

TRAVEL AS PRODUCTION: ENDOGENOUS IN-TRIP OUTPUT, SCHEDULING, AND THE BOTTLENECK

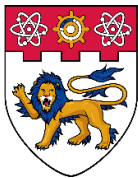
Zach J.L. Lee

July 2026

EGC Report No: 2026/03

HSS-04-90A
Tel: +65 67906073
Email: D-EGC@ntu.edu.sg

Economic Growth Centre Working Paper Series



**NANYANG
TECHNOLOGICAL
UNIVERSITY**
SINGAPORE

Economic Growth Centre
School of Social Sciences

The author(s) bear sole responsibility for this paper.

Views expressed in this paper are those of the author(s) and not necessarily those of the Economic Growth Centre, NTU.

Travel as Production: Endogenous In-Trip Output, Scheduling, and the Bottleneck

Zach J.L. Lee

Nanyang Technological University, Singapore

11 July 2026

[Latest version](#)

Abstract

Automation, ride-hailing, and other passenger services can make travel time usable for work or rest. This paper models travel as production. Scheduled activity access is the final commodity; trips are intermediate inputs; and mode, route, departure-time, and vehicle-technology choices are production plans. At a reference allocation, active travel time costs the resource value of time plus effort, whereas passive time costs that value net of reclaimed output. A constant-returns implicit-price rule appears as a special case. Embedding these time values in a Vickrey–Arnott–de Palma–Lindsey bottleneck distinguishes three effects. In a homogeneous fixed-demand pure point bottleneck, productive in-vehicle time increases waiting and queue stock but leaves money-metric generalised user cost and the optimal time-varying toll unchanged within the model’s cost accounting. A lower effective early-arrival penalty shifts arrivals earlier and reduces maximum waiting; with elastic demand, lower free-flow cost induces trips and raises the optimal toll. The production structure also yields candidate cross-equation restrictions that generic choice models need not satisfy. Stylised simulations illustrate the results.

JEL: R41, R48, D13, D61, L91.

Keywords: household production; value of travel time; bottleneck model; scheduling; congestion pricing; autonomous vehicles; revealed preference restrictions.

Email: zjllee@nus.edu.sg. Current affiliation: Lloyd’s Register Foundation Institute for the Public Understanding of Risk, National University of Singapore. This paper is available on SSRN and in the NTU Economic Growth Centre Working Paper Series. Refine.ink was used to check the paper for consistency and clarity.

1 Introduction

Automation, ride-hailing, and other passenger services can make travel time usable for work, rest, or other activity. For these trips, transport choice depends on money, clock time, and what travellers can produce while moving.

Two views of travel are especially useful for modelling this problem. In the first, from random utility and engineering demand models, a route, mode, departure time, or fare class is an *alternative* carrying attributes. The analyst writes an indirect utility or a generalised cost and estimates its coefficients. The second view draws on Becker’s allocation of time and Lancaster’s characteristics approach. It treats travel as part of a household *technology*. Time, money, goods, attention, and effort are combined to produce valued activities such as work, care, schooling, shopping, social contact, recreation, and rest (Becker, 1965; Lancaster, 1966). In these derived-demand cases, people do not value a trip in itself.¹ They value what the trip makes possible.

I develop the second view because it treats in-trip output as part of the traveller’s production problem, rather than as an exogenous attribute of a mode, and then connect it to bottleneck models of scheduling and congestion. The distinction is mostly one of accounting. Once what the trip makes possible is treated as the final commodity, choosing a trip and choosing a production plan are the same cost-minimisation problem. The final commodity is **scheduled activity access**, meaning the completion of an activity at an acceptable time and with acceptable comfort, reliability, and safety. A *trip* is an *intermediate input*. A mode, route, departure time, fare class, or vehicle technology is a **production plan** that combines priced inputs (money, active and passive clock time, physical and cognitive effort, schedule slack) to deliver the access. The traveller chooses the plan that minimises the conditional cost of the access she needs.

The framework has three main implications.

The value of time. Following DeSerpa (1971) and Jara-Díaz (2007), the value of travel-time savings (VTTS) reflects the shadow value of scarce time, as captured by binding time constraints. Recent time-use models already make onboard work and leisure explicit, and recent theory derives VTTS from those activities (Jara-Díaz, 2024; Pudāne & Correia, 2020; Pudāne et al., 2018). I use a deliberately parsimonious within-trip allocation to carry that logic into a bottleneck equilibrium. If passive travel time can be turned into paid-equivalent work or into rest, the opportunity cost of that time is reduced; if

¹Recreational and scenic travel are exceptions in which the journey itself is partly consumed.

the traveller must actively drive, the cost is the resource value of time *plus* effort. The resulting time coefficients are restricted by the shadow value of time, driving effort, and net reclaimed output rather than estimated as unrelated mode-time coefficients. Those underlying objects remain primitives to be measured, and the closed-form mapping is a local approximation rather than a claim to be the first model of travel-time use.

Scheduling and congestion. A choice model that leaves travel times exogenous is a partial-equilibrium consumer model. Building on bottleneck work on autonomous vehicles and heterogeneous values of time, the production technology is embedded in the canonical Vickrey–Arnott–de Palma–Lindsey bottleneck model (Arnott et al., 1990; Small, 1982; Vickrey, 1969), where in-trip productivity and driving effort enter through *effective* scheduling parameters. In a homogeneous, fixed-demand pure point bottleneck, making in-vehicle time productive does **not** change equilibrium money-metric user cost or the optimal time-varying money toll; it lengthens maximum waiting time and queue stock. This is the standard result that equilibrium bottleneck cost is independent of the marginal value of queueing time, interpreted through the production model. Model social cost and pricing respond to a lower effective *early-arrival* penalty, which shifts arrivals earlier and lowers the maximum wait, and to cheaper in-vehicle time under *elastic demand*, which induces trips and raises the optimal toll. Automation is the running case in the paper, but the same logic applies to delegated passenger services, such as taxis, ride-hailing, and chauffeur services, whenever they turn in-vehicle time from active time into passive, productive, or restorative time. Table 2 collects the distinct parameter–environment combinations.

Testability. A production model has empirical content only if it rules out some choice patterns. The paper therefore states three candidate cross-equation restrictions that a generic random-utility or multiple-discrete-continuous model with free time coefficients need not impose. One is that telework feasibility lowers the value of passive travel time but not active driving time. Another is a non-equivalence between money and time shocks once time use is observed. A fourth item is retained only as a joint revealed-preference diagnostic, not as a production-specific restriction.

The paper proceeds as follows. Section 2 places the model among four literatures. Section 3 sets up the travel-production economy. Section 4 derives the net-reclaimed-output VTTS decomposition and shows how a constant-returns implicit-price rule arises as a special case. Section 5 develops the production-bottleneck equilibrium and its comparative statics, including the neutrality, peak-shifting, induced-demand, and sorting results, and the robustness of the neutrality result to smooth scheduling and a duration-dependent

marginal waiting cost. Section 6 extends the value of reliability. Section 7 states the empirical restrictions and diagnostic. Section 8 draws the model cost and pricing implications. Section 9 summarises the main implication for productive travel time and discusses limitations and extensions. Notation is compiled in [Appendix A](#), proofs are in [Appendix B](#), and illustrative calibration details are in [Appendix C](#).

2 Related literature

The model sits at the intersection of four strands.

Household production and time allocation. Becker ([1965](#)) treated time as a scarce input and utility as a function of produced commodities; Lancaster ([1966](#)) shifted utility onto characteristics; DeSerpa ([1971](#)) showed that the value of saving time derives from binding time and activity constraints, not from the wage mechanically. Pollak and Wachter ([1975](#)) and Muellbauer ([1974](#)) warned that joint household production and endogenous implicit prices can weaken the empirical content of the approach. The paper therefore uses conditional cost as its organising object and does not assume that implicit prices are constant; the scalar constant-price case is derived only under explicit restrictions. Jara-Díaz ([2007](#); [2003](#)) gave the modern synthesis, decomposing the value of travel time into the value of leisure minus the value of time assigned to travel. Jara-Díaz ([2024](#)) subsequently derives that second term when work and leisure can be performed in the vehicle. The present model adopts a compact net-reclaimed-output representation so that the resulting effective parameters can be taken into a bottleneck.

Value of travel time and travel-time use. Small's ([2012](#)) survey is the reference point. The empirical value of time is heterogeneous and context-dependent, often a fraction of the wage. Pudāne et al. ([2018](#)) model the reallocation of daily activities when automated vehicles permit onboard activities; Pudāne and Correia ([2020](#)) clarify the time-allocation implications of work- and leisure-equipped vehicles. Within-person choice evidence finds lower values of time when travellers can conduct onboard activities (Molin et al., [2020](#)), while an empirical meta-synthesis finds substantial heterogeneity across private and shared automated-vehicle settings (Huda et al., [2023](#)). The current review evidence likewise treats travel-time use as affecting time allocation, mode choice, VTTS, travel experience, and second-order travel responses (de Sá et al., [2025](#)). Relative to this literature, the paper does not claim that productive travel time or activity-dependent VTTS is new. Its narrower contribution is to express the local value as a resource value of time net of

reclaimed output and effort and then trace separately how queue-time valuation, early-arrival valuation, free-flow cost, and capacity enter a bottleneck.

Scheduling, bottlenecks, and equilibrium. Vickrey (1969) introduced the bottleneck formulation with endogenous departure-time choice and schedule-delay costs. Small (1982) provided the empirical work-trip scheduling specification. Arnott–de Palma–Lindsey (1990, 1993) developed the canonical equilibrium and congestion-pricing analysis. Smith (1984) proved existence of a time-dependent bottleneck equilibrium, Daganzo (1985) proved uniqueness in the single-class case, and Lindsey (2004) extended existence, uniqueness, and trip-cost results to multiple user classes. Fosgerau and Small (2017) endogenise the *shape* of scheduling preferences from activity-utility profiles; Fosgerau, Kim and Ranjan (2018) embed the bottleneck in a monocentric city. Fosgerau and Karlström (2010) characterise the value of reliability. Most directly, van den Berg and Verhoef (2016) study autonomous vehicles in the bottleneck and separate their capacity, value-of-time, and preference-heterogeneity effects. Earlier work characterises sorting by the value of time in the heterogeneous bottleneck (Arnott et al., 1988; V. van den Berg & Verhoef, 2011). The contribution here is not a new bottleneck equilibrium, the standard independence of homogeneous bottleneck cost from α , or a first model of travel-time use. It is the common accounting that maps net reclaimed output and effort into effective queue and schedule parameters, separates parameter–environment combinations with different cost and pricing effects, and states candidate cross-equation restrictions.

Several closely related activity-and-bottleneck studies already model onboard home or work activities, mixed conventional and automated users, supply, tolling, capacity, or uncertainty. Table 1 summarises their overlap with the present paper and the specific focus of this analysis.

Table 1: Overlap and differences between the closest related studies and the present analysis.

Study	Main overlap with this paper	Main difference from this paper
Pudāne (2020)	Onboard home/work activities alter departure-time choice and bottleneck congestion.	Uses activity-specific scheduling efficiencies; the present paper instead starts from a local net-reclaimed-output/effort accounting and separates queue-time, early-arrival, and free-flow channels.
Yu et al. (2022)	Activity-based bottleneck with normal and automated cars, alternative vehicle-supply regimes, road pricing, and welfare comparisons.	The present paper does not model vehicle supply or endogenous joint adoption; it gives a conditional scheduled-access accounting and candidate cross-equation restrictions.
Wu and Li (2023)	Mixed manned/automated bottleneck, endogenous onboard activity choice, segregation, and differentiated or anonymous tolling.	The present sorting result is conditional on fixed class masses and common schedule penalties; it is not a new mixed-fleet tolling or adoption result.
Jara-Díaz (2024)	Derives the value of travel time when work and leisure can be performed in the vehicle.	The present paper uses a deliberately compact local mapping to carry net reclaimed output and effort into bottleneck parameters.
Liu et al. (2025)	Mixed automated/human-driven traffic with lower value of time, capacity enhancement, and stochastic bottleneck capacity.	The present analysis uses a deterministic pure-point setting; its capacity, reliability, and external-cost limits are kept separate.

Study	Main overlap with this paper	Main difference from this paper
Wang et al. (2026)	Values labour time in automated vehicles and traces wider economic effects.	The present paper is a partial-equilibrium bottleneck account and does not claim a complete employer-output, operator-cost, or economy-wide welfare closure.

The paper therefore does not claim to provide the first activity-based autonomous-vehicle bottleneck. Its specific contribution is a common scheduled-access accounting that links a local resource value of time, net reclaimed onboard output, and active effort to effective bottleneck parameters, distinguishes the effects of queue-time value, early-arrival value, free-flow cost, and capacity, and yields the candidate empirical restrictions in Section 7.

Activity-based and positive-utility travel. Activity-based and discrete-continuous models operationalise derived demand (Bhat, 2005, 2008; Bowman & Ben-Akiva, 2001; Kockelman, 2001; Pinjari & Bhat, 2010); latent-variable models incorporate unobserved perceptions of service quality and traveller attitudes (Bhat & Dubey, 2014; Vij & Walker, 2016). The positive-utility-of-travel literature shows travel time can produce work, rest, or enjoyment (Jain & Lyons, 2008; Malokin et al., 2019; Mokhtarian, 2019; Mokhtarian & Salomon, 2001; Singleton, 2019), a choice that becomes especially important when travel time becomes passive, as in automated vehicles, taxis, or ride-hailing services (Krueger et al., 2019). The conditional-cost interpretation used here clarifies which time-use components enter each bottleneck parameter; its empirical content rests on the joint restrictions in Section 7, not on the label “production” alone.

3 A travel-production economy

3.1 Commodities, technologies, and the production plan

A traveller must complete a scheduled activity (reach a destination to work, study, shop, or receive care) by a desired clock time t^* . To do so, she selects a **technology** j from

a finite set J (a mode, route, fare class, vehicle type, or compound itinerary) and its associated input requirements. The final commodity is the access itself, valued at B (the benefit of completing the activity). Its net value depends on the production plan because timing, comfort, reliability, safety, and effort vary across modes and departure choices.

Let the traveller's time endowment $\bar{T}$ be divided among market work $W \geq 0$, non-travel leisure or home time $L \geq 0$, active travel time T_j^A (driving, cycling, navigating, and other attention-intensive travel), and passive travel time T_j^P (riding without operating the vehicle). During passive travel, she allocates $\tau^b \in [0, T_j^P]$ to in-trip work and the remainder $\tau^r = T_j^P - \tau^b \geq 0$ to rest. Assume $A_j \equiv \bar{T} - T_j^A - T_j^P \geq 0$, so the time allocation is feasible. Feasibility requires $0 \leq W \leq A_j$; a fixed- or bounded-hours variant adds a nonempty interval $\underline{W} \leq W \leq \bar{W}$ within that range. Technology j determines how travel time is transformed while access is produced. Passive time can generate marketable work or leisure-equivalent rest, while active time requires effort. The technology is described by four production parameters:

- $\phi_j \in [0, 1]$: net effective in-trip work productivity, expressed as the fraction of a wage-equivalent return retained after task frictions and incremental cognitive effort; $\phi_j = 0$ for active driving.
- $\psi_j \in [0, 1]$: comfort/rest substitution (the leisure-equivalent value of in-trip rest).
- $e_j \geq 0$: marginal effort cost per unit of active travel time.
- the schedule consequences (earliness, lateness) implied by departure timing, developed in Section 5.

The money budget is quasilinear in a numeraire G :

$$G = I + w(W + \phi_j \tau^b) - p_j, \quad (1)$$

where I is non-labour income, w the wage, and p_j the money price of the technology (fare, fuel, tolls, parking). The time constraint is

$$W + L = \bar{T} - T_j^A - T_j^P \equiv A_j. \quad (2)$$

Preferences excluding the common benefit B are quasilinear:

$$U_j = G + u_L(L + \psi_j \tau^r) + u_W(W) - e_j T_j^A, \quad (3)$$

with u_L increasing, continuously differentiable, and strictly concave (the value of discretionary time), and u_W continuously differentiable and concave, capturing direct (dis)util-

ity from market work performed outside the trip. These conditions and the compact feasible set ensure a within-plan optimum. The term $L + \psi_j \tau^r$ says in-trip rest substitutes, at rate ψ_j , for ordinary leisure; the term $w\phi_j \tau^b$ in Equation 1 is the *net money-metric value* of marketable in-trip output, such as cash for hourly work or the avoided cost of overtime or cleared backlog for salaried work, not necessarily a contemporaneous wage payment. For tractability, any incremental task effort or disutility from working in the vehicle is absorbed into the net effective rate $w\phi_j$ rather than entered separately in u_W . The model therefore does not separately identify gross onboard productivity and onboard work disutility. A worker with fixed compensation and no bankable output has $\phi_j \approx 0$ and reclaims passive time through rest alone, $\rho_j = \psi_j \omega$, so the framework nests the rigid-hours salaried case.

Over the relevant trip, ϕ_j and ψ_j are treated as constant marginal rates. For empirical work, ϕ_j is a net effective rate after incremental task effort, task indivisibilities, setup time, and switching costs. This is a maintained aggregation, not a claim that work and rest are literally linear or that the two components can be separately recovered from choice data alone. If those frictions are modelled explicitly, reclaimed work is nonlinear in passive travel time; the conditional-cost formulation still applies, but the closed-form VTTS expression in Proposition 1 becomes a local marginal expression.

For a given technology and its travel times, define $V_j \equiv \max_{(W, \tau^b) \text{ feasible}} U_j$ and let V_0 be the value of the outside option in which the scheduled access is not produced. Completing the access gives total utility $B + V_j$. Quasilinearity makes

$$C_j \equiv V_0 - V_j$$

an exact money-metric **conditional user cost** relative to the outside option. Conditional technology choice and participation are therefore

$$j \in \arg \min_{k \in J} C_k, \quad B \geq \min_{k \in J} C_k.$$

Thus B cancels from technology choice once participation is fixed. C_j values the inputs consumed by plan j , net of reclaimed output; it is not yet a complete social resource account because wages, employer output, operator costs, and physical externalities are not all modelled. In this sense, each travel option is a technology for producing scheduled activity access.

Under elastic demand, $P(N)$ in Proposition 5 is the marginal willingness to pay for scheduled activity access. What is specific to treating access as the final output, rather

than each trip as a final good, is that the value of travel time is the shadow price of the binding time constraint net of in-trip output ([Proposition 1](#)) and that automation acts through the effective scheduling parameters, not through B .

3.2 The resource value of time

Let $X^* = L^* + \psi_j \tau^{r*}$ denote optimised effective leisure, and let $\omega \equiv u'_L(X^*)$ be the local shadow value of discretionary time. With flexible labour, W is interior and the work first-order condition gives $\omega = w + u'_W$. If work has no direct marginal (dis)utility, this reduces to $\omega = w$, as in Jara-Díaz ([2007](#)), and the constant- ω mapping is exact while that interior condition continues to hold.

With fixed or bounded work hours, the same object is defined by the multiplier on the time constraint. In that case ω need not equal the wage, consistent with empirical values of time often being a fraction of the wage (Small, [2012](#)). The VTTS expressions in [Proposition 1](#) are evaluated at this reference allocation. In the bottleneck analysis of [Section 5](#), they are then held fixed within each class. Except in cases such as the flexible-labour case just described, this is a first-order, constant-marginal-value approximation: it does not re-optimize W , L , and τ^b at every realised wait. [Proposition 8](#) later allows the cumulative waiting cost to be nonlinear without claiming to solve that full allocation problem. Because preferences are quasilinear in the numeraire, the model does not capture income effects on the value of time.

A running example. Throughout, take a commuter who must reach work by $t^* = 9:00$ and can either drive or use an automated, passive service in which she can work or rest. The calibration in [Table 4](#) uses a resource value of time $\omega = \$20/\text{h}$, a driving-effort cost $e = \$6/\text{h}$, and an automated mode whose in-trip output reclaims $\rho = \$10/\text{h}$. The example is used to illustrate each result below.

4 The value of travel time and Mun’s rule

4.1 An endogenous in-trip output decomposition

The traveller’s within-trip allocation of passive time obeys a simple rule. A marginal unit of passive travel time can be spent working, yielding the net return $w\phi_j$, or resting, worth

$\psi_j \omega$ in leisure-equivalent terms. She picks the larger. Define the **net reclaimed value of passive time**

$$\rho_j \equiv \max\{w\phi_j, \psi_j \omega\}. \quad (4)$$

Proposition 1 (value of travel time with net reclaimed in-trip output). At a reference optimum, with ω equal to the relevant local shadow value of discretionary time and with ρ_j defined by Equation 4, the value of travel-time savings is

$$\text{VTTS}_j^A = \omega + e_j, \quad \text{VTTS}_j^P = \omega - \rho_j.$$

Active travel time costs the resource value of time *plus* effort; passive travel time costs the resource value of time *net of* the output (work or rest) reclaimed during travel. Equivalently, in the notation of Jara-Díaz, the value of time assigned to travel is $\text{VTAT}_j = \rho_j$ for passive modes and $\text{VTAT}_j = -e_j$ for active modes. The coefficients are linked by the maintained allocation structure rather than specified as unrelated mode-time parameters; ω , e_j , ϕ_j , and ψ_j still require measurement or estimation.

The proof is in [Appendix B](#). The proof uses the Kuhn–Tucker conditions for the in-trip allocation and an envelope argument for the marginal value of travel time. Equation 4 is the maximum over the two uses of passive time: a quiet, work-friendly vehicle (high ϕ_j) reclaims net value through work, and a comfortable but not work-friendly one (high ψ_j , low ϕ_j) through rest. The resulting passive VTTS is below ω and may be zero or negative if net reclaimed value equals or exceeds the opportunity cost of time. The bottleneck propositions use the stronger domain restriction $\alpha_{\text{eff}} = \omega - \rho > \beta_{\text{eff}} > 0$; Proposition 1 itself does not impose it. In the running example, the commuter’s driving time is worth $\text{VTTS}^A = \omega + e = \$26/\text{h}$ and her automated time $\text{VTTS}^P = \omega - \rho = \$10/\text{h}$, so the same minutes carry a lower conditional user cost when she can use them.

Figure 1 displays the components of VTTS for an illustrative set of modes ($\omega = w = \$20/\text{h}$). Active driving carries a value of time *above* the wage; work-friendly passive modes carry a value well below it.

A direct consequence is that mode shares respond to in-trip productivity even holding money price and clock time fixed. Figure 2 traces an illustrative binary logit between driving and a robotaxi as the robotaxi’s net work-productivity parameter ϕ rises, holding $\psi = 0.3$: its simulated share climbs from 59% to 83%, and the fleet-average value of time falls. These are scenario outputs, not estimated adoption shares. The same minutes of travel have a different conditional user cost as ϕ changes. The numerical assumptions are reported in [Appendix C](#).

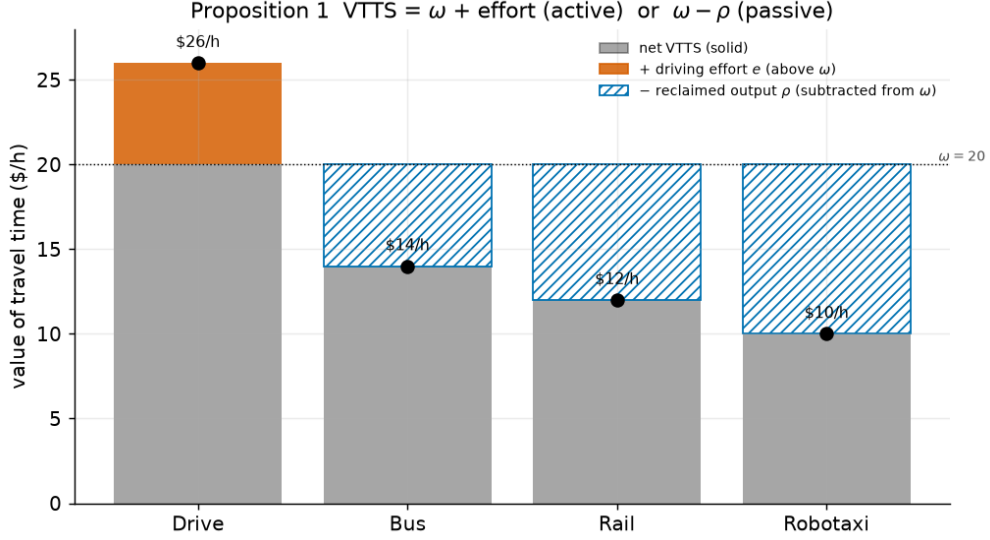


Figure 1: Value-of-travel-time components. The dotted line is the resource value ω . Driving adds an effort cost e above ω ; for passive modes the hatched cap is the reclaimed in-trip output $\rho = \max\{w\phi, \psi\omega\}$ subtracted from ω , leaving the solid net VTTS (markers). Passive-mode illustrations use $\psi = 0.3$; the robotaxi case matches the baseline $\phi = 0.5$ in Table 4.

4.2 Mun’s implicit-price rule as a corollary

Mun (2007) proposed that a trip option jointly produces a “trip output” and hedonic byproducts under a homogeneous-of-degree-one technology, and that the chosen option minimises an implicit unit price. In the present framework, Mun’s “trip output” is best read as an intermediate input into scheduled activity access, not as the final commodity itself. His implicit-price rule is recovered as a restricted special case, with the assumptions it requires made explicit.

Suppose the required scheduled activity access is y , produced with constant returns from q_j trips, $y = \zeta_j q_j$; travel involves a single time t_j per trip; hedonic byproducts $z_k = b_{kj} t_j q_j$ accrue linearly in travel time and can otherwise be bought at constant unit cost c_k ; and labour is flexible ($\omega = w$). Then the conditional cost of one unit of access through j is μ_j , where:

Corollary 1 (Mun’s rule). Under constant returns, fixed required output, linear hedonic co-production, and separable substitute production, the cost-minimising technology solves

$$\min_j \mu_j, \quad \mu_j = \frac{p_j + v_j t_j}{\zeta_j}, \quad v_j = w - \sum_k c_k b_{kj}.$$

Proposition 1 Endogenous in-trip output shifts mode choice

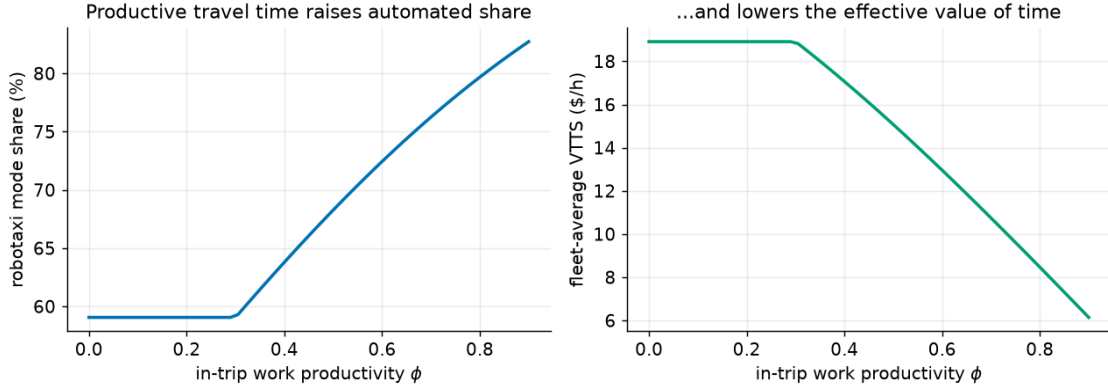


Figure 2: As in-trip work productivity ϕ rises for the automated mode, holding $\psi = 0.3$, the automated share rises and the effective value of time falls.

The common object is the marginal value reclaimed from a passive travel minute. In [Proposition 1](#), the traveller uses that minute either for work or for rest, so the reclaimed value is $\rho_j = \max\{w\phi_j, \psi_j\omega\}$. In Mun’s constant-returns case, the byproducts are jointly produced and the reclaimed value is additive, $\sum_k c_k b_{kj}$. Equivalently, if $R_j(T^P)$ denotes the optimised value of within-trip production, then $\rho_j = R'_j(T^P)$, with $R_j(T^P) = \max\{w\phi_j, \psi_j\omega\}T^P$ in Proposition 1 and $R_j(T^P) = (\sum_k c_k b_{kj})T^P$ in Mun. The “value of service time” is therefore $v_j = w - R'_j(T^P)$. Mun’s homogeneity assumption is the restriction under which a scalar unit implicit price exists.

Outside this case (fixed costs, increasing or decreasing returns, nonlinear scheduling, stochastic reliability, or congestion feedback), a scalar unit price need not summarise choice. The relevant object is then the full conditional cost function. The rest of the paper keeps this conditional-cost interpretation while specialising the timing problem to the bottleneck.

5 The production-bottleneck equilibrium

The production framing in [Section 3](#) and the value-of-time results in [Section 4](#) apply to any scheduled trip whose timing matters. This section makes travel times endogenous by specialising to the canonical single-destination commute, while leaving trip chains, multi-stop tours, and non-work travel for future research. In those cases, in-trip production may take a different form. The bottleneck comparative statics take as given both the technology classes and the number of travellers in each class, N_c . Endogenous adoption

of the passive technology is illustrated only in the partial-equilibrium mode-choice figure (Figure 2). The paper does not solve a joint adoption–congestion equilibrium in which fares or price wedges, queues, class costs, and adoption shares are determined together.

Here, consider the classic single bottleneck of capacity s (the maximum rate at which travellers can pass) serving a population that wishes to arrive at a common t^* (normalised to 0). A traveller who departs so as to arrive at time a incurs queueing time $T(a)$. The schedule penalty is β per unit of earliness if $a < 0$ and γ per unit of lateness if $a > 0$. The cost of a unit of *queueing* time is the value of that time from Section 4, denoted α_{eff} . A continuous FIFO queue forms when $\alpha_{\text{eff}} > \beta_{\text{eff}} > 0$ and $\gamma > 0$; if $\beta_{\text{eff}} \leq 0$, early arrival is no longer costly and the standard queue structure does not apply. The calibration maintains the common empirical ordering $\gamma > \alpha_{\text{eff}} > \beta_{\text{eff}}$, but only $\alpha_{\text{eff}} > \beta_{\text{eff}} > 0$ (with $\gamma > 0$) is required for the queue-forming equilibrium.

In this paper, a **pure point bottleneck** has fixed demand N , a single point bottleneck of capacity s , a common desired arrival time, FIFO service, *no* free-flow travel-time term, *no* spillback to upstream links, and *no* vehicle-operating, energy, emissions, crash-exposure, or curb/space externalities. The model distinguishes private generalised user cost from two model social-cost objects: without a toll, schedule and waiting costs are counted as real user losses; with the optimal toll, toll revenue is a transfer and only schedule cost remains. Throughout, the “optimal time-varying toll” is the fine toll normalised to zero at the edges of the arrival window; with fixed demand, adding a common constant would not change departure timing. Neither object is a complete economy-wide welfare measure. Section 8 separates the parameter–environment combinations that relax or qualify these assumptions (Table 2). The numerical illustrations use the stylised assumptions in Table 4 (Appendix C).

5.1 Reduction to an effective α - β - γ bottleneck

Let α denote the value of queueing time. By Proposition 1, a driver has $\alpha^{\text{drive}} = \omega + e$, since queueing is active and effortful. An automated or passive traveller has $\alpha^{\text{auto}} = \omega - \rho$, since queue time can partly be reclaimed.

Early arrival can also be productive when time before t^* can be used at the destination. In a reduced-form mapping, let a unit of early-arrival time yield net reclaimable value $\rho_c^E = \max\{\omega\phi_c^E, \psi_c^E\omega\}$, defined analogously to ρ_c but for early-arrival time at the destination. Then the effective earliness penalty is $\beta_{\text{eff}}^c = \beta^c - \rho_c^E$. Here β^c is the gross earliness penalty before the separately modelled reclaimed value; an empirically estimated schedule

coefficient that already nets out useful early time should not be adjusted again. The paper does not solve a separate destination-activity allocation problem. The maintained queueing regime permits $\rho_c^E = 0$ and requires $0 \leq \rho_c^E < \beta^c$. Late arrival remains feasible and carries the finite, non-reclaimable penalty γ^c . The production model therefore maps into an α - β - γ bottleneck with class-specific effective parameters. With several technologies or traveller groups, the model is a multi-class bottleneck.

Proposition 2 (reduced-form mapping). A traveller of technology c behaves in the bottleneck as an α - β - γ commuter with $\alpha_{\text{eff}}^c = \omega + e_c$ (active) or $\omega - \rho_c$ (passive), $\beta_{\text{eff}}^c = \beta^c - \rho_c^E$, and unchanged γ^c . Effort and net reclaimed output map into the effective scheduling parameters; the mapping does not by itself identify their separate components.

For a homogeneous population, the no-toll equilibrium is the textbook one. With $\delta_{\text{eff}} \equiv \beta_{\text{eff}}\gamma/(\beta_{\text{eff}} + \gamma)$, let $t_e > 0$ and $t_l > 0$ denote the earliness and lateness durations, so the arrival window is $[-t_e, t_l]$. The equilibrium generalised user cost, those two durations, and the waiting-time profile are

$$C^* = \frac{\delta_{\text{eff}}N}{s}, \quad t_e = \frac{\gamma}{\beta_{\text{eff}} + \gamma} \frac{N}{s}, \quad t_l = \frac{\beta_{\text{eff}}}{\beta_{\text{eff}} + \gamma} \frac{N}{s}, \quad T(a) = \frac{C^* - \text{sched}(a)}{\alpha_{\text{eff}}}. \quad (5)$$

Lemma 1 (boundary cost in the pure point bottleneck). In a single-class pure point bottleneck with fixed demand N , capacity s , FIFO service, no spillback, and a continuous queue-forming busy period, the first and last arrivals face no queue. Hence the arrival window and the common money cost are determined by schedule costs alone. In the piecewise-linear case,

$$\beta_{\text{eff}}t_e = \gamma t_l = C^*, \quad t_e + t_l = N/s.$$

With a smooth schedule cost $g(a)$, the corresponding boundary equations are $g(a_e) = g(a_l) = C^*$ and $a_l - a_e = N/s$. The value assigned to queue time determines only how a fixed money queue cost, $C^* - \text{sched}(a)$ or $C^* - g(a)$, is converted into physical waiting.

Assumption B1 (bottleneck environment). There are finitely many technology classes $c = 1, \dots, C$ with traveller masses $N_c > 0$, $\sum_c N_c = N$; a single point bottleneck of deterministic capacity s ; a common desired arrival time t^* ; first-in-first-out service and no spillback to upstream links; deterministic (non-stochastic) travel times; and class parameters satisfying $\alpha_{\text{eff}}^c > \beta_{\text{eff}}^c > 0$ and $\gamma^c > 0$.

Remark 1 (known equilibrium result). Lindsey (2004) gives conditions for heterogeneous-user bottlenecks under which generalised class costs and the aggregate

arrival pattern are unique. The inequality $\alpha_{\text{eff}}^c > \beta_{\text{eff}}^c$ is the single-class queue-forming condition used below; it is not asserted to be the full set of Lindsey’s heterogeneous-user conditions. The paper does not claim uniqueness of individual departure times within a class, nor uniqueness outside Lindsey’s conditions. Smith (1984) proves existence of the bottleneck equilibrium, and Daganzo (1985) proves uniqueness in the homogeneous case. This paper does not introduce a new equilibrium existence proof. Its contribution is to derive the effective parameters $(\alpha_{\text{eff}}^c, \beta_{\text{eff}}^c)$ from the traveller’s production problem.

5.2 Money-metric neutrality of in-vehicle productivity

Proposition 3 is the first comparative-statics result and the main neutrality result. It qualifies the intuition that productive in-vehicle time directly reduces congestion costs. The result compares the no-toll equilibrium, where waiting time is a real money-metric user loss in the model social-cost account, with the equilibrium under the optimal time-varying (fine) toll, where the toll removes the queue and toll revenue is treated as a transfer.

Proposition 3 (money-metric user-cost neutrality in the pure point bottleneck). Take the pure point bottleneck of Assumption B1 with a single homogeneous class, fixed demand, and $\alpha_{\text{eff}} > \beta_{\text{eff}}$. Value waiting time at α_{eff} and exclude spillback, vehicle-hours, energy, emissions, crash exposure, and curb/space. Then the following objects are independent of α_{eff} : private equilibrium generalised user cost $C^* = \delta_{\text{eff}}N/s$; the untolled model social cost $\delta_{\text{eff}}N^2/s$; the tolled model social cost $\frac{1}{2}\delta_{\text{eff}}N^2/s$ when toll revenue is treated as a transfer; and the optimal time-varying money toll. Lowering α_{eff} by making in-vehicle time more productive leaves these objects unchanged and raises maximum waiting time to $T_{\text{max}} = C^*/\alpha_{\text{eff}}$. The corresponding peak queue stock is $Q_{\text{max}} = sT_{\text{max}}$.

Boundary. The lower boundary is the relevant one. As α_{eff} falls, maximum waiting time rises to the finite value C^*/β_{eff} at the boundary. The upper ordering $\alpha_{\text{eff}} < \gamma$ is not required for neutrality; the result holds for every $\alpha_{\text{eff}} > \beta_{\text{eff}}$. At equality strict home-departure ordering reaches its boundary; below it the maintained orderly FIFO profile is invalid. The paper does not characterise the alternative equilibrium regime. Section 5.5 makes the corresponding distinction under a smooth scheduling cost.

Lemma 1 explains. The equilibrium cost is set by the *schedule* parameters $\beta_{\text{eff}}, \gamma$ and the demand-capacity ratio N/s ; α_{eff} governs only the conversion between money-metric cost and waiting *time*. A lower value of waiting time means travellers tolerate a longer wait to reach the same generalised user cost, so waiting grows just enough to hold cost constant.

Proposition 3 Neutrality of in-vehicle productivity in the pure bottleneck

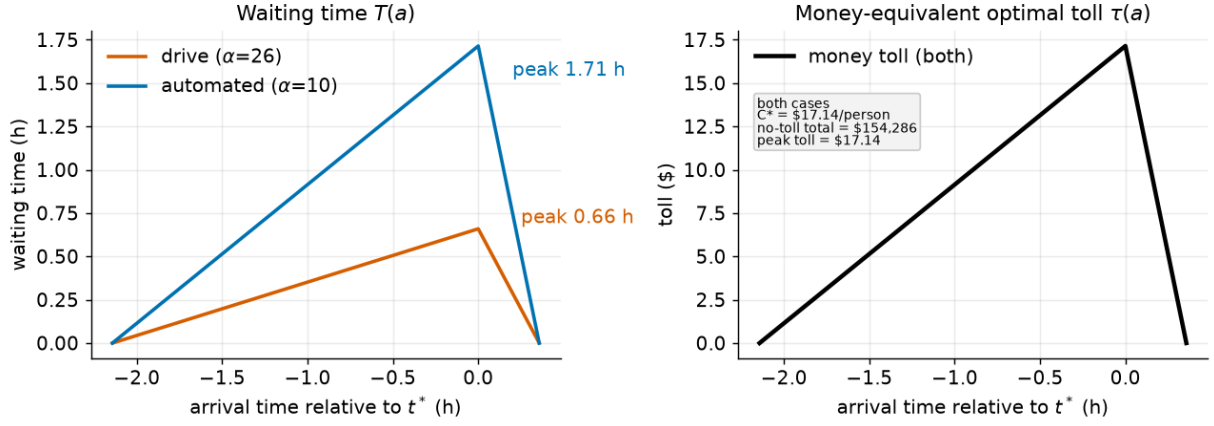


Figure 3: **Proposition 3.** Left: a lower value of in-vehicle time ($\alpha_{\text{eff}} = 10$, automated) produces a longer maximum wait than driving ($\alpha_{\text{eff}} = 26$); since $\beta_{\text{eff}} = 8 < \alpha_{\text{eff}}$, the orderly-queue regime holds for both homogeneous counterfactuals. Right: the optimal time-varying money toll is identical. Private generalised cost, the two model social-cost objects, and the toll are unchanged; waiting time and queue stock grow.

The money toll that replaces the wait is $\tau(a) = \alpha_{\text{eff}}T(a)$, and since $T(a) \propto 1/\alpha_{\text{eff}}$, the money toll is invariant. In the illustrative example the commuter’s automated trip waits 1.71 h at the maximum against 0.66 h driving, yet both face the same \$17.1 private generalised cost, the same \$154,300 untolled model social cost, and the same \$17.1 peak toll (Figure 3); only waiting time and queue stock differ. These magnitudes are not corridor estimates.

The α_{eff} -independence of C^* is a standard property of the deterministic bottleneck. The production accounting adds two things. First, α_{eff} is linked to net reclaimed output and driving effort at the maintained reference allocation (Section 4), so a homogeneous fixed-demand counterfactual can change waiting time and queue stock without changing equilibrium money-metric user cost. Second, the model separates three objects that can otherwise look similar: the value of waiting time α_{eff} , the effective early-arrival penalty β_{eff} , and the free-flow cost term under elastic demand. These objects have different model cost and pricing implications, as summarised in Table 2. The paper’s contribution is this common mapping and channel classification, together with the candidate restrictions in Section 7; a new bottleneck equilibrium or the first activity-dependent VTTS formula is not claimed.

The result limits what in-vehicle productivity alone can imply for automation. In the homogeneous fixed-demand case, if the relevant benefit is *in-vehicle productivity*, rather

Proposition 4 Productive early arrival shifts the fixed-duration window earlier

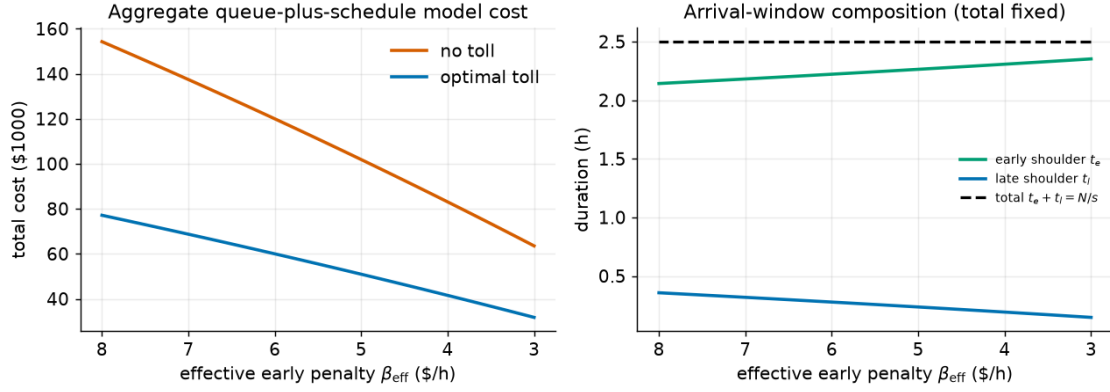


Figure 4: **Proposition 4.** A lower effective early-schedule penalty (useful early time, to the right) reduces both model social-cost objects and lengthens the early duration while shortening the late duration. The fixed-length arrival window shifts earlier and maximum waiting falls.

than cheaper schedule adjustment or added capacity, automation can increase waiting time, queue stock, and vehicle-hours without changing the model’s money-metric user cost or corrective money toll. Once the excluded physical and network externalities are priced, the full welfare effect need not be neutral and may turn negative. Physical congestion and money-metric user cost can therefore move in opposite directions. For example, spillback into upstream links would create delays outside the pure-point model cost object, so the neutrality result should not be read as a mixed-fleet or network-welfare claim.

5.3 Where model cost and pricing actually respond

The relevant model cost, surplus, and pricing effects instead run through two other parameter–environment combinations.

Proposition 4 (a lower effective early-arrival penalty shifts the arrival window earlier and lowers maximum waiting). A reduction in the reduced-form effective earliness penalty β_{eff} lowers $\delta_{\text{eff}} = \beta_{\text{eff}}\gamma/(\beta_{\text{eff}} + \gamma)$, and hence lowers both the untolled model social cost $\delta_{\text{eff}}N^2/s$ and the tolled model social cost $\frac{1}{2}\delta_{\text{eff}}N^2/s$. The early duration t_e lengthens, the late duration t_l shortens, and the total arrival-window duration remains $t_e + t_l = N/s$. For fixed α_{eff} , maximum waiting time falls with C^* , and the window shifts earlier while retaining its duration.

Unlike in-vehicle productivity in Proposition 3’s fixed-demand pure bottleneck, a lower effective early-arrival penalty reduces schedule-delay cost. Figure 4 shows the untolled

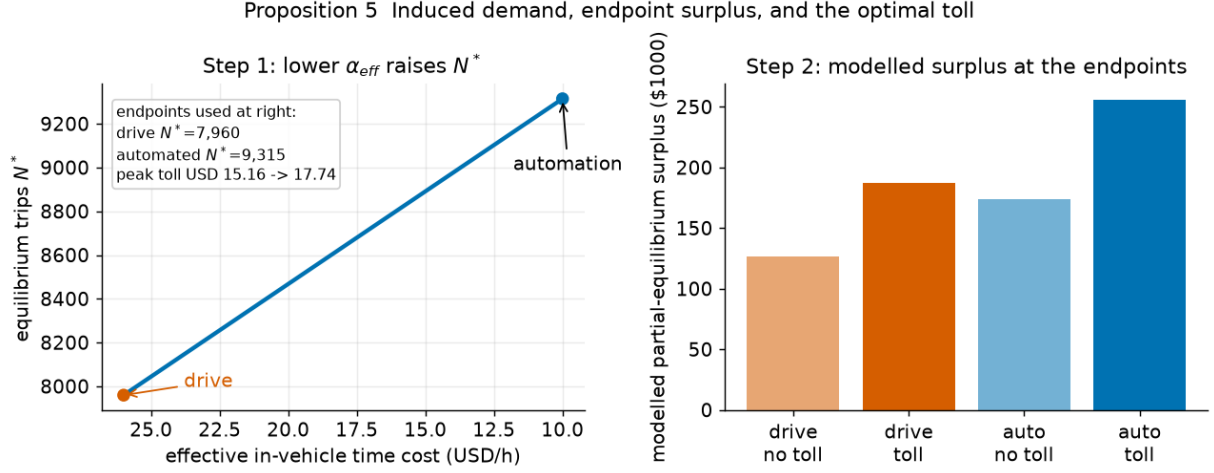


Figure 5: **Proposition 5.** Left: lower effective in-vehicle time cost induces equilibrium trips; the drive and automation endpoints are the illustrative cases evaluated on the right. The optimal peak toll rises with the induced volume. Right: model surplus at those endpoint equilibria, with and without the optimal toll.

model social cost falling by more than half, from about \$154,300 to \$63,500, as β_{eff} falls from 8 to 3. If the commuter's employer lets her start useful work when she arrives early, the reduced-form β_{eff} may fall and her contribution to the peak shifts earlier. This gives a production interpretation for flexible-arrival policies and destination amenities, such as workspaces or early-access facilities, provided the baseline β has not already netted out their value.

Proposition 5 (induced demand and the optimal toll under lower in-vehicle time cost). Let $b > 0$, $T_f \geq 0$, and $\bar{P} > \alpha_{\text{eff}} T_f$, and define the bottleneck component $C_B(N) \equiv \delta_{\text{eff}} N / s$. With inverse demand $P(N) = \bar{P} - bN$, generalised cost is $C(N) = C_B(N) + \alpha_{\text{eff}} T_f$. The interior no-toll equilibrium trip volume is $N^* = (\bar{P} - \alpha_{\text{eff}} T_f) / (b + \delta_{\text{eff}} / s)$, weakly increasing as α_{eff} falls and strictly increasing when $T_f > 0$. When $T_f > 0$, cheaper in-vehicle time *induces trips*, and the peak of the optimal time-varying toll, $C_B(N^*)$, *rises* (the schedule is $\tau(a) = C_B(N^*) - \text{sched}(a)$). The induced-demand effect runs through the free-flow term $\alpha_{\text{eff}} T_f$. The pure-queue effect in Proposition 3 does not induce trips. This is the standard elastic-demand mechanism operating through free-flow travel cost; with $T_f = 0$ automation leaves equilibrium volume unchanged.

Figure 5 (left) shows illustrative trip volume rising from about 7,960 to 9,315 as the in-vehicle time-cost parameter moves from the driving to the automated endpoint; the optimal peak toll rises from \$15.2 to \$17.7 because it is proportional to the induced volume. The right panel evaluates model surplus at those two endpoint equilibria. The time-varying toll removes the queue deadweight (the Vickrey result that the fine toll

leaves trip volume unchanged while toll revenue is a transfer), while the lower free-flow time cost raises gross trip surplus and total congestion. Together with [Proposition 3](#), the figure shows why the effect depends on the parameter–environment combination. The magnitudes are illustrative, not an automation forecast.

5.4 Sorting in the peak

When production technologies differ in the value of queue time but share schedule penalties, they sort within the peak.

Proposition 6 (interior sorting with fixed class masses). Under [Assumption B1](#), take common β, γ , two classes satisfying $\beta < \alpha_{\text{lo}} < \alpha_{\text{hi}}$, fixed masses $0 < N_{\text{lo}} < N$, and a single first-in-first-out queue. Let

$$a_{\min} = -\frac{\gamma}{\beta + \gamma} \frac{N}{s}, \quad a_{\max} = \frac{\beta}{\beta + \gamma} \frac{N}{s},$$

and define

$$x_e = \frac{\gamma}{\beta + \gamma} \frac{N_{\text{lo}}}{s}, \quad x_l = \frac{\beta}{\beta + \gamma} \frac{N_{\text{lo}}}{s}.$$

Then $0 < x_e < |a_{\min}|$ and $0 < x_l < a_{\max}$ automatically: the low- α class (automated/passive, with the lower value of waiting time) occupies the central interval $[-x_e, x_l]$, where waiting is longest and schedule delay is smallest, and the high- α class occupies the two shoulders. The mean absolute arrival deviation from t^* is strictly smaller for the low- α class. No further class-share condition is needed when both fixed class masses are positive; as $N_{\text{lo}} \downarrow 0$ both central pieces vanish, and as $N_{\text{lo}} \uparrow N$ both high- α shoulders vanish.

The waiting-time profile is the familiar piecewise-linear profile of a heterogeneous bottleneck ([Arnott et al., 1988](#); [Lindsey, 2004](#); [V. van den Berg & Verhoef, 2011](#)) ([Figure 6](#)). The production model supplies one interpretation of the class difference. With both fixed class masses positive, automated/passive travellers have a lower effective value of waiting time and occupy the centre, while active drivers take the shoulders. In the illustrative half-and-half case, the low- α class arrives on average 0.47 h from t^* against 1.42 h for the high- α class. These are conditional class outcomes, not the result of endogenous mode adoption.

Correlated schedule flexibility. The sorting result holds with common β, γ . This isolates the value-of-queue-time effect. If in-vehicle productivity is correlated with schedule flexibility, the late side is unchanged when γ is common. The early side depends on

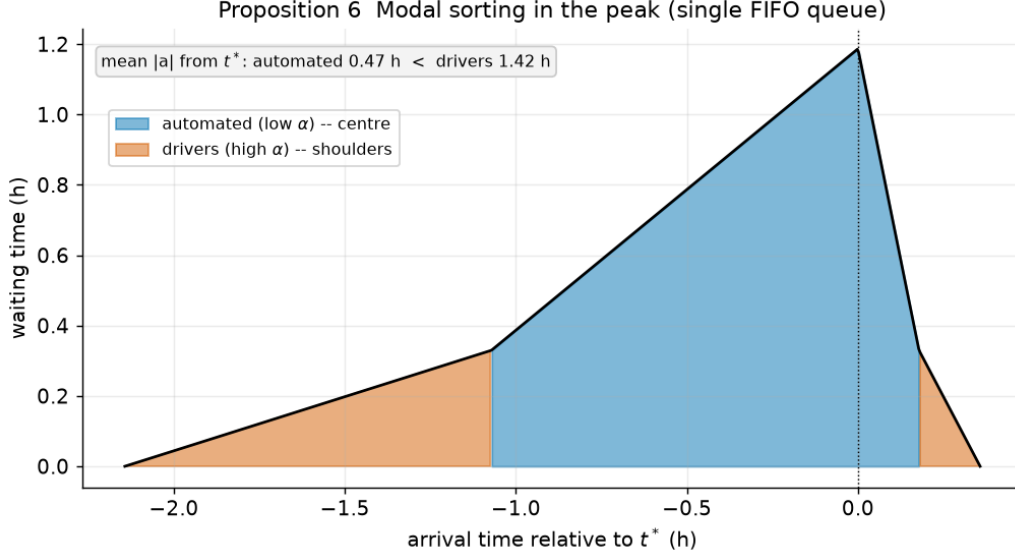


Figure 6: **Proposition 6.** With positive fixed masses and common schedule penalties, automated travellers (low value of waiting time) sort to the centre of the peak and drivers take the shoulders. Illustrative mean $|arrival from t^* |$: 0.47 h versus 1.42 h.

the slope ratio $\beta_{\text{eff}}^c / \alpha_{\text{eff}}^c$. The low- α class remains central on the early side if and only if $\beta_{\text{lo}} / \alpha_{\text{lo}} > \beta_{\text{hi}} / \alpha_{\text{hi}}$, that is, if in-vehicle productivity lowers the value of queue time proportionally more than flexibility lowers the early penalty; when flexibility dominates, the early shoulder inverts. At the calibration, this is the difference between, say, $\beta_{\text{lo}} = 6$ (slopes 0.60 vs 0.31, sorting intact) and $\beta_{\text{lo}} = 2$ (slopes 0.20 vs 0.31, early shoulder inverted). The central-band prediction therefore holds when β_{eff} is held fixed. Empirical tests must distinguish the value of queueing time from schedule flexibility before attributing sorting to α_{eff} alone.

5.5 Neutrality beyond the pure point bottleneck

Proposition 3 is stated for the piecewise-linear α - β - γ bottleneck and for the queue-forming regime $\alpha_{\text{eff}} > \beta_{\text{eff}}$. Two questions follow. Is the neutrality an artefact of the kink in the scheduling cost? What governs the boundary of the regime, where in-vehicle time becomes productive enough that its value falls below the early-schedule penalty? For automation applications, this boundary marks the limit of the strictly ordered continuous-departure regime. The section answers both questions by relaxing the assumptions one at a time. First, the kinked schedule cost is replaced with a smooth scheduling cost. Second, the constant marginal value of waiting time is replaced with a duration-dependent cumulative

waiting cost.

Replace the kinked penalty with a strictly convex, C^2 scheduling cost $g(a)$ with $g(0) = 0$, $g'(0) = 0$, and $g'' > 0$. This is the smooth “slope” form of Vickrey and of Fosgerau and Small (2017), of which the α - β - γ kink is the piecewise-linear limit. The early-schedule slope is then the local marginal cost $-g'(a)$, largest at the early edge of the arrival window; the constant β is recovered in the piecewise-linear limit. Lemma 1 still determines the window and money cost from the boundary arrivals; the new issue is whether the implied waiting-time profile respects FIFO.

Proposition 7 (neutrality under smooth scheduling; a local-slope boundary). In the pure point bottleneck with strictly convex C^2 scheduling cost g , let $a_e < 0 < a_l$ be the edges of the arrival window and define $\alpha_{\text{crit}} \equiv -g'(a_e)$. For $\alpha_{\text{eff}} > \alpha_{\text{crit}}$ a queue-forming equilibrium exists in which the window and the private equilibrium cost $C^* = g(a_e) = g(a_l)$ are determined by the boundary equations $g(a_e) = g(a_l)$ and $a_l - a_e = N/s$ alone. Hence C^* , the untolled model social cost NC^* , the tolled model social cost $s \int_{a_e}^{a_l} g(a) da$, and the zero-boundary normalised time-varying toll are **independent of** α_{eff} , while maximum waiting time is C^*/α_{eff} , rising as α_{eff} falls up to the *finite* value C^*/α_{crit} . The threshold $\alpha_{\text{crit}} = -g'(a_e)$ is the local early-schedule slope at the window edge, the smooth counterpart of the role β_{eff} plays in the kinked model. Unlike the primitive β_{eff} , it depends on the demand-capacity ratio; for a piecewise-quadratic schedule, $\alpha_{\text{crit}} = \sqrt{k_e k_l} (N/s) / (1 + \sqrt{k_l/k_e})$. The strict-ordering boundary therefore rises with corridor congestion. At $\alpha_{\text{eff}} = \alpha_{\text{crit}}$, the home-departure map has zero derivative at the early edge: strict ordering reaches its boundary, but weak FIFO is not reversed. For $\alpha_{\text{eff}} < \alpha_{\text{crit}}$, the maintained single-pass profile reverses home-departure ordering near that edge and is not an equilibrium. The proposition does not characterise the alternative regime or rule out other equilibria.

The smooth case gives two implications. First, the neutrality result does not depend on the kinked α - β - γ schedule. Second, smoothing does not eliminate the strict-ordering boundary. The boundary depends on the demand-capacity ratio, and a larger N/s raises α_{crit} . At equality the continuous-departure representation reaches a weak-ordering boundary; below it the maintained single-pass profile fails. Bunching or another non-orderly representation is possible, but neither its existence nor its welfare is established here. In-vehicle automation can push α_{eff} toward α_{crit} , and the orderly-queue neutrality result applies only above it.

The boundary-cost argument also covers the case in which the marginal cost of a traveller’s wait rises with accumulated waiting duration. This captures the idea that a long

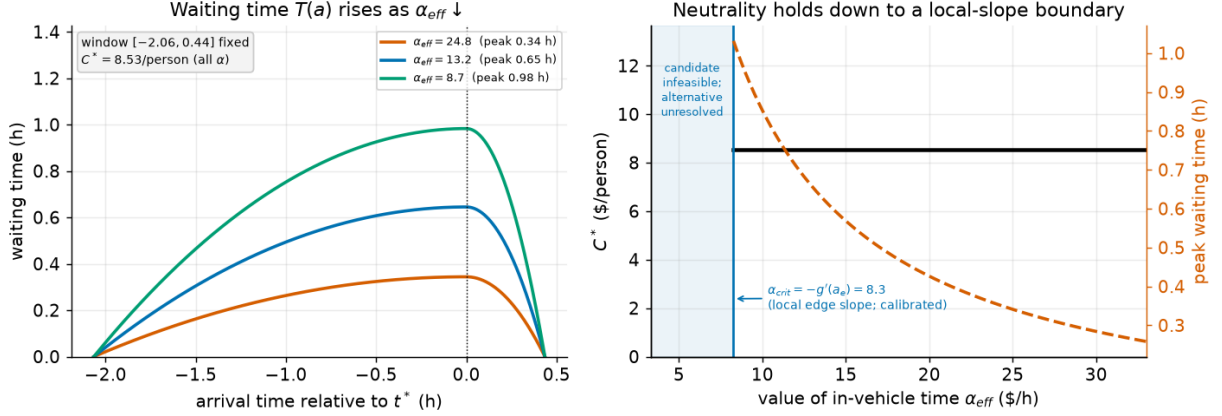


Figure 7: **Proposition 7.** Left: with a smooth, strictly convex schedule, a lower value of in-vehicle time lengthens waiting while the arrival window and C^* remain fixed. Right: user cost is flat across the strictly ordered regime while maximum waiting rises as α_{eff} falls to the local-slope boundary. The shaded lower region marks only where the maintained candidate profile is infeasible; the replacement equilibrium is not characterised.

stop-and-go wait may be more effortful than a short wait; it is not a model of the marginal cost depending on system queue stock.

Proposition 8 (conditional neutrality under duration-dependent waiting cost).

Let $K : [0, \infty) \rightarrow [0, \infty)$ be a continuously differentiable cumulative money cost of a wait of length T , with $K(0) = 0$ and $K'(T) > \beta_{\text{eff}}$ for every $T \geq 0$. Suppose a queue-forming FIFO equilibrium satisfying the boundary conditions of Lemma 1 exists. In the kinked bottleneck the private equilibrium cost $C^* = \delta_{\text{eff}}N/s$, the untolled and tolled model social costs, and the zero-boundary normalised time-varying toll are unchanged from Proposition 3; waiting time is $T(a) = K^{-1}(C^* - \text{sched}(a))$, generally nonlinear in arrival time. If K is also convex and $\alpha_0 \equiv K'(0)$, then $K(T) \geq \alpha_0 T$ and waiting is no longer than under the constant-marginal-cost baseline $K_0(T) = \alpha_0 T$, with a strict reduction wherever $K(T) > \alpha_0 T$.

By Lemma 1, the equilibrium cost is fixed by the boundary arrivals, who face *no* wait, so $\beta_{\text{eff}}t_e = \gamma t_l = C^*$ holds whatever the cumulative cost of waiting. The strict derivative condition gives $T'(a) = \beta_{\text{eff}}/K'(T) < 1$ on the early side and therefore preserves strict home-departure ordering. Under the additional convexity condition, the tangent inequality $K(T) \geq K(0) + K'(0)T = \alpha_0 T$ gives the stated comparison with the constant-marginal-cost baseline. In the illustrative case, with the marginal cost rising linearly in the traveller's own accumulated wait, $K(T) = \alpha_0 T + \frac{1}{2}\eta T^2$ ($\alpha_0 = 20$, $\eta = 8$), the max-

imum wait falls from 0.86 to 0.75 h relative to $K_0(T) = \alpha_0 T$. The private generalised cost, model social-cost objects, and corrective toll remain as in [Proposition 3](#). Taken together, [Proposition 7](#) and [Proposition 8](#) show that the neutrality argument rests on the boundary structure of the bottleneck; they do not provide a general existence result outside their stated regimes.

6 The value of reliability with production

Travel time is uncertain. Write random free-flow travel time as $\tilde{T} = \mu_T + \sigma_T X$, where the standardised delay X has mean zero, unit variance, and a fixed cumulative distribution Φ satisfying the regularity conditions in Fosgerau and Karlström (2010). With optimal departure, their result makes expected cost linear in μ_T and σ_T . Substituting the effective parameters of Section 5:

Corollary 2 (value of reliability with production). Under the Fosgerau–Karlström conditions with technology-specific effective parameters, let $\mathcal{C}_R$ denote minimised expected trip cost and let $\mathcal{R}_\Phi(\beta_{\text{eff}}, \gamma)$ denote their reliability coefficient for the fixed standardised distribution Φ . Then

$$\mathcal{C}_R = \alpha_{\text{eff}} \mu_T + \sigma_T \mathcal{R}_\Phi(\beta_{\text{eff}}, \gamma),$$

so in-trip productivity lowers the value assigned to mean travel time through α_{eff} . Holding Φ fixed, a lower effective early-arrival penalty lowers the value of reliability $\mathcal{R}_\Phi(\beta_{\text{eff}}, \gamma)$ through β_{eff} .

The corollary is a direct parameter substitution into the Fosgerau–Karlström result, not a new reliability theorem. If the in-trip production decision is made *after* the delay is realised, so the traveller works because she is delayed, the linear mean–standard-deviation form no longer holds exactly because recourse creates an option value. That case is outside the scope of the corollary.

7 Testable restrictions

A production model has empirical content only if it rules out some patterns in observed choice. The production model is a constrained special case of a random-utility or multiple-discrete-continuous model. R1–R3 are candidate cross-equation restrictions that a generic

such model, carrying free per-mode time coefficients, need not satisfy. They are not identification theorems: each requires suitable variation and controls. R4 is a joint revealed-preference diagnostic for constructed prices and rationality, not a production-specific restriction.

R1–R3 can be read as cross-equation restrictions in a mixed-logit or multiple-discrete-continuous model. R1 restricts how telework feasibility interacts with passive and active time. R2 restricts time-use responses to fare and travel-time shocks. R3 restricts the arrival-time distribution by mode, conditional on schedule flexibility. R4 asks a different question: whether a specified construction of full prices and bundles is jointly consistent with utility maximisation.

R1 (active/passive asymmetry). Because $\text{VTTS}_j^P = \omega - \max\{w\phi_j, \psi_j\omega\}$, the value of passive travel time weakly falls with effective in-trip productivity. Because $\text{VTTS}_j^A = \omega + e_j$, active driving time is unaffected by that productivity. This sign-and-zero pattern is the restriction. Under rigid hours, where ω is fixed, VTTS^P falls with the wage only in the work-dominant regime; under flexible hours, where $\omega = w$, the net comparative static need not be negative. The relevant empirical object is effective in-trip productivity, the interaction of job-level telework feasibility with mode-level work-friendliness. The production model predicts a weakly negative interaction for passive time, strictly negative only where in-trip work binds ($w\phi_j > \psi_j\omega$) and zero for rest-dominant modes, and a zero interaction for active driving time. This test should condition on wage, schedule flexibility, comfort, and selection into modes; without excluded variation, the sign pattern does not identify ϕ_j separately.

R2 (money–time non-equivalence). In a generic quasilinear generalised-cost model, a fare change and a time change of equal money value are interchangeable. In this model, a travel-time increase expands passive time. In a work-dominant regime ($w\phi_j > \psi_j\omega$) it raises in-trip work, and in a rest-dominant regime it raises in-trip rest instead, while a pure *fare* shock does not directly change either use of time. With time-use data, the two shocks are distinguishable. The test is therefore a restriction on time-use responses, with the work-hours response informative only where $w\phi_j > \psi_j\omega$; absence of a work-hours response alone, in a rest-dominant regime, does not reject the mechanism. This restriction rests on quasilinearity together with the time-budget bound on τ^b ; under income effects the model itself could predict some fare response, so the test is against a generic quasilinear generalised-cost model.

R3 (peak sorting). [Proposition 6](#) predicts that low-value-of-queue-time technologies arrive closer to t^* . The prediction can be tested with departure and arrival microdata that

include mode. Tests must account for schedule flexibility and selection into automated modes, for example with traveller fixed effects. Otherwise, sorting by a low value of queueing time can be confused with sorting by flexible schedules. As the discussion after [Proposition 6](#) shows, the central-band prediction holds only when the fall in α_{eff} dominates the fall in β_{eff} .

R4 (joint revealed-preference diagnostic). If an empirical application first defines comparable produced bundles, linear feasible budget sets, and full price vectors, the resulting choices can be checked against the generalised axiom of revealed preference and associated Afriat inequalities (Afriat, [1967](#); Varian, [1982](#)). The present paper does not derive those objects from its discrete technology set, so R4 is not asserted as a theorem or a distinctive implication of travel production. It is a joint diagnostic of the application’s constructed prices, maintained budget representation, and rationality. With few modes or little independent price variation, such a check has limited power, and non-rejection is weak evidence.

8 Model cost, surplus, and pricing

Pricing and model-cost effects depend jointly on the parameter that changes and the environment in which it changes. [Table 2](#) therefore separates the fixed-demand pure-point experiment from the elastic-demand experiment with free-flow time. “User cost” means the traveller’s money-metric generalised cost. “Model social cost” aggregates schedule and waiting losses across travellers and treats toll revenue as a transfer, while still excluding operator resources, employer-side output accounting, spillback, vehicle-hours, energy, emissions, crashes, and space use. “Model surplus” adds the inverse-demand benefit used in [Proposition 5](#). These are transparent partial-equilibrium accounting objects, not complete economy-wide welfare measures. With income effects, the signs are local at the maintained marginal utility of money.

Table 2: Parameter–environment combinations and their model cost, surplus, and pricing implications.

Parameter– environment case	Parameter	Equilibrium effect	Demand effect	Toll effect	Model cost/surplus effect
Waiting-time value, fixed $N, T_f = 0$	$\alpha_{\text{eff}} \downarrow$	waiting and queue stock rise; user cost unchanged	none	unchanged	both model social-cost objects unchanged; excluded physical effects can overturn neutrality
Effective early-arrival penalty, fixed N	$\beta_{\text{eff}} \downarrow$	arrivals shift earlier; maximum wait and user cost fall	none in fixed- N case	peak toll falls	model social cost falls
In- vehicle/free- flow cost with elastic demand, $T_f > 0$	$\alpha_{\text{eff}} T_f \downarrow$	lower fixed trip cost	trips rise	peak toll rises with induced volume	model surplus rises in the stated partial equilibrium; congestion also rises
Capacity	$s \uparrow$	user cost falls	trips rise under elastic demand	toll per user falls	model surplus improves; omitted capacity costs remain outside the account

First, the corrective toll prices the congestion externality one traveller imposes on others within the model’s user-cost account. Waiting minutes are valued net of reclaimed output through α_{eff} ; schedule effects enter separately. The neutrality entry in Table 2 applies only to the homogeneous fixed-demand pure-point objects. Once excluded physical externalities or employer-side output are added, the full sign may differ. In the pure bottleneck, the money toll is invariant to α_{eff} (Proposition 3); under elastic demand with free-flow time it rises as the lower fixed trip cost induces volume (Proposition 5).

Second, productive in-vehicle time is not a substitute for capacity expansion. Capacity expansion lowers N/s and thus C^* ; a lower α_{eff} in the pure bottleneck leaves this cost unchanged while lengthening waiting and increasing queue stock. A planner who treats productive travel time as equivalent to capacity will overstate that channel’s congestion benefit.

Third, policies that reduce the effective schedule-delay cost (flexible arrival windows or destination amenities that make early arrival useful, Proposition 4) can complement pricing. The result is a reduced-form β_{eff} comparative static, not a complete model of workplace scheduling reform.

Fourth, an economy-wide evaluation of automation requires more than the objects solved here. Within the partial-equilibrium account, a lower free-flow time cost raises trip surplus (Proposition 5), a lower value of waiting can increase physical congestion (Proposition 3), and elastic demand raises volume. Full welfare additionally depends on employer output, operator resources, technology costs, capacity, occupancy and empty travel, physical externalities, and the distributional incidence of prices.

9 Conclusion

Automation, ride-hailing, and other passenger services can make travel time usable for work or rest. This paper models that fact by treating travel as production. The final commodity is scheduled activity access, a trip is an intermediate input, and a mode, route, departure time, or vehicle technology is a production plan that combines money, active and passive time, effort, schedule delay, and in-trip output.

The production formulation separates components of the value of travel time at a reference allocation. Active time is valued at the resource value of time plus effort. Passive time is valued at the resource value of time less net work or rest reclaimed during travel. Mun’s implicit-price rule is recovered under constant returns and additive byproducts.

Embedding these effective parameters in a Vickrey–Arnott–de Palma–Lindsey bottleneck then links the accounting to known bottleneck pricing and sorting results. In a homogeneous fixed-demand pure point bottleneck, a lower value of in-vehicle time is neutral for private generalised user cost, the stated model social-cost objects, and the optimal time-varying toll, although it lengthens waiting and increases queue stock. A lower reduced-form early-arrival penalty lowers schedule-delay cost and shifts the peak earlier. Lower free-flow travel cost under elastic demand induces trips and raises the optimal toll. With positive fixed class masses and common schedule penalties, travellers with lower values of waiting time sort toward the centre of the peak.

The analysis has several main limitations. The bottleneck results use the canonical single-destination commute, a common desired arrival time, FIFO service, and no spillback. The pure point model also excludes operating, energy, emissions, safety, and curb-space externalities. It holds technology classes and their masses fixed and therefore does not solve joint mode adoption and congestion. The money-metric model social-cost objects do not close the employer-output, operator-resource, or distributional accounts.

The preference structure is also narrow. Model surplus and testability are stated under quasilinear money-metric preferences, so the analysis does not capture income effects. Outside the flexible-labour case that keeps ω constant exactly, the bottleneck uses a local constant-marginal-value approximation. The neutrality result is stated only for the strictly ordered queue-forming regime; at the boundary strict ordering fails, and below it the paper neither characterises the alternative regime nor proves that it is bunching. The in-trip productivity rate ϕ_j is a net marginal rate that absorbs incremental onboard work effort and task frictions rather than identifying them separately. The reliability corollary is a parameter substitution into Fosgerau–Karlström and does not cover recourse after delays are realised.

These limitations point to natural extensions. First, the alternative regime below the strict-ordering boundary should be characterised rather than assumed to be bunching. Second, the pure-point bottleneck should be embedded in a corridor or network model with spillback and physical externalities. Third, the linear net in-trip productivity rate should be replaced, where data permit, with a time-use technology that allows setup costs, task indivisibilities, and separately measured onboard effort. Fourth, R1–R3 should be estimated with time-use and mode data; R4 requires a separately justified bundle and budget construction. The main implication is that productive travel time is not a generic congestion benefit. Its model cost and pricing effects depend on which parameter changes and in which environment.

Acknowledgements

This work was supported by the Singapore National Research Foundation under its Campus for Research Excellence and Technological Enterprise (CREATE) programme. The author gratefully acknowledges the institutional support of the Economic Growth Centre, Nanyang Technological University, Singapore, where this work was carried out.

Declaration of Generative AI Use

The author used generative AI tools from the OpenAI GPT/Codex and Anthropic Claude model families for language editing, manuscript organisation, and coding assistance. These tools were used as aids to author-led writing and programming. They were not used to generate research data, fabricate references, or make autonomous theoretical or empirical claims. Refine.ink was used to check the paper for consistency and clarity. The author reviewed, revised, and verified all AI-assisted material and is solely responsible for the content of the paper.

Appendix A: Notation

Table 3 collects the main notation used in the paper. Class-specific versions use a class superscript or subscript where needed.

Table 3: Main notation.

Symbol	Meaning	Units or domain
$j \in J$	technology, mode, route, fare class, vehicle type, or compound itinerary	finite set
B	gross benefit of completing the scheduled activity	money
V_j, V_0, C_j	optimised plan value, outside-option value, and conditional user cost $V_0 - V_j$	money
y, q_j, ζ_j	scheduled-access output, number of trips, and access output per trip in Corollary 1	output; trips; output/trip
G	numeraire consumption	money
I, w, p_j	non-labour income, wage, and monetary price of technology j	money; money/time; money
$\bar{T}, W, L$	time endowment, market work, and non-travel leisure or home time	time
T_j^A, T_j^P	active and passive travel time under technology j	time
τ^b, τ^r	passive travel time allocated to in-trip work and rest	time
ϕ_j	net effective in-trip work-productivity factor during passive travel	$[0, 1]$
ψ_j	leisure-equivalent value of in-trip rest	$[0, 1]$

Symbol	Meaning	Units or domain
e_j	marginal effort cost of active travel time	money/time
ω	resource value of discretionary time	money/time
ρ_j	net value reclaimed from a unit of passive travel time, $\max\{w\phi_j, \psi_j\omega\}$	money/time
ϕ_c^E, ψ_c^E	early-arrival work productivity and rest substitution for class c	$[0, 1]$
ρ_c^E	net value reclaimed from a unit of early-arrival time for class c , $\max\{w\phi_c^E, \psi_c^E\omega\}$	money/time
$\text{VTTS}_j^A, \text{VTTS}_j^P$	value of active and passive travel-time savings	money/time
$t^*, a, T(a)$	desired arrival time, arrival deviation from t^* , and waiting time	clock time; time; time
$Q(a) = sT(a)$	point-bottleneck queue stock associated with waiting time $T(a)$	travellers
N, N_c, s	total demand, class- c demand, and bottleneck capacity	travellers; travellers; travellers/time
α_{eff}	effective value of queueing time	money/time
$\beta_{\text{eff}}, \gamma$	net effective earliness penalty and lateness penalty	money/time
δ_{eff}	bottleneck schedule-cost coefficient, $\beta_{\text{eff}}\gamma/(\beta_{\text{eff}} + \gamma)$	money/time
C^*	private equilibrium generalised user cost per traveller in the bottleneck	money
$\tau(a)$	time-varying money toll at arrival time a	money

Symbol	Meaning	Units or domain
T_f	free-flow travel time	time
$P(N) = \bar{P} - bN$	inverse demand for trips, with intercept $\bar{P}$ and slope b	money
$C_B(N)$	bottleneck component of private generalised cost, $\delta_{\text{eff}}N/s$	money
$g(a)$	smooth scheduling cost at arrival deviation a	money
a_e, a_l	early and late edges of the smooth arrival window	time
α_{crit}	smooth-scheduling boundary for orderly queueing, $-g'(a_e)$	money/time
$K(T)$	cumulative money cost of a wait of duration T in Proposition 8	money
η	curvature of cumulative waiting cost in the Proposition 8 illustration	money/time ²
μ_T, σ_T, X, Φ	mean, standard deviation, standardised delay, and its fixed distribution	time; time; unitless; distribution
$\mathcal{C}_R, \mathcal{R}_\Phi(\beta_{\text{eff}}, \gamma)$	minimised expected reliability-case cost and reliability coefficient	money; money/time

Appendix B: Proofs

Proof of Proposition 1. Substitute Equation 1 and Equation 2 into Equation 3:

$$U_j = I - p_j + w(W + \phi_j \tau^b) + u_L(\bar{T} - T_j^A - T_j^P - W + \psi_j(T_j^P - \tau^b)) + u_W(W) - e_j T_j^A.$$

The Kuhn–Tucker condition for τ^b equates the net onboard-work return $w\phi_j$ and $\psi_j u'_L(\cdot)$ when both in-trip work and in-trip rest are used; otherwise the optimum is at the work

or rest corner. At the optimum, therefore, the marginal unit of passive travel time is assigned to the higher-valued use. When W is interior, the first-order condition is $w - u'_L(\cdot) + u'_W(W) = 0$, giving $\omega \equiv u'_L(X^*) = w + u'_W$, where $X^* = L^* + \psi_j \tau^{r*}$ is the optimised effective-leisure argument. When work hours are fixed or bounded, the corresponding multiplier on the labour constraint defines the same local shadow value ω ; the envelope calculation below uses that shadow value and does not require an interior choice of W . Separate the direct opportunity cost $-\omega$ of removing a unit from non-travel time from the value of assigning that added passive minute within the trip. The latter one-sided envelope is $\max\{w\phi_j, \psi_j\omega\} = \rho_j$: at the work corner the upper bound on τ^b expands and contributes $w\phi_j$, while at the rest corner the contribution is $\psi_j\omega$. The two values coincide at the switching condition, so the local VTTS is continuous there. Hence $\partial U_j / \partial T_j^P = -\omega + \rho_j$, and the money-metric user cost of passive travel time is $\text{VTTS}_j^P = \omega - \rho_j$. For active time, no in-trip output is possible and effort is incurred, so $\partial U_j / \partial T_j^A = -\omega - e_j$ and $\text{VTTS}_j^A = \omega + e_j$. This is a reference-allocation envelope; treating it as constant in the bottleneck is exact only under the conditions stated in Section 3 and otherwise is a first-order approximation. ■

Proof of Corollary 1. With $y = \zeta_j q_j$ fixed, the number of trips is $q_j = y / \zeta_j$. Each trip costs p_j in money and t_j in time valued at w (flexible labour), and produces hedonic byproducts $z_k = b_{kj} t_j q_j$ that would otherwise be bought at unit cost c_k . Producing one trip therefore *saves* $\sum_k c_k b_{kj} t_j$ of substitute production. The net cost of the time input per trip is $(w - \sum_k c_k b_{kj}) t_j = v_j t_j$. The conditional cost of y units of access is $C_j(y) = (p_j + v_j t_j) q_j = \frac{p_j + v_j t_j}{\zeta_j} y$, and the cost-minimising technology minimises the unit cost $\mu_j = (p_j + v_j t_j) / \zeta_j$. ■

Proof of Proposition 2. Immediate from Proposition 1: the cost of a unit of waiting time is $\omega + e_c$ if active and $\omega - \rho_c$ if passive. By the reduced-form definition in Section 5, separately modelled net reclaimed early-arrival value reduces gross β^c by ρ_c^E ; lateness is assumed non-reclaimable. Substituting these effective values into the α - β - γ cost function yields a standard bottleneck for each class. ■

Proof of Lemma 1. In a pure point bottleneck with no initial queue, no spillback, and a continuous busy period, waiting is zero at the start and at the end of the busy period. Thus the first and last arrivals pay only schedule cost. Since all used arrival times have the same private generalised cost C^* in equilibrium, the boundary schedule costs equal C^* . In the piecewise-linear case this gives $\beta_{\text{eff}} t_e = \gamma t_l = C^*$, and serving N travellers at capacity s gives $t_e + t_l = N/s$. With smooth schedule cost g , the same boundary condition is $g(a_e) = g(a_l) = C^*$ and the service-duration condition is $a_l - a_e = N/s$.

Interior waiting fills the gap between boundary and local schedule cost: with constant waiting-time value, $\alpha_{\text{eff}}T(a) = C^* - \text{sched}(a)$ or $\alpha_{\text{eff}}T(a) = C^* - g(a)$; with duration-dependent cumulative waiting cost, the same money gap is delivered by $K(T(a))$. In all cases the boundary conditions determine money-metric user cost before the waiting-cost function converts it into a duration and associated queue stock. ■

Proof of Proposition 3. By Lemma 1, $\beta_{\text{eff}}t_e = \gamma t_l = C^*$ and $t_e + t_l = N/s$, giving $t_e = \frac{\gamma}{\beta_{\text{eff}} + \gamma} \frac{N}{s}$, $t_l = \frac{\beta_{\text{eff}}}{\beta_{\text{eff}} + \gamma} \frac{N}{s}$, and $C^* = \frac{\beta_{\text{eff}}\gamma}{\beta_{\text{eff}} + \gamma} \frac{N}{s} = \delta_{\text{eff}}N/s$. None of t_e, t_l, C^* depends on α_{eff} . Waiting time solves $\alpha_{\text{eff}}T(a) + \text{sched}(a) = C^*$, so $T(a) = (C^* - \text{sched}(a))/\alpha_{\text{eff}}$, with maximum $T_{\text{max}} = C^*/\alpha_{\text{eff}}$ at $a = 0$. Thus waiting duration and queue stock $Q(a) = sT(a)$ scale as $1/\alpha_{\text{eff}}$. The optimal time-varying toll replaces the waiting-cost component, $\tau(a) = \alpha_{\text{eff}}T(a) = C^* - \text{sched}(a)$, which is independent of α_{eff} ; the peak toll is C^* . Private generalised cost is C^* per traveller in both regimes. The untolled model social cost is $\delta_{\text{eff}}N^2/s$, because waiting is a real money-metric user loss in this account; with the optimal toll it is $\frac{1}{2}\delta_{\text{eff}}N^2/s$, because the queue is eliminated, toll revenue is a transfer, and schedule delay remains. These objects are independent of α_{eff} but are not a complete economy-wide welfare account. The claim is made for $\alpha_{\text{eff}} > \beta_{\text{eff}}$; at equality strict ordering reaches its boundary, and below it this profile is invalid. ■

Proof of Proposition 4. $\delta_{\text{eff}} = \beta_{\text{eff}}\gamma/(\beta_{\text{eff}} + \gamma)$ is increasing in β_{eff} (its derivative is $\gamma^2/(\beta_{\text{eff}} + \gamma)^2 > 0$). The untolled model social cost $\delta_{\text{eff}}N^2/s$ and the tolled model social cost $\frac{1}{2}\delta_{\text{eff}}N^2/s$ are increasing in δ_{eff} , hence decreasing as β_{eff} falls. The early duration $t_e = \frac{\gamma}{\beta_{\text{eff}} + \gamma} \frac{N}{s}$ lengthens, while $t_l = \frac{\beta_{\text{eff}}}{\beta_{\text{eff}} + \gamma} \frac{N}{s}$ shortens by the same amount because $t_e + t_l = N/s$. Since $T_{\text{max}} = C^*/\alpha_{\text{eff}}$, maximum waiting also falls for fixed α_{eff} . ■

Proof of Proposition 5. Let gross surplus be the area under inverse demand, $\int_0^N P(n) dn = \bar{P}N - \frac{1}{2}bN^2$. The private generalised cost of a trip is $C(N) = C_B(N) + \alpha_{\text{eff}}T_f$ (waiting + schedule + free-flow). *No-toll equilibrium:* travellers enter until $P(N) = C(N)$, i.e. $\bar{P} - bN = \delta_{\text{eff}}N/s + \alpha_{\text{eff}}T_f$, giving the stated positive interior solution. Its derivative is $\partial N^*/\partial \alpha_{\text{eff}} = -T_f/(b + \delta_{\text{eff}}/s) \leq 0$, strictly negative when $T_f > 0$. *Model system optimum:* the model social cost is schedule delay $\frac{1}{2}\delta_{\text{eff}}N^2/s$ plus free-flow cost $\alpha_{\text{eff}}T_fN$ and, untolled, waiting deadweight $\frac{1}{2}\delta_{\text{eff}}N^2/s$. The planner maximises $\bar{P}N - \frac{1}{2}bN^2 - \frac{1}{2}\delta_{\text{eff}}N^2/s - \alpha_{\text{eff}}T_fN$, with marginal model social cost $\text{MSC}(N) = \delta_{\text{eff}}N/s + \alpha_{\text{eff}}T_f$. *Fine toll:* a time-varying toll equal to the would-be waiting-cost component removes the queue and makes the private price equal MSC; since $\text{MSC}(N)$ coincides with no-toll $C(N)$, equilibrium volume is **unchanged**, $N^{\text{toll}} = N^{\text{nt}} \equiv N^*$ (the classic bottleneck result). Model surplus at N^* is gross surplus minus free-flow cost and schedule delay, and minus waiting deadweight only when untolled;

toll revenue is a transfer. The optimal peak toll is $C_B(N^*)$. Thus the model-surplus gain from tolling is the eliminated waiting deadweight $\frac{1}{2}\delta_{\text{eff}}N^{*2}/s$. When $T_f > 0$, a lower α_{eff} raises N^* , gross surplus, the toll, waiting, and queue stock; when $T_f = 0$, N^* , gross surplus, and the toll are unchanged, while waiting and queue stock still rise through [Proposition 3](#). ■

Proof of Proposition 6. Consider the candidate ordering with the high- α class on both shoulders and the low- α class in the centre. The extreme arrivals face no wait, so equality of high-class boundary costs gives $a_{\min} = -\frac{\gamma}{\beta+\gamma}\frac{N}{s}$ and $a_{\max} = \frac{\beta}{\beta+\gamma}\frac{N}{s}$. Within a class, equal cost implies $T'(a) = \beta/\alpha$ on the early side and $T'(a) = -\gamma/\alpha$ on the late side. Because $\beta < \alpha_{\text{lo}} < \alpha_{\text{hi}}$, both early slopes preserve strict FIFO and the low- α segments are steeper. Writing $A_1 = \alpha_{\text{lo}}\beta/\alpha_{\text{hi}}$ and $A_2 = \alpha_{\text{lo}}\gamma/\alpha_{\text{hi}}$, continuity of $T(a)$ at the class boundaries and equality of low-class cost across its early and late entries give

$$A_1 |a_{\min}| + (\beta - A_1) x_e = A_2 a_{\max} + (\gamma - A_2) x_l, \quad x_e + x_l = N_{\text{lo}}/s,$$

Because $A_1 |a_{\min}| = A_2 a_{\max}$ and $A_1/\beta = A_2/\gamma = \alpha_{\text{lo}}/\alpha_{\text{hi}}$, the first equation reduces to $\beta x_e = \gamma x_l$. Together with the mass equation it gives the explicit x_e, x_l in the proposition. Moreover, $x_e/|a_{\min}| = x_l/a_{\max} = N_{\text{lo}}/N \in (0, 1)$, so both central pieces and both shoulders are interior for every pair of positive fixed class masses. To verify incentives, on the low-class central interval the high-class cost rises moving from the early boundary toward 0 and falls moving from 0 toward the late boundary; its minima are therefore the two boundary points, where it equals its shoulder cost. On either high-class shoulder, low-class cost is minimised at the adjacent central boundary. Hence neither class gains by entering the other's interval. The class masses are correct because the central interval has service duration $x_e + x_l = N_{\text{lo}}/s$. Let $r = N_{\text{lo}}/N$ and $M = (|a_{\min}|^2 + a_{\max}^2)/(2(|a_{\min}| + a_{\max}))$. Direct integration over the capacity-rate service intervals gives mean absolute deviation rM for the low- α class and $(1+r)M$ for the high- α class, so the former is strictly smaller. In the illustration, $x_e = 1.07$, $x_l = 0.18$ h and mean $|a|$ is 0.47 h for the low- α class versus 1.42 h for the high- α class. ■

Proof of Proposition 7. By Lemma 1, $g(a_e) = g(a_l) = C^*$ and $a_l - a_e = N/s$ determine (a_e, a_l) , and hence C^* , uniquely and independently of α_{eff} . Cost equalisation gives $T(a) = (C^* - g(a))/\alpha_{\text{eff}} \geq 0$ with maximum $T(0) = C^*/\alpha_{\text{eff}}$. Strict FIFO requires the home-departure map $d(a) = a - T(a)$ to satisfy $d'(a) > 0$, equivalently $T'(a) < 1$. Since $T'(a) = -g'(a)/\alpha_{\text{eff}}$ and $-g'(a)$ is maximised at a_e , the binding condition is $\alpha_{\text{eff}} > \alpha_{\text{crit}} \equiv -g'(a_e)$. On that range costs are equalised, FIFO holds, and an arrival outside the window faces no wait and strictly higher schedule cost. The toll $\tau(a) = C^* - g(a)$,

untolled model social cost NC^* , and tolled model social cost $s \int_{a_e}^{a_l} g(a) da$ are therefore independent of α_{eff} . At equality, $d'(a_e) = 0$: strict ordering reaches its boundary, but ordering is not reversed. Below it, $d'(a_e) < 0$ and this single-pass profile is invalid. This argument does not establish what alternative equilibrium, if any, replaces it. ■

Proof of Proposition 8. By Lemma 1, $\beta_{\text{eff}} t_e = \gamma t_l = C^*$ with $t_e + t_l = N/s$, exactly as in Proposition 3, so $C^* = \delta_{\text{eff}} N/s$. The optimal toll replaces waiting by the money charge $\tau(a) = C^* - \text{sched}(a)$; private cost and both model social-cost objects are unchanged. The interior wait solves $K(T(a)) = C^* - \text{sched}(a)$, so the assumed range and monotonicity of K give a unique $T(a) = K^{-1}(C^* - \text{sched}(a))$. On the early side, implicit differentiation gives $T'(a) = \beta_{\text{eff}}/K'(T) < 1$, which preserves strict home-departure ordering; the late-side derivative is negative. If K is convex and $\alpha_0 = K'(0)$, then $K(T) \geq K(0) + K'(0)T = \alpha_0 T$. Applying the increasing inverse to any required money gap therefore gives $K^{-1}(x) \leq x/\alpha_0$, strictly where $K(T) > \alpha_0 T$. The result remains conditional on existence of the stated FIFO equilibrium. ■

Proof of Corollary 2. Under the Fosgerau–Karlström (2010) assumptions (optimal departure, piecewise-linear scheduling cost, and the fixed location-scale distribution $\tilde{T} = \mu_T + \sigma_T X$), minimised expected cost is $\mathcal{C}_R = \alpha \mu_T + \sigma_T \mathcal{R}_\Phi(\beta, \gamma)$. Replacing α, β by the effective parameters of Proposition 2 yields the stated expression. The result is a substitution into an established theorem and does not cover ex-post work/rest recourse. ■

Appendix C: Illustrative calibration

The numerical illustrations are stylised. All figures and quoted numbers are produced by replication scripts available on the author’s website alongside the working paper, using Table 4 in US dollars and hours. The values are selected for transparent, regime-consistent illustration and are not fitted to a corridor; none of the mode shares, waits, surplus levels, or tolls is an empirical estimate. Capacity is measured in travellers per hour, equivalent here to vehicles per hour only under the maintained one-traveller-per-vehicle assumption, so $N/s = 2.5$ h. The queue-forming condition $\alpha_{\text{eff}} > \beta_{\text{eff}}$ holds for both homogeneous counterfactuals ($48 > 26 > 8$ for driving and $48 > 10 > 8$ for automated); this ordering is an illustrative maintained assumption, not an estimated empirical restriction. The simultaneous two-class case is used only for Proposition 6’s conditional sorting illustration. The mode-choice figure uses one-way travel time 0.6 h, prices $p_{\text{drive}} = \$4$ and $p_{\text{auto}} = \$9$, no alternative-specific constant, and a logit scale of \$6. The smooth-scheduling illustration

uses $g(a) = \frac{1}{2}k_e a^2$ early and $\frac{1}{2}k_l a^2$ late, with $k_e = 4$ and $k_l = 90$, placing $\alpha_{\text{crit}} \approx 8$; that proximity to $\beta_{\text{eff}} = 8$ is a numerical choice, not a structural identity. The duration-dependent waiting-cost illustration for [Proposition 8](#) uses $K(T) = \alpha_0 T + \frac{1}{2}\eta T^2$ with $\alpha_0 = 20$ and $\eta = 8$.

Table 4: Stylised parameters for the numerical illustrations.

Parameter	Value	Units
ω resource value of time	20	\$/h
w wage (flexible-labour calibration, $w = \omega$)	20	\$/h
β early-schedule penalty	8	\$/h
γ late-schedule penalty	48	\$/h
$\delta_{\text{eff}} = \beta\gamma/(\beta + \gamma)$	6.86	\$/h
N commuters	9,000	persons
s bottleneck capacity	3,600	travellers/h
N/s peak duration	2.5	h
e driving-effort cost	6	\$/h
$\alpha_{\text{drive}} = \omega + e$	26	\$/h
ϕ in-trip work productivity	0.5	fraction
ψ comfort/rest substitution	0.3	fraction
$\rho = \max(w\phi, \psi\omega)$	10	\$/h
$\alpha_{\text{auto}} = \omega - \rho$	10	\$/h
T_f free-flow time (Proposition 5)	0.5	h
$\bar{P}$ inverse-demand intercept (Proposition 5)	60	\$
b inverse-demand slope (Proposition 5)	0.004	\$/person
α_0 queue-time value intercept (Proposition 8 illustration)	20	\$/h
η accumulated-wait curvature (Proposition 8 illustration)	8	\$/h ²

References

- Afriat, S. N. (1967). The construction of utility functions from expenditure data. *International Economic Review*, 8(1), 67–77. <https://doi.org/10.2307/2525382>
- Arnott, R., de Palma, A., & Lindsey, R. (1990). The economics of a bottleneck. *Journal of Urban Economics*, 27(1), 111–130. [https://doi.org/10.1016/0094-1190\(90\)90028-L](https://doi.org/10.1016/0094-1190(90)90028-L)
- Arnott, R., de Palma, A., & Lindsey, R. (1993). A structural model of peak-period congestion: A traffic bottleneck with elastic demand. *American Economic Review*, 83(1), 161–179.
- Arnott, R., de Palma, A., & Lindsey, R. (1988). Schedule delay and departure time decisions with heterogeneous commuters. *Transportation Research Record*, 1197, 56–67.
- Becker, G. S. (1965). A theory of the allocation of time. *The Economic Journal*, 75(299), 493–517. <https://doi.org/10.2307/2228949>
- Bhat, C. R. (2005). A multiple discrete-continuous extreme value model: Formulation and application to discretionary time-use decisions. *Transportation Research Part B: Methodological*, 39(8), 679–707. <https://doi.org/10.1016/j.trb.2004.08.003>
- Bhat, C. R. (2008). The multiple discrete-continuous extreme value model: Role of utility function parameters, identification considerations, and model extensions. *Transportation Research Part B: Methodological*, 42(3), 274–303. <https://doi.org/10.1016/j.trb.2007.06.002>
- Bhat, C. R., & Dubey, S. K. (2014). A new estimation approach to integrate latent psychological constructs in choice modeling. *Transportation Research Part B: Methodological*, 67, 68–85. <https://doi.org/10.1016/j.trb.2014.04.011>
- Bowman, J. L., & Ben-Akiva, M. E. (2001). Activity-based disaggregate travel demand model system with activity schedules. *Transportation Research Part A: Policy and Practice*, 35(1), 1–28. [https://doi.org/10.1016/S0965-8564\(99\)00043-9](https://doi.org/10.1016/S0965-8564(99)00043-9)
- Daganzo, C. F. (1985). The uniqueness of a time-dependent equilibrium distribution of arrivals at a single bottleneck. *Transportation Science*, 19(1), 29–37. <https://doi.org/10.1287/trsc.19.1.29>
- de Sá, A. L. S., Lavieri, P. S., Pawlak, J., Sivakumar, A., & Thompson, R. G. (2025). The effects of travel time use on activity-travel behaviour: Knowledge consolidation and research agenda for current and future transport options. *Transport Reviews*, 45(6), 869–896. <https://doi.org/10.1080/01441647.2025.2517210>
- DeSerpa, A. C. (1971). A theory of the economics of time. *The Economic Journal*, 81(324), 828–846. <https://doi.org/10.2307/2230320>

- Fosgerau, M., & Karlström, A. (2010). The value of reliability. *Transportation Research Part B: Methodological*, 44(1), 38–49. <https://doi.org/10.1016/j.trb.2009.05.002>
- Fosgerau, M., Kim, J., & Ranjan, A. (2018). Vickrey meets alonso: Commute scheduling and congestion in a monocentric city. *Journal of Urban Economics*, 105, 40–53. <https://doi.org/10.1016/j.jue.2018.02.003>
- Fosgerau, M., & Small, K. A. (2017). Endogenous scheduling preferences and congestion. *International Economic Review*, 58(2), 585–615. <https://doi.org/10.1111/iere.12228>
- Huda, F. Y., Currie, G., & Kamruzzaman, M. (2023). Understanding the value of autonomous vehicles: An empirical meta-synthesis. *Transport Reviews*, 43(6), 1058–1082. <https://doi.org/10.1080/01441647.2023.2189324>
- Jain, J., & Lyons, G. (2008). The gift of travel time. *Journal of Transport Geography*, 16(2), 81–89. <https://doi.org/10.1016/j.jtrangeo.2007.05.001>
- Jara-Díaz, S. R. (2007). *Transport economic theory*. Emerald Group Publishing Limited.
- Jara-Díaz, S. R. (2024). The value(s) of travel time savings considering in-vehicle activities. *Transportation Research Part A: Policy and Practice*, 184, 104092. <https://doi.org/10.1016/j.tra.2024.104092>
- Jara-Díaz, S. R., & Guevara, C. A. (2003). Behind the subjective value of travel time savings: The perception of work, leisure, and travel from a joint mode choice activity model. *Journal of Transport Economics and Policy*, 37(1), 29–46. <https://doi.org/10.3828/jtep.2003.37.1.29>
- Kockelman, K. M. (2001). A model for time- and budget-constrained activity demand analysis. *Transportation Research Part B: Methodological*, 35(3), 255–269. [https://doi.org/10.1016/S0191-2615\(99\)00050-8](https://doi.org/10.1016/S0191-2615(99)00050-8)
- Krueger, R., Rashidi, T. H., & Dixit, V. V. (2019). Autonomous driving and residential location preferences: Evidence from a stated choice survey. *Transportation Research Part C: Emerging Technologies*, 108, 255–268. <https://doi.org/10.1016/j.trc.2019.09.018>
- Lancaster, K. J. (1966). A new approach to consumer theory. *Journal of Political Economy*, 74(2), 132–157. <https://doi.org/10.1086/259131>
- Lindsey, R. (2004). Existence, uniqueness, and trip cost function properties of user equilibrium in the bottleneck model with multiple user classes. *Transportation Science*, 38(3), 293–314. <https://doi.org/10.1287/trsc.1030.0045>
- Liu, Q., Liu, W., Jiang, R., & Han, X. (2025). On the morning commute problem with mixed autonomous and human-driven traffic under stochastic bottleneck capacity. *Transportation Research Part B: Methodological*, 195, 103203. <https://doi.org/10.1016/j.trb.2025.103203>

- Malokin, A., Circella, G., & Mokhtarian, P. L. (2019). How do activities conducted while commuting influence mode choice? *Transportation Research Part A: Policy and Practice*, 121, 82–102. <https://doi.org/10.1016/j.tra.2018.12.015>
- Mokhtarian, P. L. (2019). Subjective well-being and travel: Retrospect and prospect. *Transportation*, 46(2), 493–513. <https://doi.org/10.1007/s11116-018-9935-y>
- Mokhtarian, P. L., & Salomon, I. (2001). How derived is the demand for travel? some conceptual and measurement considerations. *Transportation Research Part A: Policy and Practice*, 35(8), 695–719. [https://doi.org/10.1016/S0965-8564\(00\)00013-6](https://doi.org/10.1016/S0965-8564(00)00013-6)
- Molin, E. J. E., Adjenughwure, K., de Bruyn, M., Cats, O., & Warffemius, P. (2020). Does conducting activities while traveling reduce the value of time? evidence from a within-subjects choice experiment. *Transportation Research Part A: Policy and Practice*, 132, 18–29. <https://doi.org/10.1016/j.tra.2019.10.017>
- Muellbauer, J. (1974). Household production theory, quality, and the hedonic technique. *American Economic Review*, 64(6), 977–994.
- Mun, D. J. (2007). A production-based approach to travel choice modeling. *Journal of Korean Society of Transportation*, 25(5), 209–227.
- Pinjari, A. R., & Bhat, C. R. (2010). A multiple discrete-continuous nested extreme value model: Formulation and application to non-worker activity time-use and timing behavior on weekdays. *Transportation Research Part B: Methodological*, 44(4), 562–583. <https://doi.org/10.1016/j.trb.2009.08.001>
- Pollak, R. A., & Wachter, M. L. (1975). The relevance of the household production function and its implications for the allocation of time. *Journal of Political Economy*, 83(2), 255–278. <https://doi.org/10.1086/260322>
- Pudāne, B. (2020). Departure time choice and bottleneck congestion with automated vehicles: Role of on-board activities. *European Journal of Transport and Infrastructure Research*, 20(4), 306–334. <https://doi.org/10.18757/ejtir.2020.20.4.4801>
- Pudāne, B., & Correia, G. H. d. A. (2020). On the impact of vehicle automation on the value of travel time while performing work and leisure activities in a car: Theoretical insights and results from a stated preference survey—a comment. *Transportation Research Part A: Policy and Practice*, 132, 324–328. <https://doi.org/10.1016/j.tra.2019.11.019>
- Pudāne, B., Molin, E. J. E., Arentze, T. A., Maknoon, Y., & Chorus, C. G. (2018). A time-use model for the automated vehicle-era. *Transportation Research Part C: Emerging Technologies*, 93, 102–114. <https://doi.org/10.1016/j.trc.2018.05.022>

- Singleton, P. A. (2019). Discussing the “positive utilities” of autonomous vehicles: Will travellers really use their time productively? *Transport Reviews*, 39(1), 50–65. <https://doi.org/10.1080/01441647.2018.1470584>
- Small, K. A. (1982). The scheduling of consumer activities: Work trips. *American Economic Review*, 72(3), 467–479.
- Small, K. A. (2012). Valuation of travel time. *Economics of Transportation*, 1(1–2), 2–14. <https://doi.org/10.1016/j.ecotra.2012.09.002>
- Smith, M. J. (1984). The existence of a time-dependent equilibrium distribution of arrivals at a single bottleneck. *Transportation Science*, 18(4), 385–394. <https://doi.org/10.1287/trsc.18.4.385>
- van den Berg, V., & Verhoef, E. T. (2011). Winning or losing from dynamic bottleneck congestion pricing? *Journal of Public Economics*, 95(7–8), 983–992. <https://doi.org/10.1016/j.jpubeco.2010.12.003>
- van den Berg, V. A. C., & Verhoef, E. T. (2016). Autonomous cars and dynamic bottleneck congestion: The effects on capacity, value of time and preference heterogeneity. *Transportation Research Part B: Methodological*, 94, 43–60. <https://doi.org/10.1016/j.trb.2016.08.018>
- Varian, H. R. (1982). The nonparametric approach to demand analysis. *Econometrica*, 50(4), 945–973. <https://doi.org/10.2307/1912771>
- Vickrey, W. S. (1969). Congestion theory and transport investment. *American Economic Review*, 59(2), 251–260.
- Vij, A., & Walker, J. L. (2016). How, when and why integrated choice and latent variable models are latently useful. *Transportation Research Part B: Methodological*, 90, 192–217. <https://doi.org/10.1016/j.trb.2016.04.021>
- Wang, J., Jian, S., Robson, E. N., Dixit, V. V., & Rashidi, T. H. (2026). Assessing the economic impacts of labour time in autonomous vehicles. *Multimodal Transportation*, 5(1), 100235. <https://doi.org/10.1016/j.multra.2025.100235>
- Wu, S., & Li, Z.-C. (2023). Managing a bi-modal bottleneck system with manned and autonomous vehicles: Incorporating the effects of in-vehicle activity utilities. *Transportation Research Part C: Emerging Technologies*, 152, 104179. <https://doi.org/10.1016/j.trc.2023.104179>
- Yu, X., van den Berg, V. A. C., & Verhoef, E. T. (2022). Autonomous cars and activity-based bottleneck model: How do in-vehicle activities determine aggregate travel patterns? *Transportation Research Part C: Emerging Technologies*, 139, 103641. <https://doi.org/10.1016/j.trc.2022.103641>