

Submitted to *Transportation Science*

Inductive Transportation Origin-Destination Demand Forecasting via Bayesian Destination Choice Modeling

Xiaoxu Chen

Department of Logistics and Operations Management, HEC Montréal, Montreal, QC H3T 2A7, Canada;
xiao.xu.chen@hec.ca

Yafeng Yin

Department of Civil and Environmental Engineering, University of Michigan, Ann Arbor, Michigan 48109, United States;
yafeng@umich.edu

Lijun Sun

Department of Civil Engineering, McGill University, Montreal, QC H3A 0C3, Canada;
lijun.sun@mcgill.ca

Appendix A: Proof of Proposition 1 (Origin-Factor Gradient)

We derive the likelihood gradient directly in matrix form. For time t , write

$$\mathbf{G}_t = \sum_{s=1}^R w_{ts} \mathbf{u}_s \mathbf{v}_s^\top, \quad \mathbf{H}_t = C(\mathbf{G}_t),$$

where $C(\cdot)$ denotes row-wise centering over feasible destinations. Let $\mathbf{\Lambda}_t$ be the matrix of destination-choice probabilities and define

$$\mathbf{Y}_t^+ = \sum_{m=1}^M \mathbf{Y}_t^m, \quad \mathbf{n}_t^+ = \sum_{m=1}^M (N_{1t}^m, \dots, N_{St}^m)^\top, \\ \mathbf{R}_t = \mathbf{Y}_t^+ - \text{diag}(\mathbf{n}_t^+) \mathbf{\Lambda}_t.$$

The row-wise multinomial softmax likelihood satisfies

$$d\mathcal{L}_{\text{like},t} = \langle \mathbf{R}_t, d\mathbf{H}_t \rangle.$$

Because each row of $\mathbf{R}_t$ sums to zero, $\mathbf{R}_t$ is orthogonal to the row-wise constants introduced by the centering operator. Hence

$$\langle \mathbf{R}_t, d\mathbf{H}_t \rangle = \langle \mathbf{R}_t, d\mathbf{G}_t \rangle.$$

For an origin-factor perturbation $d\mathbf{u}_r$,

$$d\mathbf{G}_t = w_{tr} d\mathbf{u}_r \mathbf{v}_r^\top.$$

Therefore,

$$d\mathcal{L}_{\text{like}} = \sum_{t=1}^T w_{tr} \langle \mathbf{R}_t, d\mathbf{u}_r \mathbf{v}_r^\top \rangle = \left(\sum_{t=1}^T w_{tr} \mathbf{R}_t \mathbf{v}_r \right)^\top d\mathbf{u}_r.$$

Thus,

$$\nabla_{\mathbf{u}_r} \mathcal{L}_{\text{like}} = \sum_{t=1}^T w_{tr} \mathbf{R}_t \mathbf{v}_r.$$

Adding the Gaussian-prior differential

$$d\left(-\frac{1}{2}\mathbf{u}_r^\top \mathbf{K}_u^{-1} \mathbf{u}_r\right) = -\left(\mathbf{K}_u^{-1} \mathbf{u}_r\right)^\top d\mathbf{u}_r$$

gives the full posterior gradient.

Appendix B: Proof of Proposition 2 (Destination-Factor Gradient)

For a destination-factor perturbation $d\mathbf{v}_r$,

$$d\mathbf{G}_t = w_{tr} \mathbf{u}_r d\mathbf{v}_r^\top.$$

Using the same residual identity,

$$d\mathcal{L}_{\text{like}} = \sum_{t=1}^T w_{tr} \langle \mathbf{R}_t, \mathbf{u}_r d\mathbf{v}_r^\top \rangle = \left(\sum_{t=1}^T w_{tr} (\mathbf{R}_t)^\top \mathbf{u}_r \right)^\top d\mathbf{v}_r.$$

Therefore,

$$\nabla_{\mathbf{v}_r} \mathcal{L}_{\text{like}} = \sum_{t=1}^T w_{tr} (\mathbf{R}_t)^\top \mathbf{u}_r.$$

Adding

$$d\left(-\frac{1}{2}\mathbf{v}_r^\top \mathbf{K}_v^{-1} \mathbf{v}_r\right) = -\left(\mathbf{K}_v^{-1} \mathbf{v}_r\right)^\top d\mathbf{v}_r$$

gives the full posterior gradient.

Appendix C: Matrix Proof of Proposition 3 (Temporal-Factor Gradient)

For an arbitrary temporal-factor perturbation $d\mathbf{w}_r = (dw_{1r}, \dots, dw_{Tr})^\top$,

$$d\mathbf{G}_t = \mathbf{u}_r \mathbf{v}_r^\top dw_{tr}.$$

Then

$$d\mathcal{L}_{\text{like}} = \sum_{t=1}^T \langle \mathbf{R}_t, \mathbf{u}_r \mathbf{v}_r^\top \rangle dw_{tr} = \sum_{t=1}^T \mathbf{u}_r^\top \mathbf{R}_t \mathbf{v}_r dw_{tr}.$$

Hence

$$[\nabla_{\mathbf{w}_r} \mathcal{L}_{\text{like}}]_t = \mathbf{u}_r^\top \mathbf{R}_t \mathbf{v}_r.$$

The prior differential

$$d\left(-\frac{1}{2}\mathbf{w}_r^\top \mathbf{K}_w^{-1} \mathbf{w}_r\right) = -\left(\mathbf{K}_w^{-1} \mathbf{w}_r\right)^\top d\mathbf{w}_r$$

completes the matrix derivation of the posterior gradient.

Appendix D: Proof of Proposition 4 (Log-Transform Jacobian)

We use the standard formula for a change of variables in a probability density. Let θ_j be a random variable with density $p(\theta_j)$ and let $\phi_j = h(\theta_j)$ be a one-to-one transformation, where $h(\theta_j) = \log(\theta_j)$. The inverse transformation is $\theta_j = h^{-1}(\phi_j) = \exp(\phi_j)$.

The probability density function for ϕ_j , which we denote $\tilde{p}(\phi_j)$, is given by:

$$\tilde{p}(\phi_j) = p(h^{-1}(\phi_j)) \left| \frac{dh^{-1}(\phi_j)}{d\phi_j} \right|. \quad (1)$$

We apply this formula to our specific case. The original density is the conditional posterior $p(\theta_j | \cdot)$. We can get the $\tilde{p}(\phi_j | \cdot)$ as:

$$\tilde{p}(\phi_j | \cdot) = p(\theta_j = \exp(\phi_j) | \cdot) \left| \frac{d}{d\phi_j} \exp(\phi_j) \right|. \quad (2)$$

The derivative of the inverse transformation is:

$$\frac{d}{d\phi_j} \exp(\phi_j) = \exp(\phi_j). \quad (3)$$

Since $\theta_j > 0$, the transformation $\phi_j = \log(\theta_j)$ is defined, and its inverse $\exp(\phi_j)$ is always positive. Therefore, the absolute value is redundant:

$$\left| \frac{d}{d\phi_j} \exp(\phi_j) \right| = |\exp(\phi_j)| = \exp(\phi_j). \quad (4)$$

Substituting this back, we get the density for ϕ_j :

$$\tilde{p}(\phi_j | \cdot) = p(\theta_j = \exp(\phi_j) | \cdot) \exp(\phi_j). \quad (5)$$

To obtain the target function for the slice sampler, we take the logarithm of this density:

$$\begin{aligned} \mathcal{L}(\phi_j) &= \log \tilde{p}(\phi_j | \cdot) \\ &= \log(p(\theta_j = \exp(\phi_j) | \cdot) \exp(\phi_j)) \\ &= \log p(\theta_K(\phi_j) | \cdot) + \phi_j. \end{aligned} \quad (6)$$

Appendix E: Proof of Corollary 1

Fix an existing origin o and time interval t . Let $\mathcal{J}_o$ denote the original set of feasible destinations, and suppose that a new destination $s \notin \mathcal{J}_o$, with $s \neq o$, is added to the network. The expanded feasible destination set is therefore $\mathcal{J}_o^+ = \mathcal{J}_o \cup \{s\}$. The result is conditional on the latent utilities $\{G_{odt} : d \in \mathcal{J}_o\}$ and G_{ost} , which are held fixed when the new destination is introduced.

Before the network expansion, define the mean utility over the original feasible destinations as

$$\bar{G}_{ot} = \frac{1}{|\mathcal{J}_o|} \sum_{k \in \mathcal{J}_o} G_{okt},$$

and let the corresponding centered utility be

$$g_{odt} = G_{odt} - \bar{G}_{ot}, \quad d \in \mathcal{J}_o.$$

The original destination-choice probability is

$$\lambda_{odt} = \frac{\exp(g_{odt})}{\sum_{k \in \mathcal{J}_o} \exp(g_{okt})}.$$

Because the same centering constant is subtracted from every feasible utility, it cancels from the softmax numerator and denominator. Hence,

$$\lambda_{odt} = \frac{\exp(G_{odt} - \bar{G}_{ot})}{\sum_{k \in \mathcal{J}_o} \exp(G_{okt} - \bar{G}_{ot})} = \frac{\exp(G_{odt})}{\sum_{k \in \mathcal{J}_o} \exp(G_{okt})}.$$

Define

$$Z_{ot} = \sum_{k \in \mathcal{J}_o} \exp(G_{okt}), \quad a_{ost} = \exp(G_{ost}).$$

It follows that

$$\lambda_{odt} = \frac{\exp(G_{odt})}{Z_{ot}}, \quad d \in \mathcal{J}_o.$$

After destination s is added, the mean utility over the expanded feasible set becomes

$$\bar{G}_{ot}^+ = \frac{1}{|\mathcal{J}_o| + 1} \left(\sum_{k \in \mathcal{J}_o} G_{okt} + G_{ost} \right).$$

The expanded-network probabilities are obtained by applying the softmax transformation to the centered utilities $G_{odt} - \bar{G}_{ot}^+$. Again, the common centering constant cancels. Therefore, the probability assigned to the new destination is

$$\begin{aligned} \lambda_{ost}^+ &= \frac{\exp(G_{ost} - \bar{G}_{ot}^+)}{\sum_{k \in \mathcal{J}_o} \exp(G_{okt} - \bar{G}_{ot}^+) + \exp(G_{ost} - \bar{G}_{ot}^+)} \\ &= \frac{\exp(G_{ost})}{\sum_{k \in \mathcal{J}_o} \exp(G_{okt}) + \exp(G_{ost})} \\ &= \frac{a_{ost}}{Z_{ot} + a_{ost}}. \end{aligned}$$

Similarly, for every existing destination $d \in \mathcal{J}_o$,

$$\begin{aligned} \lambda_{odt}^+ &= \frac{\exp(G_{odt})}{Z_{ot} + a_{ost}} \\ &= \frac{Z_{ot}}{Z_{ot} + a_{ost}} \frac{\exp(G_{odt})}{Z_{ot}} \\ &= \frac{Z_{ot}}{Z_{ot} + a_{ost}} \lambda_{odt}. \end{aligned}$$

Since

$$1 - \lambda_{ost}^+ = 1 - \frac{a_{ost}}{Z_{ot} + a_{ost}} = \frac{Z_{ot}}{Z_{ot} + a_{ost}},$$

we obtain

$$\lambda_{odt}^+ = (1 - \lambda_{ost}^+) \lambda_{odt}, \quad d \in \mathcal{J}_o.$$

Thus, all probabilities assigned to the existing destinations are reduced by the same proportional factor.

Finally, for any two existing destinations $d, k \in \mathcal{J}_o$,

$$\frac{\lambda_{odt}^+}{\lambda_{okt}^+} = \frac{(1 - \lambda_{ost}^+) \lambda_{odt}}{(1 - \lambda_{ost}^+) \lambda_{okt}} = \frac{\lambda_{odt}}{\lambda_{okt}}.$$

This proves the corollary.

Appendix F: Matrix Gradients and Fisher Preconditioner

Let $\mathbf{\Lambda}_t = \{\lambda_{odt}\}_{o,d=1}^S$ be the destination-choice probability matrix at time t , with zero diagonal. Define

$$\mathbf{Y}_t^+ = \sum_{m=1}^M \mathbf{Y}_t^m, \quad \mathbf{n}_t^+ = \sum_{m=1}^M (N_{1t}^m, \dots, N_{St}^m)^\top,$$

and

$$\mathbf{R}_t = \mathbf{Y}_t^+ - \text{diag}(\mathbf{n}_t^+) \mathbf{\Lambda}_t. \tag{7}$$

Because $\mathbf{R}_t \mathbf{1} = \mathbf{0}$, the row-centering operation in the logits does not add an explicit correction term to the matrix gradients. The likelihood-gradient blocks used in the implementation are

$$\nabla_{\mathbf{u}_r} \mathcal{L}_{\text{like}} = \sum_{t=1}^T w_{tr} \mathbf{R}_t \mathbf{v}_r, \quad (8)$$

$$\nabla_{\mathbf{v}_r} \mathcal{L}_{\text{like}} = \sum_{t=1}^T w_{tr} (\mathbf{R}_t)^\top \mathbf{u}_r, \quad (9)$$

$$[\nabla_{\mathbf{w}_r} \mathcal{L}_{\text{like}}]_t = \mathbf{u}_r^\top \mathbf{R}_t \mathbf{v}_r. \quad (10)$$

The full posterior gradients subtract the GP prior precision terms, e.g.,

$$\nabla_{\mathbf{u}_r} \mathcal{L} = \nabla_{\mathbf{u}_r} \mathcal{L}_{\text{like}} - \mathbf{K}_u^{-1} \mathbf{u}_r,$$

$$\nabla_{\mathbf{v}_r} \mathcal{L} = \nabla_{\mathbf{v}_r} \mathcal{L}_{\text{like}} - \mathbf{K}_v^{-1} \mathbf{v}_r,$$

$$\nabla_{\mathbf{w}_r} \mathcal{L} = \nabla_{\mathbf{w}_r} \mathcal{L}_{\text{like}} - \mathbf{K}_w^{-1} \mathbf{w}_r.$$

For each origin-time pair, let $\lambda_{ot} = (\lambda_{odt})_{d \in \mathcal{J}_o}$ and

$$\mathbf{F}_{ot} = \text{diag}(\lambda_{ot}) - \lambda_{ot} \lambda_{ot}^\top.$$

Let $\mathbf{P}_o \in \mathbb{R}^{(S-1) \times S}$ be the selection-centering matrix satisfying $\mathbf{P}_o \mathbf{v}_r = (v_{dr} - \bar{v}_{or})_{d \in \mathcal{J}_o}$. With $n_{ot}^+ = \sum_{m=1}^M N_{ot}^m$, the likelihood Fisher blocks are

$$\mathcal{I}_{u_r} = \sum_{o=1}^S \sum_{t=1}^T n_{ot}^+ \mathbf{e}_o \mathbf{e}_o^\top (w_{tr} \mathbf{P}_o \mathbf{v}_r)^\top \mathbf{F}_{ot} (w_{tr} \mathbf{P}_o \mathbf{v}_r), \quad (11)$$

$$\mathcal{I}_{v_r} = \sum_{o=1}^S \sum_{t=1}^T n_{ot}^+ (u_{or} w_{tr})^2 \mathbf{P}_o^\top \mathbf{F}_{ot} \mathbf{P}_o, \quad (12)$$

$$\mathcal{I}_{w_r} = \sum_{o=1}^S \sum_{t=1}^T n_{ot}^+ \mathbf{e}_t \mathbf{e}_t^\top (u_{or} \mathbf{P}_o \mathbf{v}_r)^\top \mathbf{F}_{ot} (u_{or} \mathbf{P}_o \mathbf{v}_r). \quad (13)$$

The Fisher-preconditioned pMALA metric for a generic block β with prior covariance $\mathbf{C}_\beta$ is

$$\mathbf{M}_\beta(\beta) = \left(\mathbf{C}_\beta^{-1} + \mathcal{I}_\beta(\beta) + \epsilon \mathbf{I} \right)^{-1}. \quad (14)$$

Appendix G: Algorithms for Model Inference and Probabilistic Forecasting

Algorithm 1 Fisher-preconditioned pMALA for a Tensor-Factor Block

- 1: **Input:** Current block $\beta^{(i)} \in \{u_r, v_r, w_r\}$, step size δ_β , jitter ϵ , prior covariance $C_\beta \in \{K_u, K_v, K_w\}$, current values of the remaining factors, and data Y .
 - 2: **Output:** New block state $\beta^{(i+1)}$.
 - 3: Compute R_t for $t = 1, \dots, T$ using Eq. (7).
 - 4: Compute the block-specific likelihood gradient $a^{(i)} = \nabla_\beta \mathcal{L}_{\text{like}}$ using Eq. (8), (9), or (10).
 - 5: Construct the corresponding likelihood Fisher block $\mathcal{I}_\beta^{(i)}$ using Eq. (11), (12), or (13).
 - 6: Compute full gradient $g^{(i)} = a^{(i)} - C_\beta^{-1} \beta^{(i)}$.
 - 7: Set $G^{(i)} = C_\beta^{-1} + \mathcal{I}_\beta^{(i)} + \epsilon I$, $M^{(i)} = (G^{(i)})^{-1}$, and compute $L^{(i)} (L^{(i)})^\top = M^{(i)}$.
 - 8: Calculate proposal mean $\mu^{(i)} = \beta^{(i)} + \frac{\delta_\beta}{2} M^{(i)} g^{(i)}$.
 - 9: Sample $\xi \sim \mathcal{N}(0, I)$ and propose $\beta^* = \mu^{(i)} + \sqrt{\delta_\beta} L^{(i)} \xi$.
 - 10: At β^* , recompute a^* , g^* , $\mathcal{I}_\beta^*$, G^* , M^* , and $\mu^* = \beta^* + \frac{\delta_\beta}{2} M^* g^*$.
 - 11: Calculate log-proposal densities:
 - $\log q(\beta^* | \beta^{(i)}) = -\frac{1}{2} \log |\delta_\beta M^{(i)}| - \frac{1}{2\delta_\beta} (\beta^* - \mu^{(i)})^\top G^{(i)} (\beta^* - \mu^{(i)})$,
 - $\log q(\beta^{(i)} | \beta^*) = -\frac{1}{2} \log |\delta_\beta M^*| - \frac{1}{2\delta_\beta} (\beta^{(i)} - \mu^*)^\top G^* (\beta^{(i)} - \mu^*)$.
 - 12: Calculate full posterior log-ratio:

$$\log \alpha = \left(\mathcal{L}(\beta^*) - \mathcal{L}(\beta^{(i)}) \right) + \left(\log q(\beta^{(i)} | \beta^*) - \log q(\beta^* | \beta^{(i)}) \right).$$
 - 13: Draw $a \sim \text{Uniform}(0, 1)$.
 - 14: **if** $\log(a) < \log \alpha$ **then**
 - 15: $\beta^{(i+1)} = \beta^*$
 - 16: **else**
 - 17: $\beta^{(i+1)} = \beta^{(i)}$
 - 18: **end if**
 - 19: **return** $\beta^{(i+1)}$.
-

Algorithm 2 Slice Sampling for Hyperparameter θ_j

- 1: **Input:** Current state $\theta_j^{(i)}$, posterior log-density $\mathcal{L}(\phi_j)$, current factors U, V, W .
 - 2: **Output:** New state $\theta_j^{(i+1)}$.
 - 3: Compute $\phi_j = \log(\theta_j^{(i)})$.
 - 4: Log-likelihood threshold: $\gamma \sim \text{Uniform}(0, 1)$, $\log c = \mathcal{L}(\phi_j) + \log \gamma$.
 - 5: Draw an initial sampling range: $\kappa \sim \text{Uniform}(0, \epsilon)$, $\phi_{\min} = \phi_j - \kappa$, $\phi_{\max} = \phi_{\min} + \epsilon$.
 - 6: **loop**
 - 7: Propose new state: $\phi_j^* \sim \text{Uniform}(\phi_{\min}, \phi_{\max})$.
 - 8: **if** $\mathcal{L}(\phi_j^*) > \log c$ **then**
 - 9: $\theta_j^{(i+1)} = \exp(\phi_j^*)$.
 - 10: **break**
 - 11: **else**
 - 12: **if** $\phi_j^* < \phi_j$ **then:** $\phi_{\min} = \phi_j^*$ **else:** $\phi_{\max} = \phi_j^*$.
 - 13: **end if**
 - 14: **end loop**
 - 15: **return** $\theta_j^{(i+1)}$.
-

Algorithm 3 Centered and Non-centered Interweaving for GP Hyperparameters

- 1: **Input:** Centered factors U, V, W , current covariance hyperparameters θ_K , data Y , and the slice sampler in Algorithm 2.
 - 2: **Output:** Updated centered factors U, V, W and hyperparameters θ_K .
 - 3: Construct current Cholesky factors L_u, L_v, L_w from $K_u(\theta_u)$, $K_v(\theta_v)$, and $K_w(\theta_w)$.
 - 4: Whiten the current centered factors: $Z_u = L_u^{-1}U$, $Z_v = L_v^{-1}V$, and $Z_w = L_w^{-1}W$.
 - 5: **for** each hyperparameter θ_j in the origin spatial covariance block **do**
 - 6: Define $\phi_j = \log(\theta_j)$ and reconstruct factors by $U(\phi_j) = L_u\{\theta_u(\phi_j)\}Z_u$.
 - 7: Define slice target $\mathcal{L}_{\text{NC}}(\phi_j) = \log p(Y | U(\phi_j), V, W) + \log p(\exp(\phi_j)) + \phi_j$.
 - 8: Update θ_j using Algorithm 2.
 - 9: **end for**
 - 10: Reconstruct $U = L_u(\theta_u)Z_u$ using the accepted origin spatial hyperparameters.
 - 11: **for** each hyperparameter θ_j in the destination spatial covariance block **do**
 - 12: Define $\phi_j = \log(\theta_j)$ and reconstruct factors by $V(\phi_j) = L_v\{\theta_v(\phi_j)\}Z_v$.
 - 13: Define slice target $\mathcal{L}_{\text{NC}}(\phi_j) = \log p(Y | U, V(\phi_j), W) + \log p(\exp(\phi_j)) + \phi_j$.
 - 14: Update θ_j using Algorithm 2.
 - 15: **end for**
 - 16: Reconstruct $U = L_u(\theta_u)Z_u$ and $V = L_v(\theta_v)Z_v$ using the accepted spatial hyperparameters.
 - 17: **for** each hyperparameter θ_j in the temporal covariance block **do**
 - 18: Define $\phi_j = \log(\theta_j)$ and reconstruct trial factors by $W(\phi_j) = L_w\{\theta_w(\phi_j)\}Z_w$.
 - 19: Define slice target $\mathcal{L}_{\text{NC}}(\phi_j) = \log p(Y | U, V, W(\phi_j)) + \log p(\exp(\phi_j)) + \phi_j$.
 - 20: Update θ_j using Algorithm 2.
 - 21: **end for**
 - 22: Reconstruct $W = L_w(\theta_w)Z_w$ using the accepted temporal hyperparameters.
 - 23: **return** U, V, W, θ_K .
-

Algorithm 4 Overall MCMC Sampling Algorithm for Model Estimation

- 1: **Input:** Observed tensors $\mathbf{Y}$, CP rank R , burn-in iterations I_1 , sampling iterations I_2 .
 - 2: **Output:** Posterior samples $\left\{ \mathbf{U}^{(i)}, \mathbf{V}^{(i)}, \mathbf{W}^{(i)}, \boldsymbol{\theta}_K^{(i)} \right\}_{i=I_1+1}^{I_1+I_2}$.
 - 3: Initialize $\mathbf{U}^{(0)}, \mathbf{V}^{(0)}, \mathbf{W}^{(0)}, \boldsymbol{\theta}_K^{(0)}$.
 - 4: **for** $i = 1$ to $I_1 + I_2$ **do**
 - 5: **for** $r = 1$ to R **do**
 - 6: Sample origin factor $\mathbf{u}_r^{(i)}, \mathbf{v}_r^{(i)}, \mathbf{w}_r^{(i)}$ using Fisher-preconditioned pMALA with Algorithm 1.
 - 7: **end for**
 - 8: Update covariance hyperparameters $\boldsymbol{\theta}_K^{(i)}$ using centered/non-centered interweaving with Algorithm 3;
the one-dimensional log-scale updates inside this step use slice sampling with Algorithm 2.
 - 9: **end for**
 - 10: **return** $\left\{ \mathbf{U}^{(i)}, \mathbf{V}^{(i)}, \mathbf{W}^{(i)}, \boldsymbol{\theta}_K^{(i)} \right\}_{i=I_1+1}^{I_1+I_2}$.
-

Algorithm 5 Posterior Predictive Inference for New Location s

- 1: **Input:** Posterior samples $\left\{ \mathbf{U}^{(i)}, \mathbf{V}^{(i)}, \mathbf{W}^{(i)}, \boldsymbol{\theta}_K^{(i)} \right\}_{i=1}^{I_2}$, network information for new location s (to compute covariances), estimated departures $\{N_{st}\}_{t=1}^T$ from new location s .
 - 2: **Output:** Predictive samples of O-D flows $\left\{ \mathbf{y}_{st}^{(i)} \right\}_{i=1, t=1}^{I_2, T}$.
 - 3: **for** $i = 1$ to I_2 **do**
 - 4: Compute $\mathbf{K}_u^{(i)}, \mathbf{K}_v^{(i)}, \mathbf{k}_{us}^{(i)}, \mathbf{k}_{vs}^{(i)}, k_{uss}^{(i)}, k_{vss}^{(i)}$ using the spatial kernel functions.
 - 5: Compute predictive variances $(\sigma_{us}^2)^{(i)}$ and $(\sigma_{vs}^2)^{(i)}$.
 - 6: **for** $r = 1$ to R **do**
 - 7: Compute predictive means $(\mu_r^{us})^{(i)}$ and $(\mu_r^{vs})^{(i)}$.
 - 8: Draw new factors $(u_r^s)^{(i)}$ and $(v_r^s)^{(i)}$ from the Gaussian predictive distributions.
 - 9: **end for**
 - 10: Construct $\mathbf{u}_s^{(i)} = \left[(u_1^s)^{(i)}, \dots, (u_R^s)^{(i)} \right]^\top$ and $\mathbf{v}_s^{(i)} = \left[(v_1^s)^{(i)}, \dots, (v_R^s)^{(i)} \right]^\top$.
 - 11: **for** $t = 1$ to T **do**
 - 12: Compute logits $G_{sdt}^{(i)} = \sum_{r=1}^R (u_r^s)^{(i)} v_{dr}^{(i)} w_{tr}^{(i)}$ for $d \in \{1, \dots, S\}$.
 - 13: Compute probabilities $\lambda_{st}^{(i)}$ using the masked softmax.
 - 14: Draw predictive counts: $\mathbf{y}_{st}^{(i)} \sim \text{Multinomial} \left(N_{st}, \lambda_{st}^{(i)} \right)$.
 - 15: **end for**
 - 16: **end for**
 - 17: **return** All predictive samples $\left\{ \mathbf{y}_{st}^{(i)} \right\}_{i=1, t=1}^{I_2, T}$.
-

Appendix H: Some Tables and Figures in the Experiments

Table 1 Inductive bike O-D forecast performance under varying observed proportions/ranks.

Observed proportion	Rank	logL (10^6)	MAE	RMSE	CRPS
10%	1	-2.087	0.022574	0.128420	0.021744
10%	2	-2.076	0.022258	0.128816	0.020776
10%	5	-2.035	0.022064	0.128837	0.020203
10%	10	-2.069	0.022070	0.129494	0.019905
10%	15	-2.063	0.021893	0.130349	0.019592
30%	1	-2.069	0.022486	0.128045	0.021557
30%	2	-2.012	0.022078	0.127885	0.020402
30%	5	-1.966	0.021694	0.127955	0.019665
30%	10	-1.946	0.021537	0.127113	0.019410
30%	15	-1.949	0.021519	0.127146	0.019368
50%	1	-2.059	0.022447	0.127984	0.021389
50%	2	-2.000	0.021978	0.127188	0.020160
50%	5	-1.948	0.021498	0.127633	0.019332
50%	10	-1.939	0.021379	0.127960	0.019256
50%	15	-1.925	0.021341	0.127515	0.019239
70%	1	-2.039	0.022385	0.126865	0.021352
70%	2	-1.953	0.021905	0.125528	0.020124
70%	5	-1.882	0.021336	0.123889	0.019196
70%	10	-1.871	0.021168	0.123435	0.019028
70%	15	-1.867	0.021081	0.123455	0.018926
90%	1	-2.029	0.022371	0.126780	0.021259
90%	2	-1.944	0.021873	0.125689	0.020029
90%	5	-1.870	0.021275	0.124456	0.019101
90%	10	-1.856	0.021094	0.124065	0.018920
90%	15	-1.849	0.020993	0.123360	0.018812

Table 2 Inductive metro O-D forecast performance with varying observed proportions/ranks.

Observed proportion	Rank	logL (10^6)	MAE	RMSE	CRPS
10%	1	-148.717	8.4009	23.7554	7.2945
10%	2	-147.944	8.2723	23.6477	7.1934
10%	5	-149.648	8.3659	24.1015	7.3290
10%	10	-148.320	8.2014	24.8707	7.1675
10%	15	-150.149	8.4064	24.6396	7.3545
10%	20	-150.178	8.4189	24.3917	7.3633
30%	1	-146.642	7.8128	21.7549	6.7623
30%	2	-145.239	7.5330	21.1748	6.5115
30%	5	-144.533	7.2463	21.0766	6.2950
30%	10	-144.598	7.1568	21.4866	6.2350
30%	15	-144.740	7.1667	21.5719	6.2470
30%	20	-145.229	7.2181	21.5942	6.2974
50%	1	-145.916	7.7158	21.9070	6.6988
50%	2	-144.349	7.4116	21.3580	6.4346
50%	5	-143.971	7.1818	23.6866	6.2871
50%	10	-142.938	6.8255	24.5503	5.9697
50%	15	-142.730	6.7415	24.0499	5.8909
50%	20	-142.229	6.6422	22.8594	5.7899
70%	1	-142.329	7.0743	20.3997	6.0725
70%	2	-139.870	6.4902	19.2653	5.5425
70%	5	-135.932	5.4306	17.0432	4.5726
70%	10	-134.781	4.9498	17.0001	4.1509
70%	15	-134.532	4.8292	16.5844	4.0400
70%	20	-134.376	4.7289	16.4594	3.9512
90%	1	-141.144	6.7820	19.7189	5.7967
90%	2	-138.360	6.1439	18.2477	5.2116
90%	5	-133.795	4.7994	14.9143	3.9736
90%	10	-131.569	3.8913	12.8672	3.1448
90%	15	-131.233	3.7382	12.1666	3.0053
90%	20	-130.988	3.5930	11.8154	2.8754

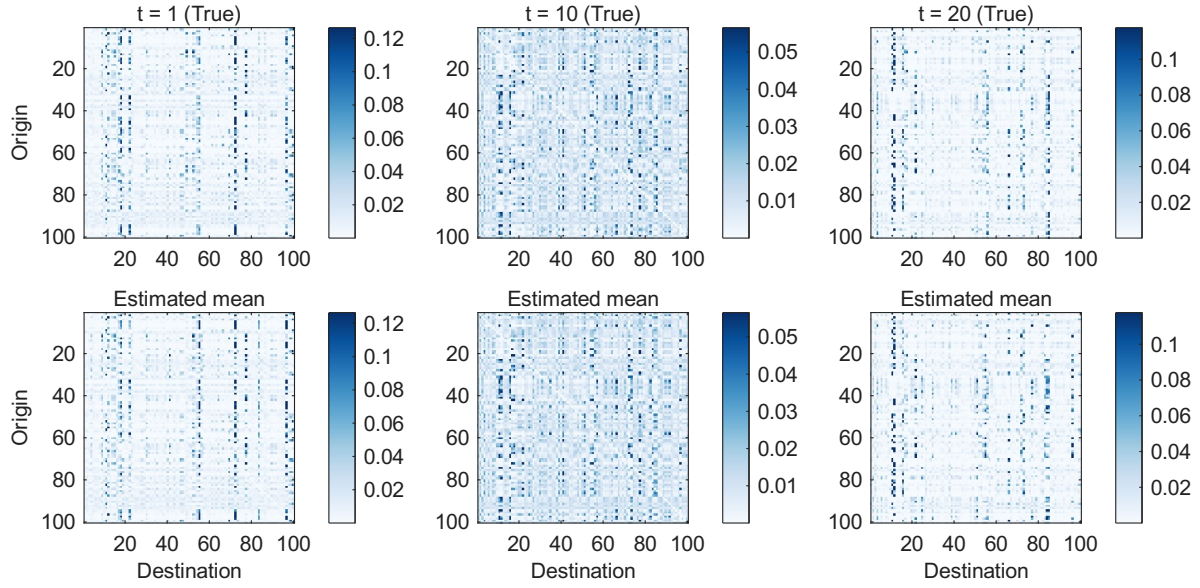


Figure 1 Visual comparison of O-D assignment probability matrices.

Note. The top row displays the ground truth probabilities for $t = 1, 10, 20$, while the bottom row displays the corresponding estimated posterior means. The color intensity represents the probability magnitude.

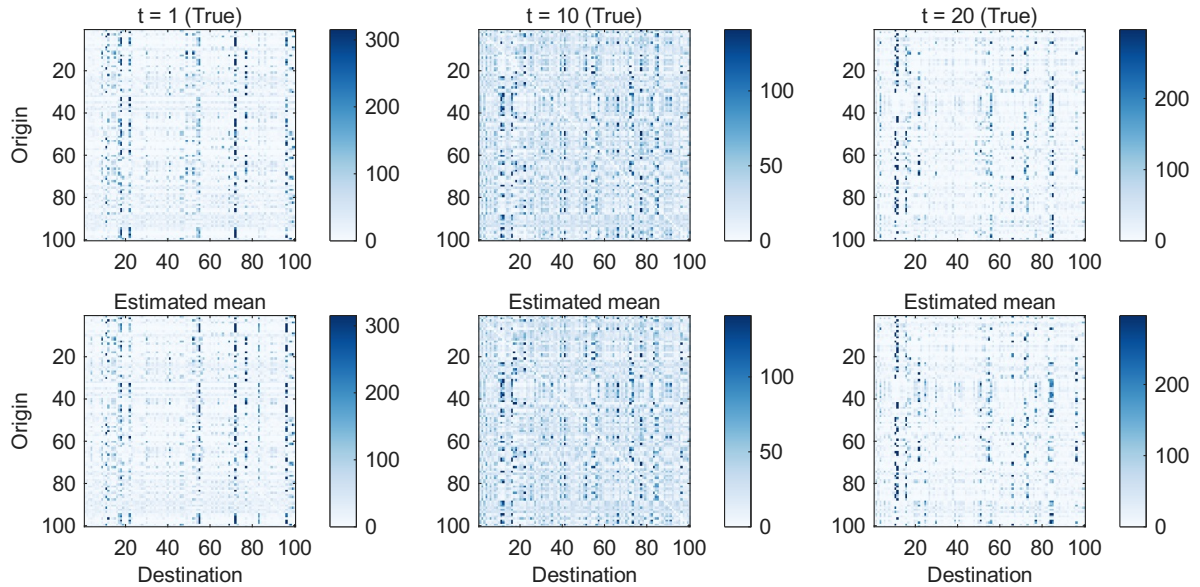


Figure 2 Visual comparison of O-D matrices.

Note. The top row displays the true O-D matrices for $t = 1, 10, 20$, while the bottom row displays the corresponding estimated posterior O-D means. The color intensity represents the O-D flow.

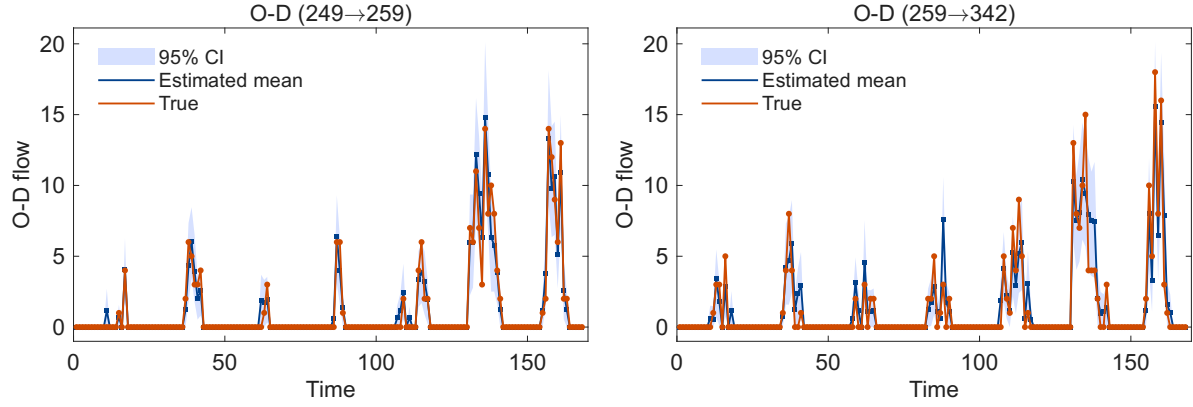


Figure 3 Estimated bike O-D flows with uncertainty quantification.

Note. The shaded regions display posterior predictive 95% credible intervals, the blue line is the posterior predictive mean, and the red points/line are the observed validation flows.

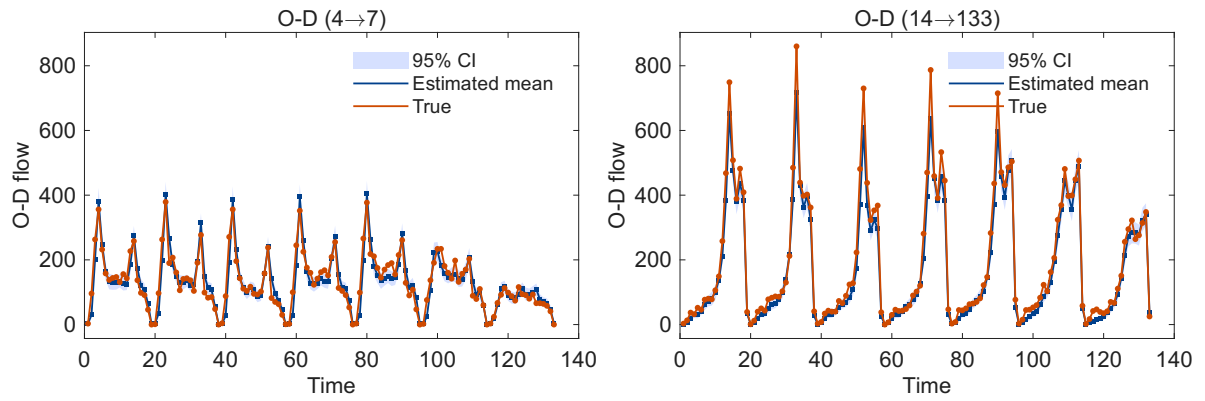


Figure 4 Estimated metro O-D flows with uncertainty quantification

Note. The shaded regions display posterior predictive 95% credible intervals, the blue line is the posterior predictive mean, and the red points/line are the observed validation flows.

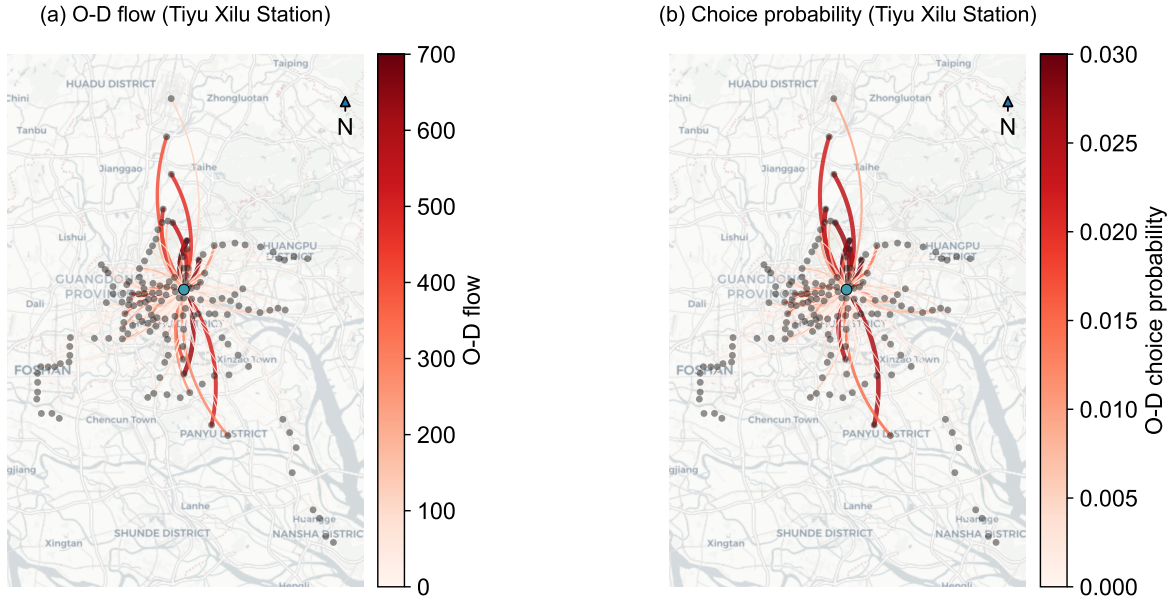


Figure 5 True O-D flow and the estimated O-D choice probabilities for TiYu Xilu Station.

Note. The left panel shows aggregate validation O-D flows from the selected origin, and the right panel shows the estimated destination-choice probabilities.

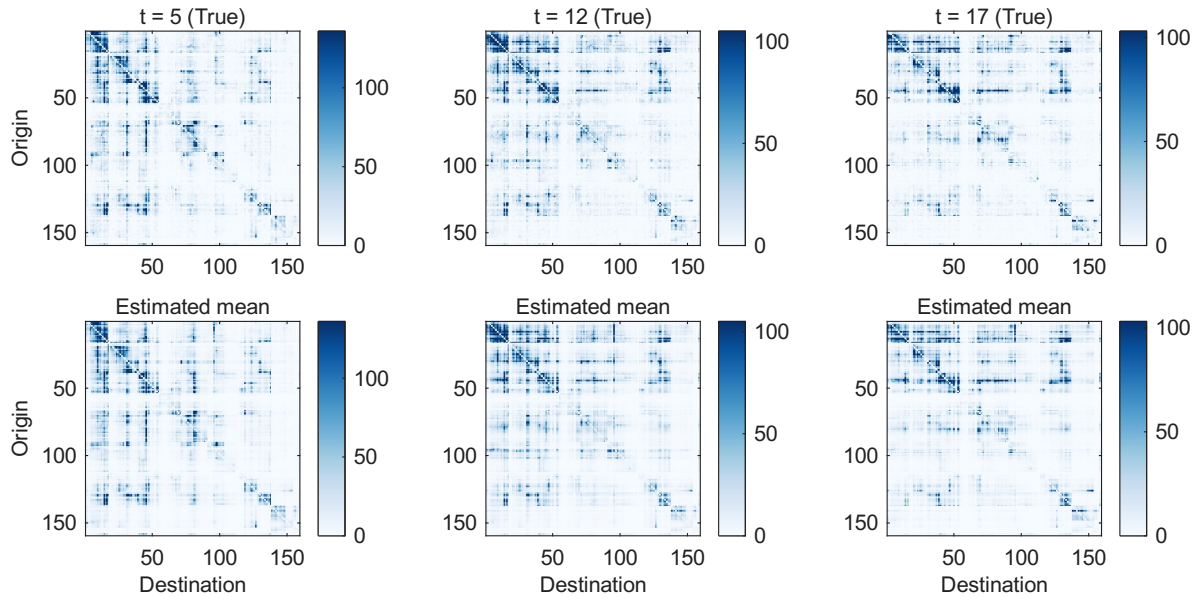


Figure 6 Visual comparison of estimated and true metro O-D matrices.

Note. The top row displays the true O-D matrices for $t = 5, 12, 17$, while the bottom row displays the corresponding estimated posterior O-D means. The color intensity represents the O-D flow on log scale.

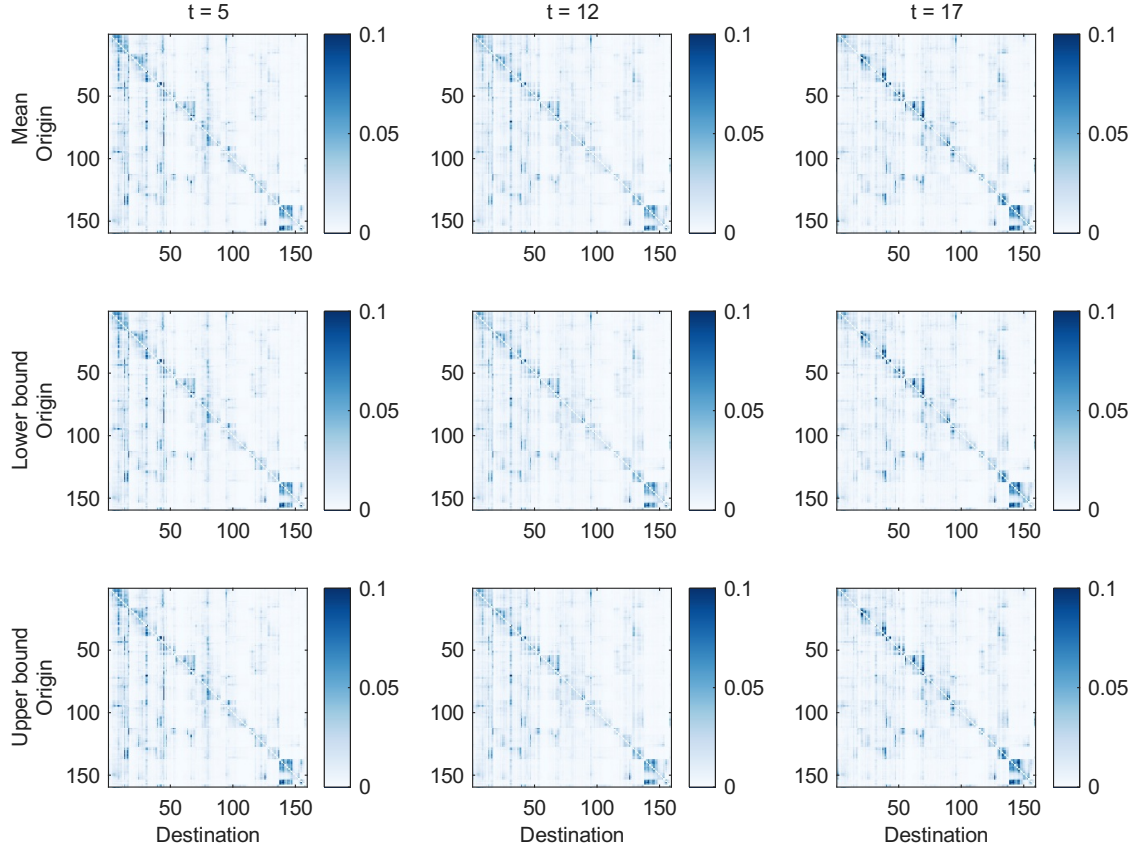


Figure 7 Estimated metro O-D assignment probability matrices.

Note. The top row displays the estimated posterior mean of probabilities for $t = 5, 12, 17$, while the second and the last rows respectively display the corresponding estimated lower and upper bounds of 95% CIs. The color intensity represents the probability magnitude.