

# RAST-MoE-RL

A Regime-Aware Spatio-Temporal Mixture-of-Experts  
framework for deep reinforcement learning in ride-hailing

---

Yuhan Tang<sup>\*1</sup> · Kangxin Cui<sup>\*2</sup> · Jung Ho Park<sup>\*3</sup> · Yibo Zhao<sup>\*1</sup> · Xuan Jiang<sup>1</sup> · Haoze He<sup>4</sup> · Jiangbo Yu<sup>5</sup> · Haris Koutsopoulos<sup>6</sup> · Jinhua Zhao<sup>1</sup>

<sup>1</sup>Massachusetts Institute of Technology · <sup>2</sup>Tongji University · <sup>3</sup>University of California, Berkeley · <sup>4</sup>Carnegie Mellon University ·

<sup>5</sup>McGill University · <sup>6</sup>Northeastern University

\*equal contribution

# The Matching Decision

Every request poses one question, tens of thousands of times an hour: match this passenger to a driver now, or wait for a better one to appear?

## Match now

Assign to the nearest driver immediately

|        |                                 |
|--------|---------------------------------|
| Match  | <b>Fast</b>                     |
| PICKUP | <b>Slow — driver may be far</b> |

vs

## Wait

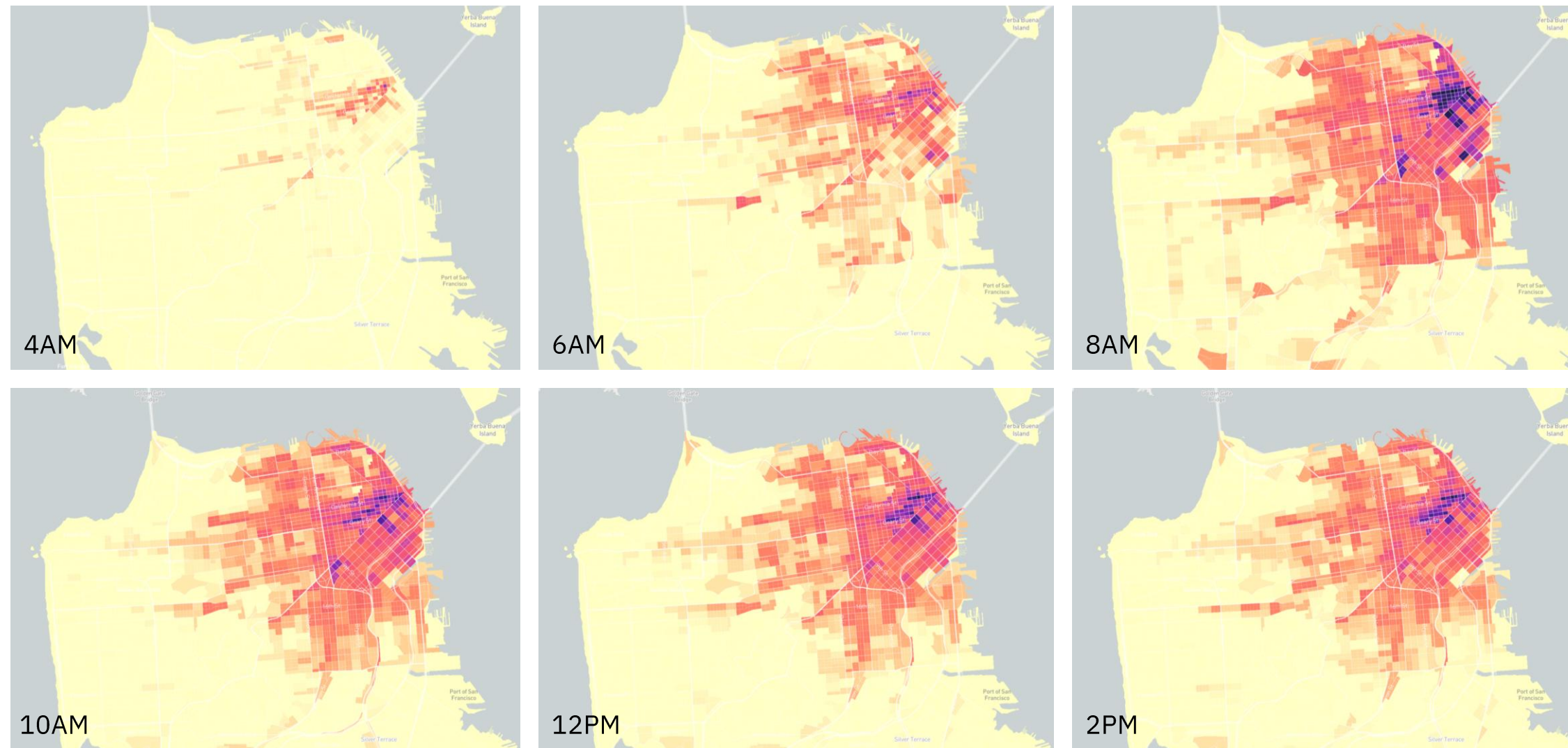
Batch requests, hold for a better match

|        |                               |
|--------|-------------------------------|
| Match  | <b>Slow — passenger waits</b> |
| PICKUP | <b>Fast — closer driver</b>   |

## PROBLEM STATEMENT

# Challenges

The city runs in distinct regimes that shift continuously — and the right strategy in one actively fails in another.



Heatmap of spatiotemporal demand patterns in San Francisco (SFTCA)

**HOLD TOO LONG** Penalize pickup too hard and the agent **holds requests indefinitely** — matching waits explode past 20 minutes.

**MATCH TOO FAST** Penalize batching too hard and it **matches everything instantly** — pickup waits past 25 minutes.

# Limitations of Prior Work

Methods using centralized coordinator do not account for non-stationary regimes

| APPROACH                                                     | HOW IT WORKS                                               | LIMITATION                   |
|--------------------------------------------------------------|------------------------------------------------------------|------------------------------|
| <div><b>Heuristics</b><br/>Fixed windows · time-of-day</div> | Simple, interpretable rules tuned for average conditions.  | Does not adapt in real time. |
| <div><b>Monolithic RL</b><br/>Single policy network</div>    | One set of weights compresses every regime into one model. | Does not generalize well     |
| <div><b>RAST-MoE-RL</b><br/>Our work</div>                   | Experts that each handle a distinct urban regime.          |                              |

# Integration of MoE Architecture

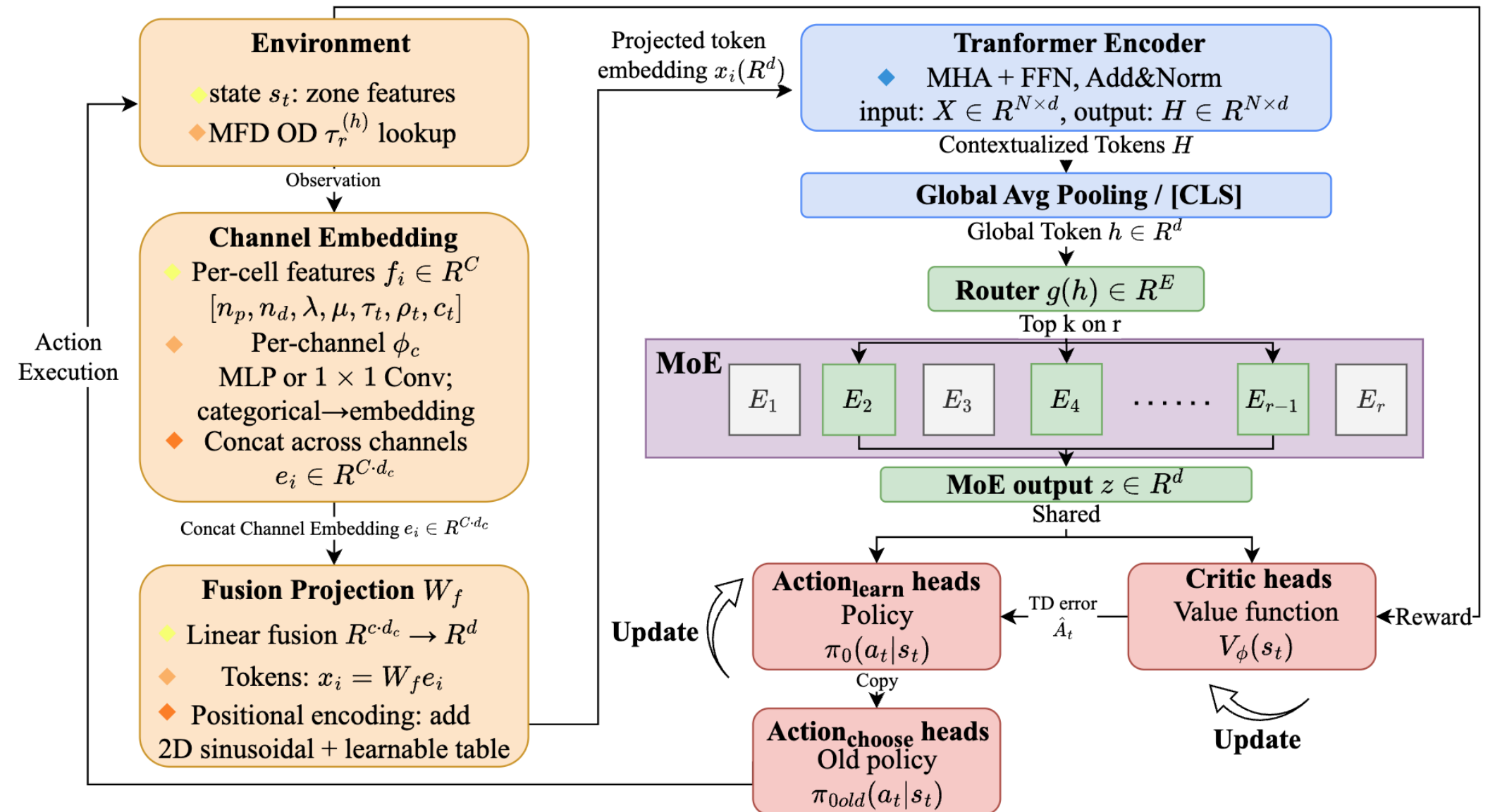
Instead of one network handling every situation, a learned gate routes each state to a few **specialist experts**. Different experts are activated at different operational regimes such as rush hour and late at night. This specialization emerges from training.

**12M**

parameters in the encoder

**4/16**

experts active per decision



# Reward Design

---

Our reward design prevents reward hacking and promotes temporal balancing of match and pickup waits.

- **Reward balancing match time and pickup time**

$$r_t = -\left[c_m \Delta W_t^{\text{match}} + c_p(t) \Delta W_t^{\text{pickup}}\right] - \lambda_t (g(s_t, a_t) - \alpha).$$

- **Service quality measure soft constraint**

$$g(s_t, a_t) \leq \alpha$$

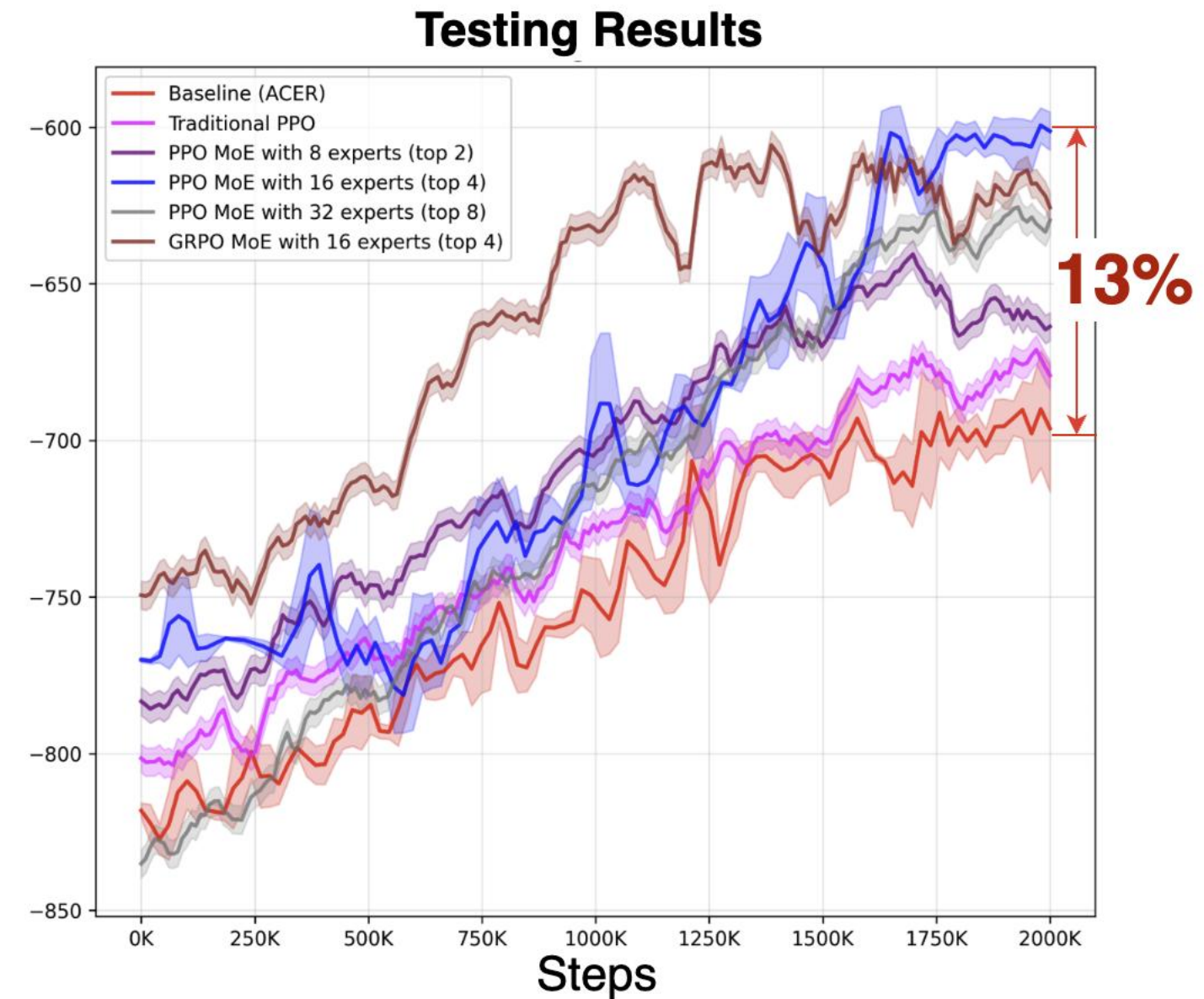
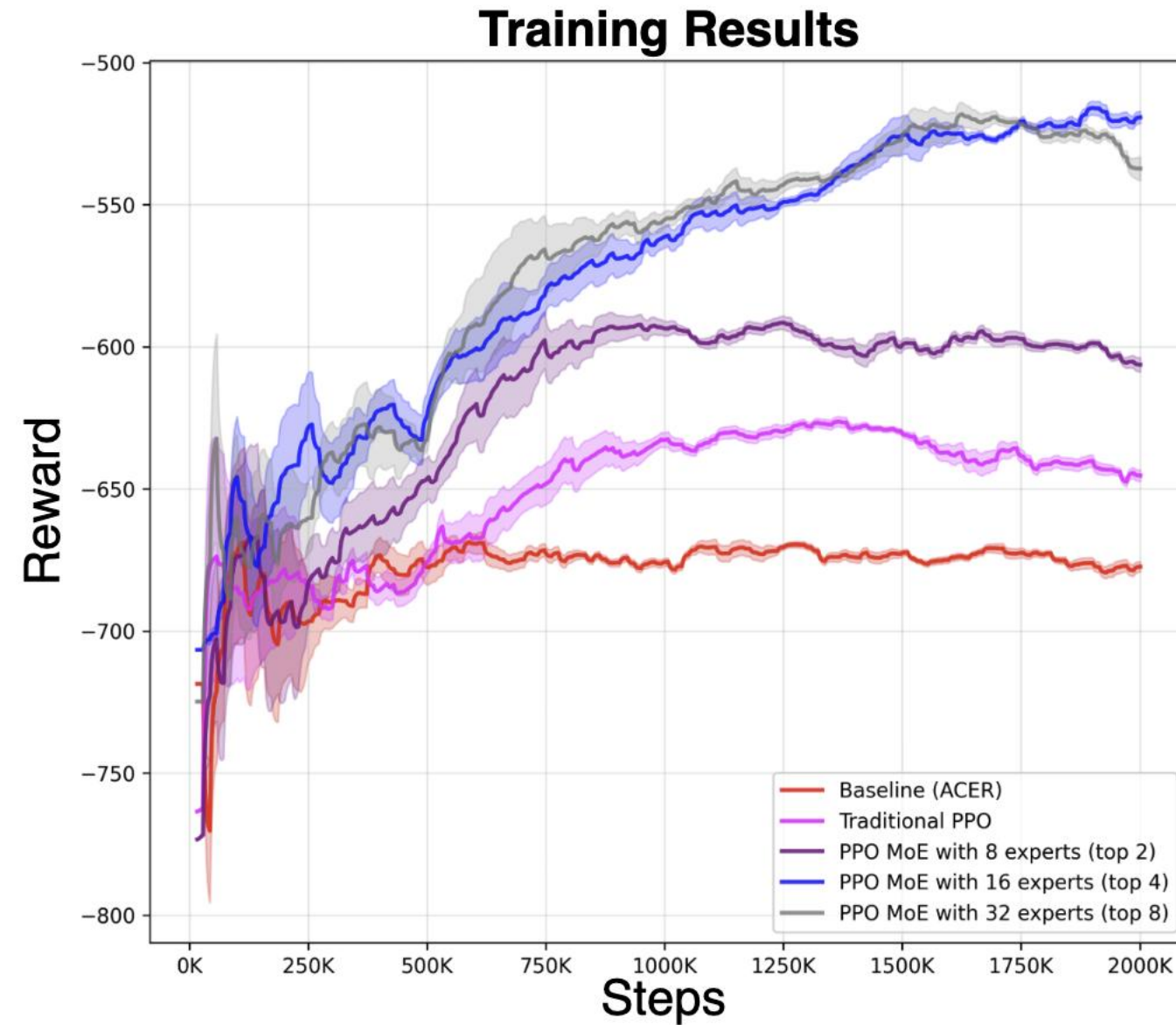
- **Adaptive multiplier update**

$$\lambda_{t+1} \leftarrow \left[\lambda_t + \xi (g(s_t, a_t) - \alpha)\right]_+$$

## RESULTS

# Training and Testing Report

Comparison of training and testing among baseline algorithms show the advantage of PPO MoE with 16 experts with top=4 expert activation.



RESULTS

# Performance Comparison

On real San Francisco TNC data, RAST-MoE-RL beats every baseline — heuristic and RL alike — reducing matching and pickup waits simultaneously.

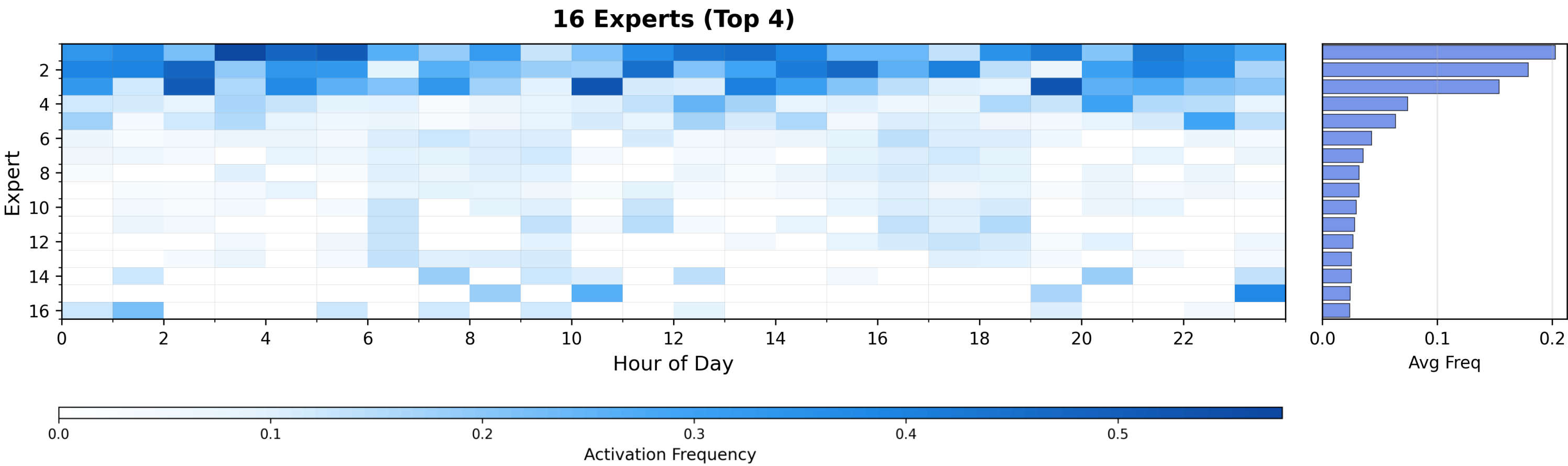
| METHOD                                                 | MATCHING WAIT          | PICKUP WAIT            |
|--------------------------------------------------------|------------------------|------------------------|
|                                                        | SECONDS · LOWER BETTER | SECONDS · LOWER BETTER |
| RAST-MoE-RL <div>OURS</div><br>MoE 16/4 · 12M · PPO    | 182                    | 523                    |
| Dense Transformer<br>12M · PPO                         | 201                    | 569                    |
| Plain MLP<br>12M · PPO                                 | 208                    | 577                    |
| Best heuristic<br>Conditioned on zone-level congestion | 203                    | 784                    |

–10% matching · –15% pickup · +13% reward. Pickup wait drops ~33% below the strongest heuristic.

RESULTS

# Evidence of Specialization by Regime

The 16 experts do not share the load equally. Each fires in distinct patterns across the 24-hour cycle, matching the city's operational regimes — rush hour, off-peak, late night.

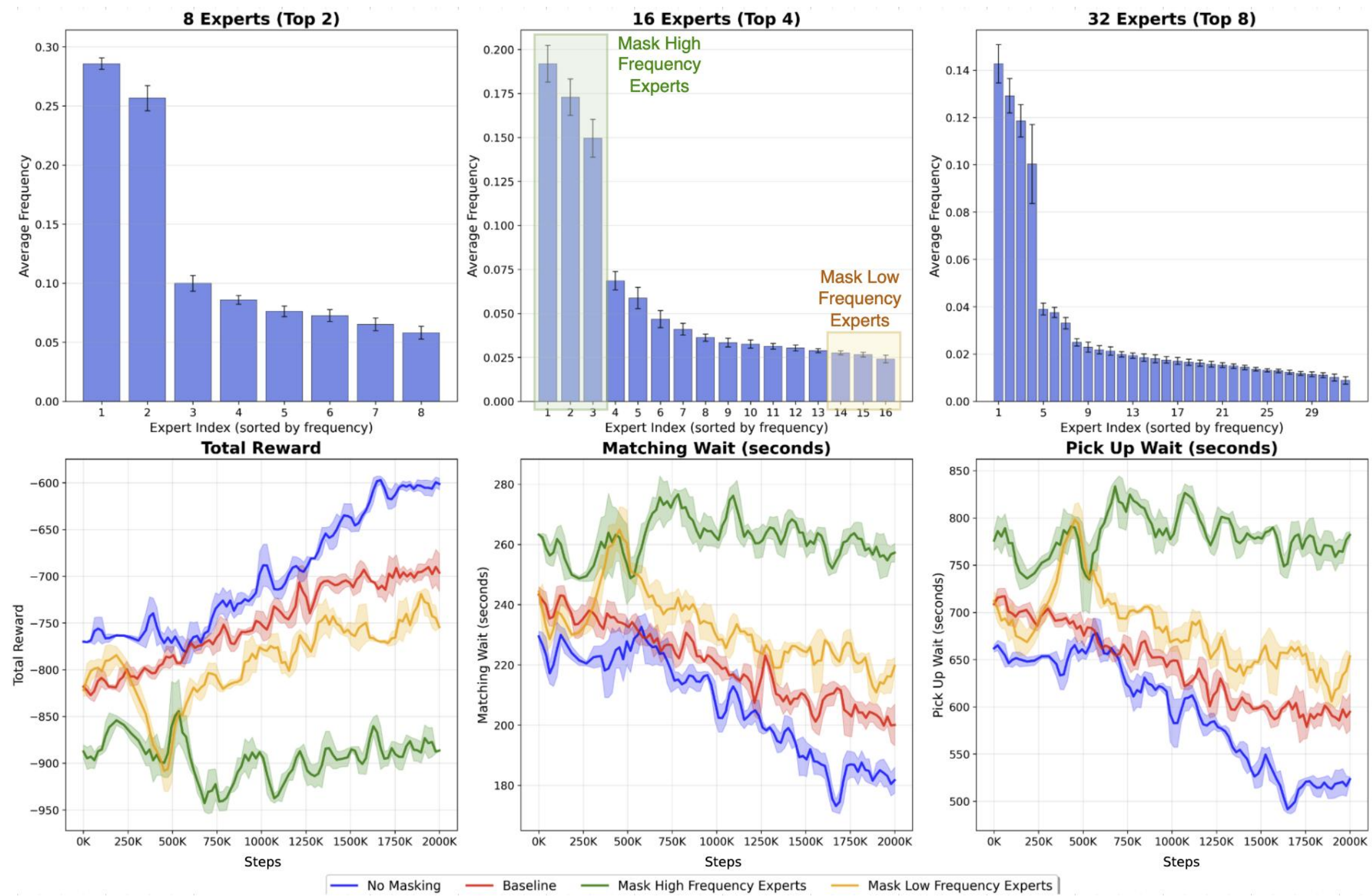


Expert activation heatmap

RESULTS

# Importance of Individual Experts

Are the experts genuinely specializing? Every expert turns out to be irreplaceable.



Expert masking

RESULTS

# Cross-City Generalization

Application of policy trained in SF city to San Diego and Alameda County reveals good generalization of the policy.

| Alameda County    |                       |                     | San Diego         |                      |                   |
|-------------------|-----------------------|---------------------|-------------------|----------------------|-------------------|
|                   | SF → ALA<br>ZERO-SHOT | ALA → ALA<br>NATIVE |                   | SF → SD<br>ZERO-SHOT | SD → SD<br>NATIVE |
| Matching wait (s) | 196                   | 191                 | Matching wait (s) | 205                  | 208               |
| Pickup wait (s)   | 595                   | 578                 | Pickup wait (s)   | 626                  | 617               |

The SF-trained policy lands within **1.5–3%** of a model trained natively in each city. The regime-aware representation captures something [general](#) about urban ride-hailing.

## CONCLUSION

# Three Contributions

### 01 A regime-aware environment.

Adaptive delayed matching as a spatio-temporal MDP, with a physics-based travel-time surrogate fast enough for millions of rollouts.

---

### 02 An adaptive reward for model training.

An online Lagrangian multiplier tightens service constraints in real time, stabilizing long-horizon training.

---

### 03 A compact 12M-parameter MoE encoder.

Experts specialize across regimes, beat every baseline on both waits, and generalize zero-shot across cities.

**Regime awareness** — in how you model the world and represent its state likely extends to any large-scale spatio-temporal resource allocation problem.

CORRESPONDENCE

[xuanj@mit.edu](mailto:xuanj@mit.edu)

[haozeh@cs.cmu.edu](mailto:haozeh@cs.cmu.edu)