.
ICML 2025.

A Black Swan Hypothesis: The Role of Human Irrationality in AI Safety.

Hyunin Lee, Chanwoo Park, [David Abel](#), Ming Jin.
ICLR 2025.

Studying the Interplay Between the Actor and Critic Representations in Reinforcement Learning.

Samuel Garcin, Trevor McInroe, Pablo Samuel Castro, Christopher G. Lucas, [David Abel](#), Prakash Panangaden, Stefano V. Albrecht.
ICLR 2025.

Three Dogmas of Reinforcement Learning.

[David Abel](#), Mark K. Ho, Anna Harutyunyan.

RLC 2024.

Pragmatic Feature Preferences: Learning Reward-Relevant Preferences from Human Feedback.

Andi Peng, Yuying Sun, Tianmin Shu, David Abel.

ICML 2024.

A Definition of Continual Reinforcement Learning.

David Abel, André Barreto, Benjamin Van Roy, Doina Precup, Hado van Hasselt, Satinder Singh.

NeurIPS 2023.

Settling the Reward Hypothesis.

Michael Bowling, John D. Martin, David Abel, Will Dabney.

ICML 2023.

Meta-Gradients in Non-Stationary Environments.

Jelena Luketina, Sebastian Flennerhag, Yannick Schroecker, David Abel, Tom Zahavy, Satinder Singh.

CoLLAs 2022.

On the Expressivity of Markov Reward

David Abel, Will Dabney, Anna Harutyunyan, Mark K. Ho, Michael L. Littman, Doina Precup, Satinder Singh.

NeurIPS 2021 (**Outstanding Paper Award**).

Revisiting Peng's $Q(\lambda)$ for Modern Reinforcement Learning

Tadashi Kozuno, Yunhao Tang, Mark Rowland, Rémi Munos, Steven Kapturowski, Will Dabney, Michal Valko, David Abel.

ICML 2021.

Lipschitz Lifelong Reinforcement Learning

Erwan Lecarpentier, David Abel, Kavosh Asadi, Yuu Jinnai, Emmanuel Rachelson, Michael L. Littman.

AAAI 2021.

What can I do here? A Theory of Affordances in Reinforcement Learning

Khimya Khetarpal, Zafarali Ahmed, Gheorghe Comanici, David Abel, Doina Precup.

ICML 2020.

Value Preserving State-Action Abstractions.

David Abel, Nathan Umbanhowar, Khimya Khetarpal, Dilip Arumugam, Doina Precup, Michael L. Littman.

AISTATS 2020.

The Efficiency of Human Cognition Reflects Planned Use of Information Processing.

Mark K. Ho, David Abel, Jonathan D. Cohen, Michael L. Littman, Thomas L. Griffiths.

AAAI 2020.

The Expected-Length Model of Options.

David Abel^{*}, John Winder^{*}, Marie desJardins, Michael L. Littman

IJCAI 2019.

Finding Options that Minimize Planning Time.

Yuu Jinnai, David Abel, D. Ellis Hershkowitz, Michael L. Littman, George Konidaris.
ICML 2019.

Discovering Options for Exploration by Minimizing Cover Time.

Yuu Jinnai, Jee Won Park, David Abel, George Konidaris.
ICML 2019.

State Abstraction as Compression in Apprenticeship Learning.

David Abel, Dilip Arumugam, Kavosh Asadi, Yuu Jinnai, Michael L. Littman, Lawson L.S. Wong.
AAAI 2019.

State Abstractions for Lifelong Reinforcement Learning.

David Abel, Dilip Arumugam, Lucas Lehnert, Michael L. Littman.
ICML 2018.

Policy and Value Transfer in Lifelong Reinforcement Learning

David Abel*, Yuu Jinnai*, Yue Guo, George Konidaris, Michael L. Littman.
ICML 2018.

Bandit-Based Solar Panel Control.

David Abel, Edward C. Williams, Stephen Brawner, Emily Reif, Michael L. Littman.
IAAI 2018.

Near Optimal Behavior via Approximate State Abstraction.

David Abel*, D. Ellis Hershkowitz*, Michael L. Littman.
ICML 2016.

Goal-Based Action Priors.

David Abel, D. Ellis Hershkowitz, Gabriel Barth-Maron, Stephen Brawner, Kevin O'Farrell, James MacGlashan, Stefanie Tellex.
ICAPS 2015.

Workshops, Symposia, and Extended Abstracts

Memory Allocation in Resource-Constrained Reinforcement Learning

Massimiliano Tamborski, David Abel.
RLDM 2025.

Agency is Frame-Dependent

David Abel, André Barreto, Michael Bowling, Will Dabney, Shi Dong, Steven Hansen, Anna Harutyunyan, Khimya Khetarpal, Clare Lyle, Razvan Pascanu, Georgios Piliouras, Doina Precup, Jonathan Richens, Mark Rowland, Tom Schaul, Satinder Singh.
RLDM 2025.

Expressing Non-Markov Reward to a Markov Agent

David Abel, André Barreto, Michael Bowling, Will Dabney, Steven Hansen, Anna Harutyunyan, Mark K. Ho, Ramana Kumar, Michael L. Littman, Doina Precup, Satinder Singh.
RLDM 2022.

On the Expressivity of Markov Reward (Extended Abstract)

David Abel, Will Dabney, Anna Harutyunyan, Mark K. Ho, Michael L. Littman, Doina Precup,

Satinder Singh.
IJCAI 2022.

Bad-Policy Density: A Measure of Reinforcement Learning Hardness.
David Abel, Cameron Allen, Dilip Arumugam, D. Ellis Hershkowitz, Michael L. Littman, Lawson L.S. Wong.
ICML Workshop on Reinforcement Learning Theory, 2021.

Skill Discovery with Well-Defined Objectives.
Yuu Jinnai, David Abel, Jee Won Park, D. Ellis Hershkowitz, Michael L. Littman, and George Konidaris.
ICLR Workshop on Structure and Priors in Reinforcement Learning, 2019.

simple_rl: Reproducible Reinforcement Learning in Python.
David Abel.
ICLR Workshop on Reproducibility in Machine Learning, 2019.

A Theory of State Abstraction for Reinforcement Learning.
David Abel.
AAAI Doctoral Consortium 2019.

Toward Good Abstractions for Lifelong Learning.
David Abel, Dilip Arumugam, Lucas Lehnert, Michael L. Littman.
NeurIPS Workshop on Hierarchical Reinforcement Learning, 2017.

Toward Improving Solar Panel Efficiency using Reinforcement Learning.
David Abel, Emily Reif, Edward C. Williams, Michael L. Littman.
EnvirolInfo 2017, early version at RLDM 2017.

Agent-Agnostic Human-in-the-Loop Reinforcement Learning.
David Abel, John Salvatier, Andreas Stuhlmüller, Owain Evans.
NeurIPS Workshop on the Future of Interactive Learning Machines, 2016.

Exploratory Gradient Boosting for Reinforcement Learning in Complex Domains.
David Abel, Alekh Agarwal, Fernando Diaz, Akshay Krishnamurthy, Robert Schapire.
ICML Workshop on Abstraction in Reinforcement Learning, 2016.

Reinforcement Learning as a Framework for Ethical Decision Making.
David Abel, James MacGlashan, Michael L. Littman.
AAAI Workshop on AI, Society, & Ethics, 2016.

Affordances as Transferrable Knowledge for Planning Agents.
Gabriel Barth-Maron, David Abel, James MacGlashan, Stefanie Tellex.
AAAI Symposium on Knowledge and Skill Transfer, 2014.

Toward Affordance-Aware Planning.
David Abel, Gabriel Barth-Maron, James MacGlashan, Stefanie Tellex.
RSS Workshop on Affordances in Vision for Cognitive Robotics, 2014.

- 2026 *The Foundations of Agency*, Philosophy and Reinforcement Learning Symposium.
The Foundations of Agency, Reinforcement Learning and Philosophy Workshop.
Savage, McCarthy, and Carse: A New Braid, RL in Big Worlds Workshop at RLC.
Tracing Learning through Agents, Continual RL Workshop at RLC.
Agents, Planning, and Generalization, GenPlan Workshop at ICAPS.
Three Dogmas of Reinforcement Learning, Berkeley Sensorimotor AI Journal Club.
Plasticity as the Mirror of Empowerment, Berkeley Sensorimotor AI Journal Club.
- 2025 *Two Elements of Agency*, London Institute of Philosophy.
The Science of Agency, London School of Economics.
The Science of Agency, Chevening AI Scholars.
The Science of Agency, UoE Student Society.
Plasticity as the Mirror of Empowerment, Bath.
Plasticity as the Mirror of Empowerment, Imperial.
Plasticity as the Mirror of Empowerment, Mila.
Plasticity as the Mirror of Empowerment, Edinburgh Cognitive Science Conference.
Plasticity as the Mirror of Empowerment, Max Planck.
The Reward Hypothesis, Edinburgh Institute for Perception, Action, and Behavior.
The Science of Agency, Base Rates Podcast.
The Science of Agency, TalkRL Podcast.
- 2024 *The Reward Hypothesis*, Gatsby's Computational Neuroscience Unit at UCL.
Three Dogmas of Reinforcement Learning, Penn State.
Three Dogmas of Reinforcement Learning, Purdue.
The Reward Hypothesis, RTX AI Seminar.
Toward a Science of Agency, RL Safety Workshop at RLC.
Task Specification for Robotics, Panel at RSS Workshop.
The Limits of Reward Functions, Panel at RLC.
The Reward Hypothesis, RSS Workshop on Task Specification.
Agency and Abstraction, Oregon State University.
Elements of Agency, Mila.
A Definition of Continual Reinforcement Learning, University of Edinburgh.
The Reward Hypothesis, University of Edinburgh.
- 2023 *Three Dogmas of Reinforcement Learning*, ICML Workshop on Interactive Learning.
On Reinforcement Learning, AI Podden.
The Reward Hypothesis, Principles of Intelligence Podcast.
On Research, Panel at New in ML Workshop at NeurIPS.
AI and the Reward Hypothesis, Fidelity AI Lab.
- 2022 *On the Expressivity of Markov Reward*, Brown.
On the Expressivity of Markov Reward, Imperial College London.
On the Expressivity of Markov Reward, U. of Witwatersrand.
On the Expressivity of Markov Reward, Australian Reinforcement Learning Group.
On the Expressivity of Markov Reward, UCSB.
Agency as Reward Maximization, Comp. Bio. Conference at ICL.
- 2021 *On the Expressivity of Markov Reward*, Mila.
Abstraction in Reinforcement Learning, Purdue.
- 2020 *Abstraction in Reinforcement Learning*, Michigan.
Abstraction in Reinforcement Learning, MSR Redmond.

Abstraction in Reinforcement Learning, NYU.
Abstraction in Reinforcement Learning, Australian Reinforcement Learning Group.

- 2019 *Abstraction and Meta Reinforcement Learning*, NeurIPS Workshop on Meta-Learning.
Abstraction in Reinforcement Learning, UC Berkeley.
Abstraction in Reinforcement Learning, U. Mass..
Abstraction in Reinforcement Learning, UT Austin.
Abstraction in Reinforcement Learning, Oxford CS.
Abstraction in Reinforcement Learning, Oxford Philosophy.
- 2018 *State Abstraction in Reinforcement Learning*, UCLA.
State Abstraction in Reinforcement Learning, USC.
State Abstraction in Reinforcement Learning, UCSD.
State Abstraction in Reinforcement Learning, CU Boulder.
State Abstraction in Reinforcement Learning, Princeton.
State Abstraction in Reinforcement Learning, Oregon State.
State Abstraction in Reinforcement Learning, Baidu Sunnyvale.
Bandit-Based Solar Panel Control, Computational Sustainability Seminar.
- 2017 *How Artificial Intelligence Should Model the World*, Research Matters.
Abstraction and Lifelong Reinforcement Learning, Carnegie Mellon University.
Artificial Intelligence; How Intelligent, and How Soon?, Nordea Markets NYC.
Abstraction and Reinforcement Learning, MSR Redmond.

SERVICE

- Reviewing 2026: JMLR, ICLR, ICML, RLC, Workshops.
2025: AAAI, AISTATS, JMLR, NeurIPS, RLDM, RLC, RSS, Workshops..
2024: ACM, AISTATS, CoLLAs, ICML, JMLR, NeurIPS, OpenMind, RLC, TMLR.
2023: AAAI, AISTATS, CoLLAs, ICML, JMLR, Nature, NeurIPS, TMLR, Workshops.
2022: AISTATS, CoLLAs, EWRL, ICML, ICLR, NeurIPS, RLDM, JMLR, TMLR.
2021: AIJ, ICML, IEEE, JMLR, NeurIPS, Workshops.
2020: AAAI, AIJ, JAIR, JMLR, MLJ, ICML, ICLR, IEEE, NeurIPS, Workshops.
2019: AAAI, ACM, JMLR, ICLR, ICML, MLJ, NeurIPS, RLDM, Workshops.
2018: JMLR, ICML, IEEE, MLJ, NeurIPS, Workshops.
2017: ICAPS, IEEE.
- Service JMLR Editorial Board, 2020, 2021, 2022, 2023, 2024, 2025.
Senior Area Chair, RLC: 2025, 2026.
Area Chair, NeurIPS: 2023, 2024, 2025.
Area Chair, ICML: 2026.
Area Chair, AAAI: 2023.
Senior PC Member, RLC: 2024.
Senior PC Member, CoLLAs: 2022, 2023, 2024.
PC for NeurIPS Workshop: 2021.
PC for IJCAI Workshop: 2021.
PC for NeurIPS Workshop: 2020.
PC for ICML Workshop: 2020.
- Events Program Chair, RLDM: 2027.
Keynote Chair, RLC: 2026.
Workshop Chair, RLDM: 2025.

Associate Program Chair, CoLLAs: 2025.
 Co-Chair of RLC Workshop on Finding the Frame: 2024, 2025, 2026.
 Co-Chair of RLDM Workshop on RL as a Model of Agency: 2022.
 Co-Chair of ICLR Social on Philosophy of AGI: 2021.
 Co-Chair of RLDM Workshop on Moral Decision Making: 2019.

Brown CS Faculty-Grad Student Liaison: 2018-2019
 Ph.D Recruiting Czar, 2018.
 Grad Student Social Czar, 2016-2017.
 Ph.D Mentorship Program Coordinator, 2017-2019.
 Co-organizer, Reinforcement Learning Reading Group, 2017-2018.

AWARDS

General Principal's Medal, University of Edinburgh, *University-wide award for outstanding service to the community*, 2025.
 Outstanding Paper Award, NeurIPS 2021, *On the Expressivity of Markov Reward*.
 Runner-up for the AAAI/ACM SIGAI Doctoral Dissertation Award.
 Long-listed for Supervisor of the Year 2025, University of Edinburgh.
 Nominated for the Victor Lesser Distinguished Dissertation Award.
 Top Reviewer, AISTATS 2022; ICML 2018-2021; NeurIPS 2019, 2020, 2025; ICLR 2026.

Brown Presidential Award for Excellence in Teaching, *University wide award presented to four graduate students per year for outstanding pedagogical achievement*.
 University Open Graduate Fellowship for Masters in Philosophy.
 Great Teaching Assistant Award, *Artificial Intelligence*.
 Great Teaching Assistant Award, *A First Byte of Computer Science*.

Carleton Distinction in Major and Thesis, *Philosophy*.
 College Distinction in Major and Capstone, *Computer Science*.