

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

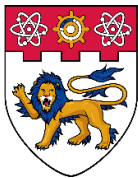
Zach J.L. Lee

June 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

2026-06-26

[Latest version](#)

Abstract

Automation, ride-hailing, and other passive services make travel time partly usable for work or rest. Transport choice then depends on what travellers can produce while moving, not only on money and clock time. I model travel as production. The final commodity is *scheduled activity access*, trips are intermediate inputs, and modes, routes, departure times, and vehicle technologies are production plans. Three results follow. First, the value of travel time is endogenous to within-trip work and rest. Active time costs the resource value of time plus effort; passive time costs that value net of reclaimed output. Under additive in-trip outputs, the model also yields a constant-returns implicit-price rule as a special case. Second, putting the production technology into a Vickrey/Arnott–de Palma–Lindsey bottleneck separates welfare and pricing effects. Productive in-vehicle time lengthens the physical queue but leaves money-metric generalised cost and the fine toll unchanged in the pure point bottleneck. This neutrality applies only to money-metric traveller cost, before spillback, vehicle-hour, energy, emissions, crash-exposure, and curb-space externalities are priced. Productive early arrival flattens the peak; lower free-flow travel cost under elastic demand induces trips and raises the optimal toll. Third, the production structure imposes falsifiable restrictions relative to generic random-utility and discrete-continuous models. 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 passive services make travel time partly 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. I add the traveller's within-trip production choice: *how much output to produce during the trip*. In the time-allocation literature, that choice is usually treated as exogenous. If passive travel time can be turned into paid-equivalent work or into rest, the opportunity cost of that time is reduced endogenously; if the traveller must actively drive, the cost is the resource value of time *plus* effort. The value assigned

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

to travel time is therefore determined by the within-trip production choice, rather than introduced as a free preference parameter. The positive-utility-of-travel literature documents this use of travel time empirically (Malokin et al., 2019; Mokhtarian, 2019; Mokhtarian & Salomon, 2001). Here it enters a production-and-equilibrium model.

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 the pure bottleneck, making in-vehicle time productive does **not** change the equilibrium cost or the optimal toll; it lengthens only the physical queue. This is the standard result that equilibrium bottleneck cost is independent of the marginal value of queueing time, interpreted through the production model. Welfare and pricing respond to productive *early* arrival, which flattens the peak, 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 1 collects the four channels through which productive travel time affects welfare and pricing.

Testability. A production model has empirical content only if it rules out some choice patterns. The paper therefore states testable restrictions imposed by the production structure, restrictions that a generic random-utility or multiple-discrete-continuous model need not impose. One restriction is that telework feasibility lowers the value of passive travel time but not active driving time. The other is a non-equivalence between money and time shocks once time use is observed.

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 endogenous-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-flattening, induced-demand, and sorting results, and the robustness of the neutrality result to smooth scheduling and a queue-dependent value of time. Section 6 extends the value of reliability. Section 7 states the empirical restrictions. Section 8 draws the welfare 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 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 response in this paper is to use cost *duality* in which implicit prices are not assumed constant; the constant-price case is derived as a special case. 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. The model makes the second term endogenous to an in-trip production choice.

Value of travel time. Small’s (2012) survey is the reference point. The empirical value of time is heterogeneous and context-dependent, often a fraction of the wage. The model in this paper accounts for this by having the relevant object be a resource value of time net of in-trip output and effort, both of which vary by technology and traveller.

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). Relative to this literature, the contribution here is not the bottleneck mechanics. It is the within-trip work/rest production foundation that makes α_{eff} and β_{eff} endogenous, the productive-early-arrival channel, and the falsifiable restrictions.

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 production-duality interpretation clarifies the meaning of generalised cost and supplies testable restrictions.

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 , non-travel leisure or home time L , 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 $\tau^r = T_j^P - \tau^b$ to rest. 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:

- $\phi_j \in [0, 1]$: in-trip work productivity (fraction of the wage earnable during passive travel); $\phi_j = 0$ 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 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 and strictly concave (the value of discretionary time) and u_W capturing any direct (dis)utility of work. 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 says in-trip work earns at rate ϕ_j of the wage. The term $w\phi_j \tau^b$ is the *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. 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.

The benchmark treats ϕ_j and ψ_j as constant marginal rates over the relevant trip. For empirical work, ϕ_j should be interpreted as a net effective productivity rate after task indivisibilities, setup time, and switching costs. 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.

The traveller chooses (W, τ^b) to maximise Equation 3 subject to Equation 1–Equation 2, taking technology j and its travel times as given, and then chooses j . The **conditional cost** of access through j is the money-metric value of the resources it consumes; the chosen technology minimises it. In this sense, each travel option is a technology for producing scheduled activity access. The traveller chooses the least-cost technology for the

access she needs.

Since the activity is fixed in the conditional problem, B is common across technologies and drops out of the technology choice; participation requires $B \geq \min_j C_j$, and 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).

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. Since preferences are quasilinear in the numeraire, income effects on the value of time are outside the benchmark.

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 3](#) 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, earning $w\phi_j$, or resting, worth $\psi_j \omega$ in

leisure-equivalent terms. She picks the larger. Define the **reclaimed value of passive time**

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

Proposition 1 (value of travel time with endogenous 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. These values are derived from the optimised within-trip production choice rather than imposed as primitive preference parameters.

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 value through work, and a comfortable but not work-friendly one (high ψ_j , low ϕ_j) through rest. Either way the value of passive travel time falls below ω . The positive-utility literature documents this use of travel time; here it is derived within the model. 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 cost less than half as much when she can use them.

Figure 1 displays the components of VTTS for a calibrated 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 a binary logit between driving and a robotaxi as the robotaxi’s work productivity ϕ rises, holding $\psi = 0.3$: its share climbs from 59% to 83%, and the fleet-average value of time falls. The same minutes of travel have a different production cost as ϕ changes. The logit calibration is reported in [Appendix C](#).

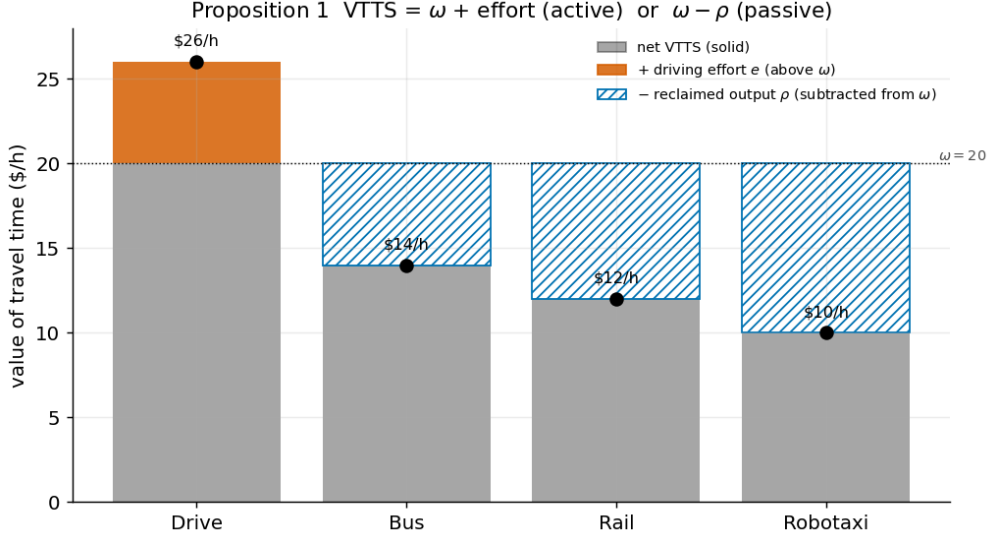


Figure 1: Value-of-travel-time components. The dotted line is the resource value ω .

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 3.

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, $y = a_j q_j$; 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

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

Proposition 1 Endogenous in-trip output shifts mode choice

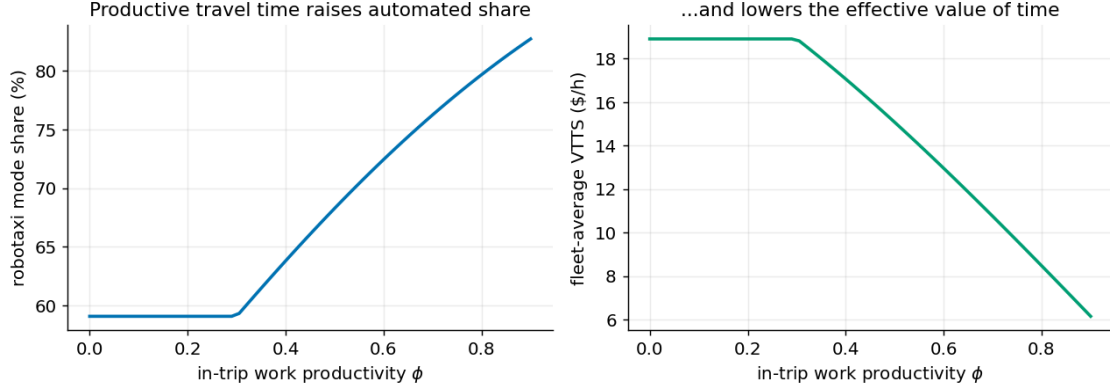


Figure 2: As in-trip work productivity ϕ rises for the automated mode, holding $\psi = 0.3$, the automated share rises and the effective value of time falls.

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 in the partial-equilibrium mode-choice figure (Figure 2), but a joint adoption-congestion equilibrium, in which travellers choose the technology given a fare or price wedge while queues, costs, and adoption shares are determined together, is left outside the benchmark.

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 in the benchmark 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.

Pure point bottleneck refers to the benchmark with 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. Section 8 separates the welfare and pricing mechanisms that relax or qualify these assumptions (Table 1). The numerical illustrations use the stylised calibration in Table 3 (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. If a unit of early-arrival time yields reclaimable value $\rho_c^E = \max\{w\phi_c^E, \psi_c^E\omega\}$, defined analogously to ρ_c but for early-arrival time at the destination, then the effective earliness penalty is $\beta_{\text{eff}} = \beta - \rho_c^E$. This requires $\rho_c^E < \beta$, so that early arrival remains costly. Late arrival remains feasible and carries the finite, non-reclaimable penalty γ . The production model therefore reduces to an α - β - γ bottleneck with technology-specific effective parameters. With several technologies or traveller groups, the model is a multi-class bottleneck.

Proposition 2 (reduction). 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), $\beta_{\text{eff}}^c = \beta - \rho_c^E$, and unchanged γ . Effort and in-trip output are not new attributes; they change the effective scheduling parameters.

For a homogeneous population, the no-toll equilibrium is the textbook one. With $\delta_{\text{eff}} \equiv \beta_{\text{eff}}\gamma/(\beta_{\text{eff}} + \gamma)$, the equilibrium generalised cost, the earliest and latest arrivals, and the queue 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 queueing time is a real resource cost, with the equilibrium under the fine time-varying toll, where the toll removes the queue and toll revenue is treated as a transfer.

Proposition 3 (money-metric production-cost neutrality in the pure point bottleneck). Take the pure point bottleneck of [Assumption B1](#) with a single class and $\alpha_{\text{eff}} > \beta_{\text{eff}}$. Let the welfare object be money-metric traveller production cost, with queue time valued at α_{eff} and excluding spillback, vehicle-hours, energy, emissions, crash exposure, and curb/space. Then the following objects are independent of α_{eff} : the equilibrium generalised cost $C^* = \delta_{\text{eff}}N/s$; the money-metric bottleneck cost, equal to $\delta_{\text{eff}}N^2/s$ untolled and $\frac{1}{2}\delta_{\text{eff}}N^2/s$ with the fine toll; and the fine money toll. Lowering α_{eff} by making in-vehicle time more productive leaves these objects unchanged and only lengthens the physical queue, whose peak rises to $T_{\text{max}} = C^*/\alpha_{\text{eff}}$.

Boundary. The lower boundary is the relevant one. As α_{eff} falls, the peak queue 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 or below β_{eff} , the orderly FIFO queue is no longer the maintained object, and the paper makes no claim for that regime. This boundary is taken up in [Section 5.5](#), where a smooth scheduling cost identifies the regime beyond it as one of bunching, not orderly queueing.

[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 cost and queue time. A lower value of queue time means travellers tolerate a longer queue to reach the same generalised cost, so the queue grows just enough to hold cost constant. The money toll that replaces the queue is $\tau(a) = \alpha_{\text{eff}}T(a)$, and since $T(a) \propto 1/\alpha_{\text{eff}}$, the money toll is invariant. In the running example the commuter’s automated trip waits 1.71 h at the peak against 0.66 h driving, yet both face the same \$17.1 per-traveller cost, the same \$154,300 untolled total, and the same \$17.1 peak toll ([Figure 3](#)); only the automated queue is longer.

The α_{eff} -independence of C^* is a standard property of the deterministic bottleneck. The production micro-foundation adds two things. First, α_{eff} is itself *endogenous* to in-trip output and driving effort ([Section 4](#)), so in-vehicle automation changes queue length without changing the equilibrium money cost in the pure bottleneck. Second, the model

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

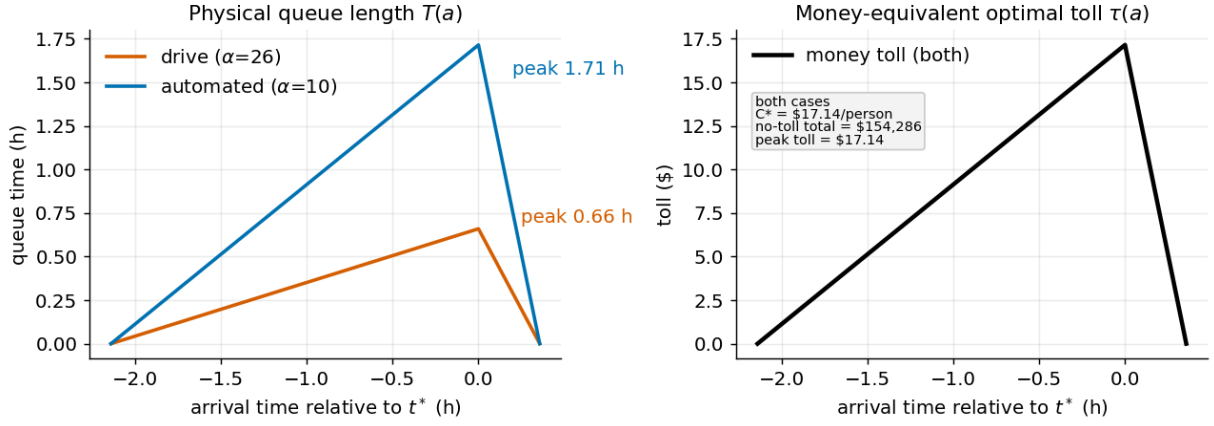


Figure 3: **Proposition 3.** Left: a lower value of in-vehicle time ($\alpha_{\text{eff}} = 10$, automated) produces a longer physical queue than driving ($\alpha_{\text{eff}} = 26$); since $\beta_{\text{eff}} = 8 < \alpha_{\text{eff}}$, the queue regime holds for both classes. Right: the fine money toll is identical. Money-metric production cost and the fine toll are unchanged; only physical congestion grows.

separates three effects that can otherwise look similar: the value of queueing time α_{eff} , the effective early-arrival penalty β_{eff} , and the free-flow cost term under elastic demand. These effects have different welfare and pricing implications, as summarised in Table 1. The new analytical results are the endogenous-output value-of-time decomposition (Proposition 1), the classification of welfare and pricing effects, and the falsifiable restrictions developed in Section 7 (R1–R4); a new bottleneck equilibrium is not claimed.

The result limits what in-vehicle productivity alone can imply for automation. If the relevant benefit is *in-vehicle productivity*, rather than cheaper schedule adjustment or added capacity, automation can lengthen queues and add vehicle-hours without changing the model’s money-metric traveller cost or corrective money toll. Once the excluded physical and network externalities are priced, the welfare effect need not be neutral and may turn negative. Physical congestion and money-metric traveller cost can therefore move in opposite directions. For example, spillback into upstream links would create delays outside the pure-point welfare object, so the neutrality result should not be read as a network-welfare claim.

5.3 Where welfare and pricing actually respond

The relevant welfare and pricing effects instead run through two other channels.

Proposition 4 Productive early arrival flattens and lengthens the peak

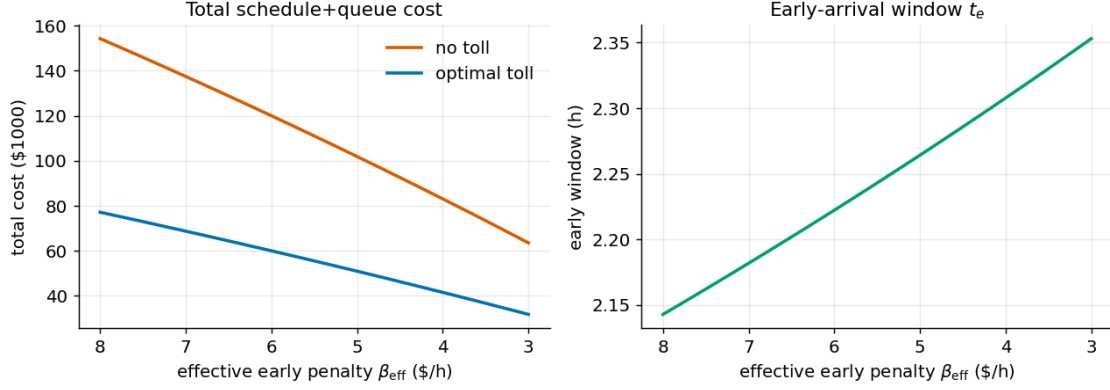


Figure 4: **Proposition 4.** Lower effective early-schedule penalty (productive early arrival, to the right) reduces total cost and lengthens the early-arrival window. The peak flattens.

Proposition 4 (productive early arrival flattens the peak). A reduction in the effective earliness penalty β_{eff} , because time before t^* can be used at the destination or because work is flexible, lowers $\delta_{\text{eff}} = \beta_{\text{eff}}\gamma/(\beta_{\text{eff}} + \gamma)$, and hence lowers both the no-toll total cost $\delta_{\text{eff}}N^2/s$ and the tolled total cost $\frac{1}{2}\delta_{\text{eff}}N^2/s$. The early shoulder t_e lengthens, so the peak spreads and flattens.

Unlike in-vehicle productivity, productive early arrival reduces real schedule-delay cost. Figure 4 shows total 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 work when she arrives early, her minutes before t^* are productive, her effective β_{eff} falls, and the peak she contributes to spreads. This gives a production-theoretic rationale for flexible-arrival policies and destination amenities, such as workspaces or early-access facilities, that make arriving early useful.

Proposition 5 (induced demand and the optimal toll under automation). Add a free-flow travel time T_f , so the generalised cost is $C(N) = \delta_{\text{eff}}N/s + \alpha_{\text{eff}}T_f$, and let demand be elastic, $P(N) = a - bN$. Then the no-toll equilibrium trip volume is $N^* = (a - \alpha_{\text{eff}}T_f)/(b + \delta_{\text{eff}}/s)$, increasing as α_{eff} falls. Cheaper in-vehicle time (automation) therefore *induces trips*, and the peak of the optimal time-varying toll, $\delta_{\text{eff}}N^*/s$, *rises* (the schedule is $\tau(a) = C^* - \text{sched}(a)$, with maximum $\delta_{\text{eff}}N^*/s$). 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.

Proposition 5 Induced demand, endpoint welfare, and the optimal toll

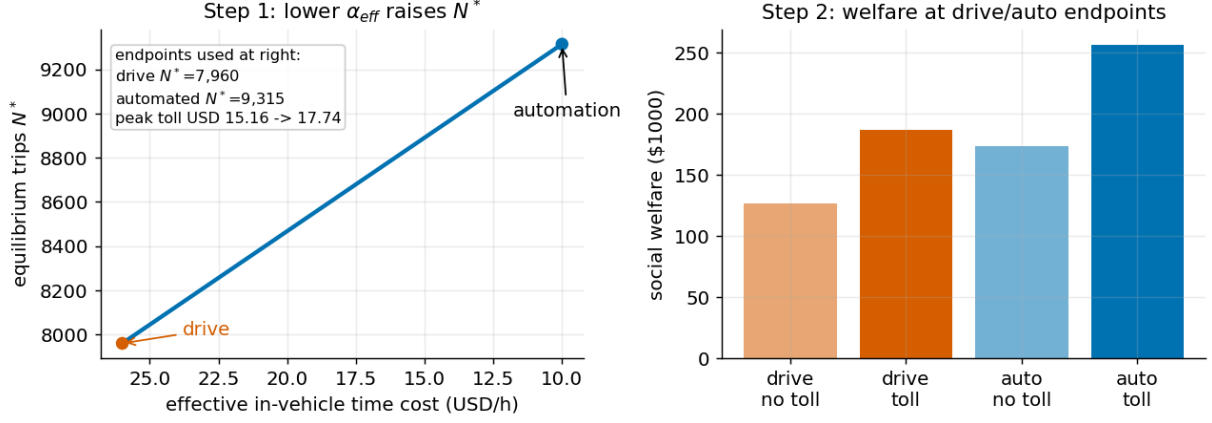


Figure 5: **Proposition 5.** Left: lower effective in-vehicle time cost induces equilibrium trips; the drive and automation endpoints are the cases evaluated on the right. The optimal peak toll rises with the induced volume. Right: welfare at those endpoint equilibria, with and without the optimal toll.

Figure 5 (left) shows trips rising from about 7,960 to 9,315 as the technology shifts from driving to automated; the optimal peak toll rises from \$15.2 to \$17.7 because it is proportional to the induced volume. The right panel evaluates welfare at those two endpoint equilibria. The time-varying toll removes the queue deadweight (the Vickrey result that the fine toll leaves trip volume unchanged but converts queue into revenue), while automation raises gross trip surplus and, with it, both welfare and total congestion. Together with Proposition 3, the figure shows that the welfare effect of automation depends on which component of travel cost changes.

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 result). With common β, γ and two classes $\alpha_{lo} < \alpha_{hi}$ in 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 $A_1 = \alpha_{lo}\beta/\alpha_{hi}$ and $A_2 = \alpha_{lo}\gamma/\alpha_{hi}$. Let (x_e, x_l) solve

$$A_1|a_{\min}| + (\beta - A_1)x_e = A_2a_{\max} + (\gamma - A_2)x_l, \quad x_e + x_l = N_{lo}/s.$$

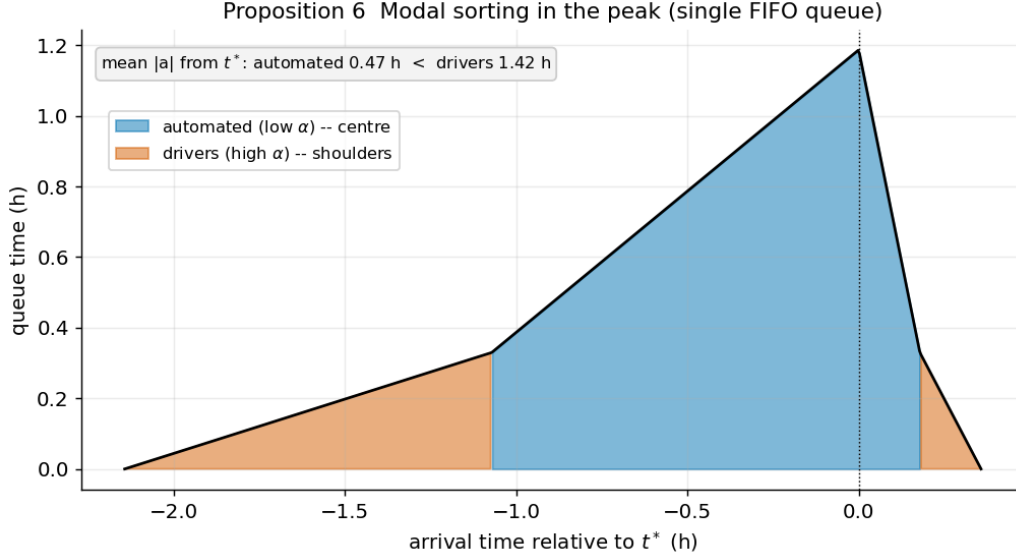


Figure 6: **Proposition 6.** Automated travellers (low value of queue time) sort to the centre of the peak; drivers take the shoulders. Mean $|$ arrival from t^* $|$: 0.47 h vs 1.42 h.

If that solution is interior, $0 < x_e < |a_{\min}|$ and $0 < x_l < a_{\max}$, then the low- α class (automated/passive, which has the lowest value of queueing time) occupies the central interval $[-x_e, x_l]$ of the peak, where queues are longest and schedule delay is smallest. The high- α class (drivers) occupies the shoulders; the mean absolute arrival deviation from t^* is then strictly smaller for the low- α class. The interior condition holds for moderate low- α shares; as the low- α share $\rightarrow 0$, the central band vanishes, and as it $\rightarrow 1$, the shoulders do.

The queue 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 identifies the class difference. Automated/passive travellers have a lower effective value of queueing time, so they occupy the centre of the peak, while active drivers take the shoulders. When the commuter switches to the automated mode, she joins the centre, arriving on average 0.47 h from t^* against 1.42 h as a driver.

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 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 $\beta_{10} = 2$ (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 ordered-queue benchmark. 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 value of queue time is replaced with a queue-dependent value.

Replace the kinked penalty with a strictly convex, C^1 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 physical queue respects FIFO.

Proposition 7 (neutrality under smooth scheduling; a local-slope boundary).

In the pure point bottleneck with strictly convex C^1 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 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 no-toll money-metric social cost NC^* , the tolled money-metric social cost $s \int_{a_e}^{a_l} g(a) da$, and the fine toll are **independent of** α_{eff} , while the peak queue 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 bunching boundary therefore rises with corridor congestion. For $\alpha_{\text{eff}} \leq \alpha_{\text{crit}}$ first-in-first-out

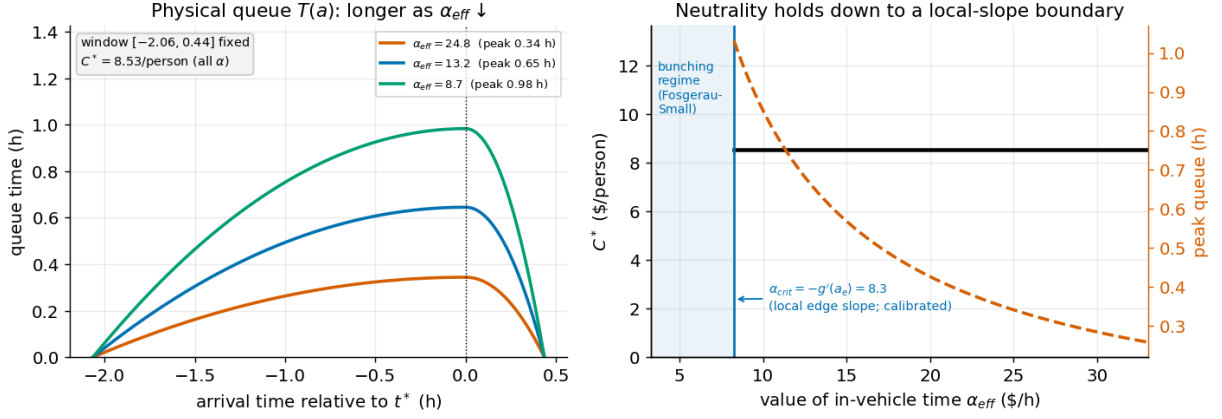


Figure 7: **Proposition 7.** Left: with a smooth, strictly convex schedule, a lower value of in-vehicle time lengthens the physical queue, but the arrival window and C^* are fixed. Right: C^* (money cost, solid) is flat across the queueing regime while the peak queue (dashed) rises as α_{eff} falls, up to the finite C^*/α_{crit} at $\alpha_{\text{crit}} = -g'(a_e)$, the local early-schedule slope at the window edge. Its value ≈ 8 here is a calibration choice, not a match to the primitive β_{eff} . Below it the system enters a bunching regime.

ordering fails at the early edge ($T'(a_e) = \alpha_{\text{crit}}/\alpha_{\text{eff}} \geq 1$) and no orderly queueing equilibrium of this form exists.

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 boundary. The location of the boundary is demand-dependent. A larger N/s raises α_{crit} and makes bunching relevant at higher values of queue time. Below the boundary, the maintained single-pass FIFO profile fails. Characterising that regime requires a bunching model, so no welfare claim is made there. The production model identifies the transition point. In-vehicle automation pushes α_{eff} toward α_{crit} , and the orderly-queue neutrality holds until it is crossed.

The boundary-cost argument also covers the case in which the value of queue time rises with the queue. This captures the idea that stop-and-go waiting may be more effortful than free-flow driving.

Proposition 8 (conditional neutrality under a queue-dependent value of time). Let $K(T)$ be the cumulative money cost of a wait of length T , with $K(0) = 0$, $K' > 0$, and $K'(T) \geq \beta_{\text{eff}}$ on the equilibrium range, and suppose a queue-forming FIFO equilibrium satisfying the boundary conditions of Lemma 1 exists. In the kinked bottleneck the equilibrium cost $C^* = \delta_{\text{eff}}N/s$, the no-toll and tolled money-metric

social costs, and the fine toll are unchanged from [Proposition 3](#); the physical queue is $T(a) = K^{-1}(C^* - \text{sched}(a))$, generally nonlinear in arrival time.

By [Lemma 1](#), the equilibrium cost is fixed by the boundary arrivals, who face *no* queue, so $\beta_{\text{eff}} t_e = \gamma t_l = C^*$ holds whatever the cost of waiting. A convex K changes how a given money cost is converted into waiting. In the calibrated illustration, with the marginal value of queue time rising linearly in the traveller’s own accumulated wait, $K(T) = \alpha_0 T + \frac{1}{2} \kappa T^2$ ($\alpha_0 = 20$, $\kappa = 8$), the peak *shortens* from 0.86 to 0.75 h relative to the constant-value baseline $K(T) = \alpha_0 T$, since travellers tolerate less of an increasingly costly queue. The money-metric cost and the corrective toll remain as in [Proposition 3](#). Taken together, [Proposition 7](#) and [Proposition 8](#) show that the neutrality result rests on the structure of the bottleneck. The convenient constants of the α - β - γ form are not essential.

6 The value of reliability with production

Travel time is uncertain. With optimal departure under a random free-flow time of mean μ_T and standard deviation σ_T , Fosgerau and Karlström (2010) show that expected cost is linear in the mean and the standard deviation. Substituting the effective parameters of [Section 5](#):

Corollary 2 (value of reliability with production). Under the Fosgerau–Karlström conditions with technology-specific effective parameters,

$$\mathbb{E}[C^*] = \alpha_{\text{eff}} \mu_T + \sigma_T \kappa(\beta_{\text{eff}}, \gamma),$$

so in-trip productivity lowers the value assigned to mean travel time through α_{eff} . Productive early arrival lowers the value of reliability $\kappa(\beta_{\text{eff}}, \gamma)$ through β_{eff} .

The corollary follows by direct substitution of the effective parameters into the Fosgerau–Karlström formula. One qualification is important. 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–R4 are over-identifying restrictions it imposes that a generic such model, carrying free per-mode time coefficients, need not satisfy.

The restrictions 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 is a revealed-preference consistency check using constructed full prices. A generic model with free per-mode time coefficients need not satisfy these restrictions.

R1 (active/passive asymmetry). Because $VTTS_j^P = \omega - \max\{w\phi_j, \psi_j\omega\}$, the value of passive travel time weakly falls with effective in-trip productivity. Because $VTTS_j^A = \omega + e_j$, active driving time is unaffected by that productivity. This sign-and-zero pattern is the restriction. The wage comparative static is as follows. Under rigid hours, where ω is fixed, $VTTS^P$ also falls with the wage; under flexible hours, where $\omega = w$, it need not. 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 the wage because telework feasibility may proxy wage differences.

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 (revealed-preference consistency). When full production-price vectors can be constructed, choices over produced bundles should satisfy the generalised axiom of revealed preference and the associated Afriat inequalities (Afriat, [1967](#); Varian, [1982](#)). This is a joint test because the full-price vector (money, time, effort, and schedule) is built from the model’s constructed shadow prices rather than observed directly. It tests the production-price construction together with rationality, not an assumption-free non-parametric restriction. With few modes or little independent price variation, such tests have limited power, so non-rejection is weak evidence. Slutsky-type restrictions in full prices should be treated as stronger, data-demanding implications.

8 Welfare and pricing

In the model, pricing effects depend on which part of the travel-production technology changes. [Table 1](#) collects four cases and their welfare and pricing implications. The classification is money-metric and uses the quasilinear benchmark. With income effects, the signs should be read as local effects evaluated at the maintained marginal utility of money, not as global additive-surplus comparisons. The rest of this section discusses the four cases.

Table 1: Production channels and their welfare and pricing implications.

Channel	Parameter	Pure-bottleneck effect	Elastic-demand effect	Toll effect	Welfare effect
In-vehicle productivity	$\alpha_{\text{eff}} \downarrow$	queue lengthens; money cost unchanged	trips rise via T_f	unchanged in pure bottleneck; higher with elastic demand	neutral in pure bottleneck; sign depends on spillback and physical externalities
Productive early arrival	$\beta_{\text{eff}} \downarrow$	peak flattens; cost falls	lower schedule cost	lower peak-toll component	improves
Free-flow / elastic demand	$\alpha_{\text{eff}} T_f \downarrow$	none (no T_f in pure bottleneck)	induced demand	higher optimal toll	trip surplus up, congestion up (net ambiguous)
Capacity	$s \uparrow$	generalised cost falls	trips rise	toll per user falls	improves

First, the corrective toll prices the congestion externality one traveller imposes on others, expressed in production terms. This includes minutes of delay, induced schedule delay, unreliability, and lost in-trip output. For in-vehicle productivity, the welfare entry in Table 1 is neutral only for the model’s money-metric traveller cost. Once the excluded physical externalities are priced, the sign may turn negative. In the pure bottleneck, the money toll is invariant to in-vehicle productivity (Proposition 3); under elastic demand with free-flow time it rises with automation (Proposition 5).

Second, productive in-vehicle time is not a substitute for capacity expansion. Capacity expansion lowers N/s and thus C^* ; in-vehicle automation lowers α_{eff} and, in the pure bottleneck, leaves this cost unchanged while lengthening queues. A planner who treats productive travel time as equivalent to capacity will over-value automation and set tolls too low.

Third, policies that reduce schedule-delay cost (flexible arrival windows, destination amenities that make early arrival useful, [Proposition 4](#)) complement pricing because they shrink the deadweight that pricing reallocates.

Fourth, the welfare evaluation of automation depends on the production effect. It raises trip surplus ([Proposition 5](#)), can worsen physical congestion ([Proposition 3](#)), and can induce demand. Net welfare depends on the parameters highlighted by the model: ρ_j , e_j , T_f , and demand elasticity.

9 Conclusion

Automation, ride-hailing, and other passive services make travel time partly 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 the components of the value of travel time. Active time is valued at the resource value of time plus effort. Passive time is valued at the resource value of time less the work or rest reclaimed during travel. Mun’s implicit-price rule is recovered under constant returns and additive byproducts. Embedding the technology in a Vickrey–Arnott–de Palma–Lindsey bottleneck then links these production effects to known bottleneck pricing and sorting results. In-vehicle productivity is neutral for money-metric cost and the fine toll in the pure point bottleneck, although it lengthens queues and may worsen welfare once the excluded physical and network externalities are priced. Productive early arrival lowers schedule-delay cost and flattens the peak. Lower free-flow travel cost under elastic demand induces trips and raises the optimal toll. Travellers with lower values of queueing 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 benchmark also excludes operating, energy, emissions, safety, and curb-space externalities, except where they are discussed as welfare qualifications.

The preference structure is also narrow. Welfare and testability are stated under quasi-linear money-metric preferences, so income effects are outside the benchmark. The neutrality result is stated only for the queue-forming regime; below the boundary, the model

identifies a transition to bunching but does not solve that regime. The in-trip productivity rate ϕ_j is a net marginal rate and abstracts from setup and teardown costs. 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 bunching regime should be characterised with smooth scheduling preferences or random queue position, and its welfare effects priced directly. Second, the pure-point bottleneck should be embedded in a corridor or network model with spillback and physical externalities. Third, the linear in-trip productivity rate should be replaced, where data permit, with a time-use technology that allows setup costs and task indivisibilities. Fourth, the revealed-preference restrictions should be estimated with time-use and mode data. The main implication is that productive travel time is not a generic congestion benefit. Its welfare and pricing effects depend on which production parameter changes.

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

Table 2: 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
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	in-trip work productivity during passive travel	$[0, 1]$
ψ_j	leisure-equivalent value of in-trip rest	$[0, 1]$
e_j	marginal effort cost of active travel time	money/time
ω	resource value of discretionary time	money/time
ρ_j	value reclaimed from a unit of passive travel time, $\max\{w\phi_j, \psi_j\omega\}$	money/time

Symbol	Meaning	Units or domain
ϕ_c^E, ψ_c^E	early-arrival work productivity and rest substitution for class c	$[0, 1]$
ρ_c^E	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 queueing time	clock time; time; time
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
$\beta_{\text{eff}}, \gamma$	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^*	equilibrium generalised cost per traveller	money
$\tau(a)$	time-varying money toll at arrival time a	money
T_f	free-flow travel time	time
$P(N)$	inverse demand for trips	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 waiting T in Proposition 8	money

Symbol	Meaning	Units or domain
μ_T, σ_T	mean and standard deviation of random travel time	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 $w\phi_j$ 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 . For passive time, define the within-trip value $m_j(T^P) = \max_{0 \leq \tau^b \leq T^P} \{w\phi_j \tau^b + u_L(L + \psi_j(T^P - \tau^b))\}$. By the constrained envelope theorem, $m'_j(T^P) = \max\{w\phi_j, \psi_j \omega\} = \rho_j$: in the work corner the bound $\tau^b \leq T^P$ binds and a marginal unit of passive time is taken for work, contributing $w\phi_j$; otherwise it is taken for rest, contributing $\psi_j \omega$. Hence $\partial U_j / \partial T_j^P = -\omega + \rho_j$, and the disutility (money-metric 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 $\text{VTTS}_j^A = \omega + e_j$. ■

Proof of Corollary 1. With $y = a_j q_j$ fixed, the number of trips is $q_j = y/a_j$. 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 $C_j(y) = (p_j + v_j t_j) q_j = \frac{p_j + v_j t_j}{a_j} y$, and the cost-minimising technology minimises the unit cost $\mu_j = (p_j + v_j t_j)/a_j$. ■

Proof of Proposition 2. Immediate from Proposition 1: the cost of a unit of queue time is the value of that time, $\omega + e_c$ if active and $\omega - \rho_c$ if passive; productive early arrival reduces the marginal earliness penalty by the reclaimed rate 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, the queue 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 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 queueing then fills the gap between the boundary cost and the local schedule cost: with constant queue-time value, $\alpha_{\text{eff}} T(a) = C^* - \text{sched}(a)$ or $\alpha_{\text{eff}} T(a) = C^* - g(a)$; with a queue-dependent value of time, the same money gap is delivered by the corresponding cumulative queue cost. In all cases the boundary conditions determine the money cost before queue-time valuation converts it into physical waiting. ■

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} . The queue time solves $\alpha_{\text{eff}} T(a) + \text{sched}(a) = C^*$, so $T(a) = (C^* - \text{sched}(a))/\alpha_{\text{eff}}$, with peak $T_{\text{max}} = C^*/\alpha_{\text{eff}}$ at $a = 0$: the queue length scales as $1/\alpha_{\text{eff}}$. The optimal time-varying toll replaces queue time, $\tau(a) = \alpha_{\text{eff}} T(a) = C^* - \text{sched}(a)$, which is independent of α_{eff} ; the peak toll is C^* . The *private* generalised cost is C^* per traveller in both regimes (the fine toll replaces queue time one-for-one, hence is a pure transfer). The *money-metric bottleneck* cost (the model's social cost, which excludes the physical externalities listed in the proposition) differs by regime but not by α_{eff} : $\delta_{\text{eff}} N^2/s$ without the toll (the queue is real time lost) and $\frac{1}{2} \delta_{\text{eff}} N^2/s$ with the optimal toll (the queue is eliminated, leaving only schedule delay). All three objects (private cost, social cost, and money toll) are independent of α_{eff} , while the peak queue $T_{\text{max}} = C^*/\alpha_{\text{eff}}$ rises as α_{eff} falls. The claim is made for $\alpha_{\text{eff}} > \beta_{\text{eff}}$, the regime in which the early side of the FIFO queue is well ordered; behaviour at or below β_{eff} is not asserted (see the Boundary remark). ■

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 no-toll total cost $\delta_{\text{eff}} N^2/s$ and the tolled total cost $\frac{1}{2} \delta_{\text{eff}} N^2/s$

are increasing in δ_{eff} , hence decreasing as β_{eff} falls. The early window $t_e = \frac{\gamma}{\beta_{\text{eff}} + \gamma} \frac{N}{s}$ is decreasing in β_{eff} , so it lengthens as β_{eff} falls. ■

Proof of Proposition 5. Let gross surplus be the area under inverse demand, $\int_0^N P(n) dn = aN - \frac{1}{2}bN^2$. The private generalised cost of a trip is $C(N) = \delta_{\text{eff}}N/s + \alpha_{\text{eff}}T_f$ (queue + schedule + free-flow). *No-toll equilibrium:* travellers enter until $P(N) = C(N)$, i.e. $a - bN = \delta_{\text{eff}}N/s + \alpha_{\text{eff}}T_f$, giving $N^{nt} = (a - \alpha_{\text{eff}}T_f)/(b + \delta_{\text{eff}}/s)$, decreasing in α_{eff} . *System optimum:* the total model social cost is schedule delay $\frac{1}{2}\delta_{\text{eff}}N^2/s$ (queue is deadweight, schedule delay is half of $\delta_{\text{eff}}N/s$ per user) plus free-flow $\alpha_{\text{eff}}T_fN$ plus, untolled, the queue deadweight $\frac{1}{2}\delta_{\text{eff}}N^2/s$. The planner maximises $aN - \frac{1}{2}bN^2 - \frac{1}{2}\delta_{\text{eff}}N^2/s - \alpha_{\text{eff}}T_fN$, with marginal social cost $\text{MSC}(N) = \delta_{\text{eff}}N/s + \alpha_{\text{eff}}T_f$. (For the fine-tolled system optimum the queue is removed, so the resource cost is schedule delay plus free-flow time; the no-toll allocation has the same volume but adds the queue deadweight $\frac{1}{2}\delta_{\text{eff}}N^2/s$.) *Fine toll:* a time-varying toll equal to the would-be queue removes the queue and makes the private price equal MSC; since $\text{MSC}(N)$ coincides with the no-toll $C(N)$, the equilibrium volume is **unchanged**, $N^{\text{toll}} = N^{nt} \equiv N^*$ (the classic bottleneck result). The welfare *decomposition* at N^* is therefore: gross surplus $aN^* - \frac{1}{2}bN^{*2}$; minus free-flow $\alpha_{\text{eff}}T_fN^*$; minus schedule delay $\frac{1}{2}\delta_{\text{eff}}N^{*2}/s$; minus, untolled only, the queue deadweight $\frac{1}{2}\delta_{\text{eff}}N^{*2}/s$, which the toll converts to revenue (a transfer). The optimal peak toll is the queue component $\delta_{\text{eff}}N^*/s$, rising as N^* rises, i.e. as α_{eff} falls. Thus the welfare gain from tolling is the eliminated queue deadweight $\frac{1}{2}\delta_{\text{eff}}N^{*2}/s$, and automation (lower α_{eff}) raises N^* , gross surplus, the toll, and physical congestion together. ■

Proof of Proposition 6. Two classes with common β, γ and $\alpha_{\text{lo}} < \alpha_{\text{hi}}$ share a single FIFO queue over an arrival window $[a_{\min}, a_{\max}]$ of length N/s (arrivals leave the bottleneck at rate s throughout the busy period). The extreme arrivals face no queue, so the high- α class, which occupies the shoulders, as shown next, has boundary cost $C_{\text{hi}} = \delta_{\text{eff}}N/s$, $\delta_{\text{eff}} = \beta\gamma/(\beta + \gamma)$, with $a_{\min} = -\frac{\gamma}{\beta + \gamma} \frac{N}{s}$ 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; since $\alpha_{\text{lo}} < \alpha_{\text{hi}}$ the low- α profile is steeper, so the low- α class tolerates the longest queues and occupies the central interval $[-x_e, x_l]$ around t^* while the high- α class takes the shoulders $[a_{\min}, -x_e] \cup [x_l, a_{\max}]$. Writing $A_1 = \alpha_{\text{lo}}\beta/\alpha_{\text{hi}}$ and $A_2 = \alpha_{\text{lo}}\gamma/\alpha_{\text{hi}}$, queue continuity and equality of the low- α cost across the two entries to the centre give the linear system

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

with a unique solution (x_e, x_l) . The allocation is *interior*, with a central low- α band between two high- α shoulders, iff $0 < x_e < |a_{\min}|$ and $0 < x_l < a_{\max}$, which holds for moderate low- α shares; as $N_{\text{lo}} \rightarrow 0$ the central band vanishes and as $N_{\text{lo}} \rightarrow N$ the shoulders vanish (the population becomes effectively homogeneous). On the interior the low- α class occupies the interval around t^* , so its mean absolute arrival deviation is strictly smaller. At the calibration the replication script returns $x_e = 1.07$, $x_l = 0.18$ h and mean $|a|$ of 0.47 h (low- α) versus 1.42 h (high- α). ■

Proof of Proposition 7. By Lemma 1, the smooth boundary equations are $g(a_e) = g(a_l) = C^*$ and $a_l - a_e = N/s$; with g strictly convex these two equations determine (a_e, a_l) , and hence C^* , uniquely and independently of α_{eff} . Interior cost equalization, $\alpha_{\text{eff}}T(a) + g(a) = C^*$, gives $T(a) = (C^* - g(a))/\alpha_{\text{eff}} \geq 0$ with peak $T(0) = C^*/\alpha_{\text{eff}}$ (since $g(0) = 0$). First-in-first-out requires the home-departure time $a - T(a)$ to be increasing, i.e. $T'(a) < 1$; since $T'(a) = -g'(a)/\alpha_{\text{eff}}$ and $-g'(a)$ is maximised at the early edge a_e , the binding condition is $-g'(a_e)/\alpha_{\text{eff}} < 1$, i.e. $\alpha_{\text{eff}} > \alpha_{\text{crit}} \equiv -g'(a_e)$. On the late side $-g'(a) < 0$, so $T'(a) < 0$ and FIFO is automatic; only the early edge binds and there is no upper bound on α_{eff} , just as the kinked bottleneck requires only $\alpha_{\text{eff}} > \beta_{\text{eff}}$ and imposes no upper bound on the value of queue time. On that range the profile is an equilibrium: costs are equalized, FIFO holds, and an arrival outside $[a_e, a_l]$ faces no queue at strictly higher schedule cost $g > C^*$, so no deviation pays. The fine toll $\tau(a) = \alpha_{\text{eff}}T(a) = C^* - g(a)$, the no-toll social cost NC^* (the queue is real time lost), and the tolled social cost $s \int_{a_e}^{a_l} g(a) da$ (the toll removes the queue) inherit independence of α_{eff} from C^* and the window. At $\alpha_{\text{eff}} = \alpha_{\text{crit}}$, $T'(a_e) = 1$ (the early entrance rate diverges); for $\alpha_{\text{eff}} < \alpha_{\text{crit}}$, $T'(a_e) > 1$, the home-departure ordering reverses near the early edge, and no single-pass FIFO equilibrium of this form exists. ■

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: $C^* = \delta_{\text{eff}}N/s$, independent of how waiting is costed. The optimal fine toll replaces queue time by a money charge that holds each traveller's generalised cost at C^* ; since the schedule penalty $\text{sched}(a)$ and C^* are unchanged, the toll $\tau(a) = C^* - \text{sched}(a)$ is unchanged, and so are the no-toll and tolled social costs, which depend only on C^* and the window. The interior wait solves $K(T(a)) = C^* - \text{sched}(a)$, so $T(a) = K^{-1}(C^* - \text{sched}(a))$ by monotonicity of K ; the early-side first-in-first-out ordering holds because $K'(T) \geq \beta_{\text{eff}}$ (the constant- α case is $K(T) = \alpha_{\text{eff}}T$, with $K' = \alpha_{\text{eff}}$). When K is convex the marginal wait cost rises, so the cumulative cost reaches $C^* - \text{sched}(a)$ at a shorter wait near the peak, and the peak queue is smaller than the constant-value baseline C^*/α_0 . ■

Proof of Corollary 2. Under the Fosgerau–Karlström (2010) assumptions (optimal departure, scheduling cost piecewise linear in clock deviation), expected cost at the optimal head start is linear in the mean and standard deviation of travel time, $\mathbb{E}[C^*] = \alpha\mu_T + \sigma_T\kappa(\beta, \gamma)$, where κ depends only on the schedule parameters and the standardised delay distribution. Replacing α, β by the effective $\alpha_{\text{eff}}, \beta_{\text{eff}}$ of Proposition 2 yields the stated expression; $\kappa(\beta_{\text{eff}}, \gamma)$ inherits monotonicity in β_{eff} , so productive early arrival lowers the value of reliability. The result is a substitution into an established theorem and does not cover the case in which the in-trip production decision is taken after the delay is realised. ■

Appendix C: Calibration

The numerical illustrations are stylised. All figures and the quoted numbers are produced by replication scripts available on the author’s website alongside the working paper, under the calibration in Table 3, in US dollars and hours. The values are chosen for clear, regime-consistent illustration and are not calibrated to a particular corridor; the magnitudes should not be read as empirical estimates of automation welfare or toll effects. Capacity is measured in travellers per hour, equivalent here to vehicles per hour under one traveller per vehicle, so $N/s = 2.5$ h. The queue-forming condition $\alpha_{\text{eff}} > \beta_{\text{eff}}$ holds for both classes (under the maintained empirical ordering $\gamma > \alpha_{\text{eff}} > \beta_{\text{eff}}$: $48 > 26 > 8$ for driving, $48 > 10 > 8$ for automated), so Proposition 3 applies to each as a single-class (homogeneous) counterfactual; the simultaneous two-class case, where both technologies share one FIFO bottleneck, is the sorting environment of Proposition 6. The mode-choice illustration (Figure 2) uses a binary logit with one-way travel time 0.6 h, money prices $p_{\text{drive}} = \$4$ and $p_{\text{auto}} = \$9$, no alternative-specific constant, and a logit scale of \$6. The smooth-scheduling illustration (Section 5.5, Figure 7) uses a piecewise-quadratic schedule $g(a) = \frac{1}{2}k_e a^2$ (early), $\frac{1}{2}k_l a^2$ (late), with $k_e = 4$ and $k_l = 90$, which places the queueing boundary at $\alpha_{\text{crit}} \approx 8$. This proximity to $\beta_{\text{eff}} = 8$ is a calibration choice, not a structural identity; α_{crit} depends on N/s and the schedule curvatures. The queue-dependent value-of-time illustration for Proposition 8 uses $K(T) = \alpha_0 T + \frac{1}{2}\kappa T^2$ with $\alpha_0 = 20$ and $\kappa = 8$.

Table 3: Stylised calibration of the numerical illustrations.

Parameter	Value	Units
ω resource value of time	20	\$/h

Parameter	Value	Units
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
a 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
κ accumulated-wait slope (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>
- 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>
- 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., & 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>
- 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)
- 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>
- 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>