



Volume 46, Issue 2

Immigration, integration pipelines, and local innovation intensity

Jinghong Gu

Faculty of Economics, Saga University, 1 Honjo-machi, Saga, 840-8502, Japan.

Abstract

This note derives a local ranking result for immigration and innovation within a two-stage integration pipeline. Immigrants first enter an unintegrated state, then transition to a partially integrated stage, and finally reach full integration. Each transition is governed by a stage-specific hazard rate that declines with congestion, reflecting capacity constraints in language training, credential recognition, licensing, and innovation-network access. At the no-immigration baseline, a proportional expansion of later-stage capacity raises the local slope of innovation intensity more than the same proportional expansion of earlier-stage capacity, if and only if fully integrated workers carry a sufficiently large innovation weight relative to their partially integrated counterparts. This ranking has no analogue in a one-stage benchmark, where proportional easing merely rescales a single transition margin. An auxiliary result shows that, under a common-congestion schedule, the model can generate a hump-shaped pattern for innovation intensity when the local-gain condition holds. An illustrative parameterization demonstrates that the late-minus-early gap widens when selection into full integration is stronger or the later bottleneck is tighter. These findings suggest that, for innovation intensity, where capacity is expanded within the integration pipeline matters as much as how much capacity is added overall.

Submitted: March 12, 2026.

1. Introduction

Immigration is central to debates on innovation and long-run growth. Recent research documents positive average effects of immigration on local innovation and related outcomes and documents the roles of inventors, firms, and knowledge diffusion in shaping those gains (Terry et al., 2026; Glennon, 2024; Lissoni and Miguelez, 2024). Much less is known about the transition from arrival to innovation effectiveness: how quickly newcomers move through language support, recognition, occupational sorting, and network entry before they become innovation-effective workers. To address that gap, this note derives a *local ranking result* (Corollary 1): at the no-immigration baseline, proportional easing of later-stage capacity can raise innovation intensity by more than equal easing of earlier-stage capacity.

The note models that transition as a two-stage integration pipeline. Language acquisition is gradual and shapes labor-market performance (Chiswick and Miller, 2001; Dustmann and Fabbri, 2003). Dustmann and Glitz (2011) emphasize that host-country human capital and skill transferability shape later occupational integration, credential recognition improves interview and hiring prospects for skilled jobs (Damelang et al., 2020), and access to language training remains a policy margin for local integration (Foged and van der Werf, 2023). Taken together, these findings suggest treating integration as a sequence of stage hazards rather than as a single instantaneous transition. They also map naturally into concrete reform margins, such as foreign-credential recognition in Canada or language-training provision in Germany.

The contribution is therefore not non-monotonicity by itself. A one-stage model can also generate a hump-shaped schedule. What the two-stage pipeline adds is a ranking across distinct policy margins. Corollary 1 shows that, at the no-immigration baseline and for proportional changes in baseline capacities, easing the later stage raises the local slope of innovation intensity by more than equal easing of the earlier stage when fully integrated workers carry the larger innovation weight. The main object is *innovation intensity*—an innovation-output-per-worker concept—rather than aggregate idea production, and the level benchmark is reported only as a scale-effect comparison. Proposition 1 is therefore auxiliary: it shows that the benchmark congestion schedule can also generate non-monotonicity. Section 2 presents the model, Section 3 states the local ranking and the auxiliary existence result, Section 4 provides an illustrative parameterization and benchmark peak comparisons, and Section 5 summarizes the policy lesson, empirical roadmap, and extensions. Appendix A records the one-stage benchmark, Appendix B reports how the threshold conditions and derivative formulas move with φ and δ , and Appendix Figure A1 collects fuller comparative statics under the benchmark schedule.

2. Model

Time is continuous. The stock of native workers is fixed at $N > 0$. Immigrants arrive at rate μN , where $\mu \geq 0$ is the inflow rate relative to the native stock. Workers exit the working-age population at rate $\delta > 0$.

New arrivals start unintegrated. Let U_t denote unintegrated immigrants, M_t partially integrated immigrants, and H_t fully integrated immigrants. The interpretation is economic and network integration: M_t corresponds to immigrants who have completed basic language or orientation support and can enter intermediate jobs, but are still waiting for the final margins that unlock full innovation relevance, such as professional licensing, full credential recognition, or stable

entry into host-country innovation networks; H_t corresponds to immigrants who have entered innovation-relevant jobs and those networks. In the applications emphasized below, low-inflow capacities satisfy $\bar{a}_1 > \bar{a}_2$ because classroom-based language or orientation support is often completed faster than later licensing, credential recognition, or professional-network entry; the latter margins are typically more selective, more incumbent-dependent, and more institution-intensive. The transition law is

$$\dot{U}_t = \mu N - (\delta + a_1)U_t, \quad (1)$$

$$\dot{M}_t = a_1 U_t - (\delta + a_2)M_t, \quad (2)$$

$$\dot{H}_t = a_2 M_t - \delta H_t. \quad (3)$$

The hazards a_1 and a_2 are stage-specific integration rates.

For tractability, the note uses a common-congestion stationary benchmark. Let $I_t \equiv U_t + M_t + H_t$ denote the total immigrant stock. Summing (1)–(3) gives

$$\dot{I}_t = \mu N - \delta I_t,$$

so stationarity implies

$$I(\mu) = \frac{\mu N}{\delta}, \quad \sigma(\mu) = \frac{I(\mu)}{N + I(\mu)} = \frac{\mu}{\mu + \delta}.$$

Thus the total immigrant stock is pinned down by inflows and exits, while the stage allocation of that stock is pinned down by the hazards. As a benchmark specification, suppose both hazards take the schedule

$$a_k(\mu) = \bar{a}_k (1 - \sigma(\mu))^p = \bar{a}_k \left(\frac{\delta}{\mu + \delta} \right)^p, \quad \bar{a}_k > 0, \quad p > 0, \quad k \in \{1, 2\}. \quad (4)$$

This specification captures congestion effects in a reduced-form manner. The same inflow wave can be expected to put simultaneous pressure on classrooms, administrative processing systems, recognition offices, and matching institutions, so a shared congestion index parsimoniously captures that joint pressure without imposing institutional equivalence across stages. Using the same curvature p at both stages is likewise a benchmark normalization: it isolates stage order rather than heterogeneous queue technologies. More generally, one could write $a_k(\mu) = \bar{a}_k f_k(\mu)$ with stage-specific queue objects and distinct curvatures p_1, p_2 . The parameter p governs how sharply capacity deterioration sets in: $p = 1$ is the proportional-crowding benchmark, $p > 1$ implies nonlinear queueing or capacity constraints, and $p < 1$ implies sub-proportional crowding.

At stationarity, per-native stage ratios satisfy

$$u(\mu) = \frac{\mu}{\delta + a_1(\mu)}, \quad m(\mu) = \frac{a_1(\mu)}{\delta + a_2(\mu)} u(\mu), \quad h(\mu) = \frac{a_2(\mu)}{\delta} m(\mu). \quad (5)$$

Effective R&D labor is

$$R = s(N + \varphi M + \psi H), \quad s > 0, \quad 0 \leq \varphi < \psi,$$

where partially integrated immigrants contribute with relative innovation weight φ and fully integrated immigrants with relative innovation weight ψ .

The main object is *innovation intensity*, defined as per-worker innovation growth:

$$g^{sf}(\mu) = \eta \frac{R}{L}, \quad L = N + U + M + H, \quad (6)$$

which yields the normalized stationary ratio

$$G^{sf}(\mu) = \frac{g^{sf}(\mu)}{g^{sf}(0)} = \frac{1 + \varphi m(\mu) + \psi h(\mu)}{1 + \mu/\delta}. \quad (7)$$

Stationarity implies $L = N + I = N(1 + \mu/\delta)$, so the denominator in (7) is exactly $1 + \mu/\delta$. It is the *dilution margin*: if total labor expands faster than innovation-effective labor, innovation intensity falls. For comparison, a level benchmark in the spirit of Romer (1990) is

$$G^{lev}(\mu) = \frac{g^{lev}(\mu)}{g^{lev}(0)} = 1 + \varphi m(\mu) + \psi h(\mu), \quad (8)$$

which removes that dilution term.

The scale-free concept is the main object for two reasons. First, local immigration shocks expand both the effective idea-producing labor force and the denominator relevant for innovation intensity, so a per-worker ratio is the natural object for comparing places with different inflow rates. Second, it avoids reading the note as a defense of first-generation scale-effects of the sort emphasized by Jones (1995). Empirically, (7) maps most naturally into outcomes such as patents per worker, inventor employment per worker, or other innovation-output-per-worker measures rather than into aggregate patent counts. The contrast with (8) is simple: G^{sf} includes the dilution margin and can therefore fall below baseline at high inflows, whereas G^{lev} removes that dilution term and behaves like a level benchmark in the spirit of Romer (1990).

3. Main results

For any differentiable hazard schedules that can be written locally as $a_k(\mu) = \bar{a}_k f_k(\mu)$ with $f_k(0) = 1$, the local slope of (7) around $\mu = 0$ depends on

$$\Lambda \equiv \frac{\varphi \bar{a}_1}{(\delta + \bar{a}_1)(\delta + \bar{a}_2)} + \frac{\psi \bar{a}_1 \bar{a}_2}{\delta(\delta + \bar{a}_1)(\delta + \bar{a}_2)}. \quad (9)$$

Λ summarizes the first-order increase in innovation-effective labor generated by a marginal inflow at the no-immigration baseline. Because Λ is linear in φ and ψ , the local gain condition can be written as

$$\psi > \psi_{\min}(\varphi) \equiv \frac{(\delta + \bar{a}_1)(\delta + \bar{a}_2)}{\bar{a}_1 \bar{a}_2} - \frac{\delta \varphi}{\bar{a}_2}. \quad (10)$$

Equations (9)–(10) use only the baseline hazard levels $a_k(0) = \bar{a}_k$; the slopes $f'_k(0)$ do not enter because the immigrant stocks are zero at the no-immigration baseline.

Before turning to the main ranking result, note that the common-congestion benchmark can also

generate non-monotonicity.

Proposition 1. *Under the benchmark congestion schedule (4), suppose $p > 0$ and $\Lambda > 1/\delta$. Then $G^{sf}(\mu)$ exceeds 1 for some small positive inflow and attains at least one interior maximizer $\mu^* > 0$. In particular, the benchmark schedule is non-monotonic.*

Proof sketch. Evaluating (7) at zero gives $G^{sf}(0) = 1$. Because $R = s(N + \varphi M + \psi H)$ and the ratio is normalized by $g^{sf}(0) = \eta s$, the positive constants η and s cancel. A first-order expansion of (5) and (7) around $\mu = 0$ therefore gives $G^{sf'}(0) = \Lambda - 1/\delta$. Hence $\Lambda > 1/\delta$ implies a local increase above 1, so there exists $\varepsilon > 0$ with $G^{sf}(\varepsilon) > 1$. As $\mu \rightarrow \infty$, the immigrant denominator in (7) continues to grow linearly while the integration hazards collapse, so the innovation-effective share of total labor shrinks; formally, $a_k(\mu) = O(\mu^{-p})$ and thus $G^{sf}(\mu) \rightarrow 0$. Therefore there exists $M > \varepsilon$ such that $G^{sf}(\mu) < 1$ for all $\mu \geq M$. Since G^{sf} is continuous on $[0, M]$, it attains a maximum there; because $G^{sf}(\varepsilon) > 1 > \sup_{\mu \geq M} G^{sf}(\mu)$, any global maximizer must lie in $(0, M)$. The congestion-curvature parameter p does not enter Λ because congestion matters only once a positive immigrant stock has accumulated; at $\mu = 0$ it therefore affects higher-order terms, not the first-order slope. For the purposes of this note, existence is the auxiliary theoretical result because the genuinely new content comes from Corollary 1 rather than from global uniqueness. Numerically, the benchmark schedules used below display a single peak throughout the reported grids, but uniqueness is not claimed analytically. $\square$

Corollary 1. *Suppose the local hazard schedules satisfy $a_k(0) = \bar{a}_k$ and $\bar{a}_1 > \bar{a}_2$, the empirically common case in which early support clears faster than later licensing, credential recognition, or network entry. At $\mu = 0$, the relevant semi-elasticities are $\partial \Lambda / \partial \ln \bar{a}_k = \bar{a}_k \partial \Lambda / \partial \bar{a}_k$. They satisfy*

$$\frac{\partial \Lambda}{\partial \ln \bar{a}_1} = \frac{\bar{a}_1(\delta\varphi + \psi\bar{a}_2)}{(\delta + \bar{a}_1)^2(\delta + \bar{a}_2)}, \quad (11)$$

$$\frac{\partial \Lambda}{\partial \ln \bar{a}_2} = \frac{\bar{a}_1\bar{a}_2(\psi - \varphi)}{(\delta + \bar{a}_1)(\delta + \bar{a}_2)^2}. \quad (12)$$

These expressions depend only on the baseline hazard levels $\bar{a}_1$ and $\bar{a}_2$; the specific shapes of the congestion schedules affect higher-order terms but not the local ranking itself. Hence easing the later stage raises the local slope by more than the same proportional easing of the earlier stage if and only if

$$\psi > \psi^\dagger(\varphi) \equiv \frac{\varphi(\bar{a}_1\bar{a}_2 + 2\delta\bar{a}_2 + \delta^2)}{\bar{a}_2(\bar{a}_1 - \bar{a}_2)}. \quad (13)$$

Interpretation. Corollary 1 is the distinctive implication of the two-stage pipeline, but it is a *local* implication: it ranks proportional changes in baseline capacities at $\mu = 0$ and does not, by itself, rank global peaks or welfare. The comparison is stated in proportional terms because it expands the integration-specific transition margin by the same percentage across stages, which is the cleanest common benchmark when competing exit δ is present. The threshold (13) has a transparent economic reading. Because $\bar{a}_1 - \bar{a}_2$ appears in its denominator, the condition is easier to satisfy when the early stage already clears materially faster than the late stage. In that case, early-stage easing mainly expands the intermediate pool M , whereas late-stage easing releases already prepared workers into the high-weight state H . The result is therefore non-trivial: Appendix A shows that the

one-stage benchmark has only one transition margin to rescale, so no analogue of the ranking exists there. In practical terms, if language training is available but licensing, credential recognition, or network entry remains slow, easing the latter margin can raise baseline innovation intensity by more than an equal proportional easing of the earlier stage.

The local ranking itself does not depend on choosing (7) over (8); the common-congestion benchmark (4) is used to state Proposition 1 compactly and to organize the numerical illustration. Under the level benchmark, the same congestion mechanism can also generate a hump-shaped schedule, but for $p > 1$ the path returns toward baseline from above rather than crossing below it. The stronger statement that very large inflows lower innovation intensity is therefore a statement about the scale-free object in (7), not about every possible growth concept. Appendix A records the one-stage benchmark, in which proportional easing simply rescales a single transition margin and no analogue of Corollary 1 arises.

4. Illustrative parameterization

The baseline parameterization is $\delta = 0.03$, $\bar{a}_1 = 0.60$, $\bar{a}_2 = 0.30$, $p = 5$, $\varphi = 0.60$, and $\psi = 2.20$. This parameterization is illustrative, not structural. The benchmark makes the early transition margin roughly twice as large as the later one, consistent with evidence that language and orientation support usually clear faster than licensing, credential recognition, or professional matching (Chiswick and Miller, 2001; Dustmann and Fabbri, 2003; Damelang et al., 2020; OECD/European Commission, 2023). With $p = 1$, crowding is proportional. Values $p > 1$ imply increasingly sharp deterioration once queues form, so $p = 5$ should be read as an upper-end queue-formation case rather than a point estimate. Appendix Figure A1(c) shows that moving from $p = 1$ to $p = 10$ mainly shifts the location of the benchmark non-monotonic schedule rather than the local ranking itself. The weights $\varphi = 0.60$ and $\psi = 2.20$ are not wage multiples or worker-level productivity estimates. They are reduced-form *innovation weights* that summarize how strongly workers in each stage are connected to frontier firms, innovation-intensive occupations, and inventor networks. The comparison is therefore between the average native worker and a *selected subset* of immigrants who have cleared both bottlenecks and sorted into innovation-relevant matches. Values of ψ exceeding unity are therefore plausible even when average worker productivity gaps are modest: they say that the average member of H is more innovation-relevant than the average native worker, not that every fully integrated immigrant is twice as productive. Terry et al. (2026) and Glennon (2024) motivate that interpretation by showing that the measured gains from immigration are concentrated in high-innovation firms and high-skill innovation matches rather than spread uniformly across all arrivals. In that sense, $\psi = 2.20$ is a conservative upper-end benchmark and $\psi = 3.0$ is a high-end scenario for stronger concentration. Under the baseline values $(\delta, \bar{a}_1, \bar{a}_2, \varphi) = (0.03, 0.60, 0.30, 0.60)$, the thresholds are $\psi_{\min} = 1.095$ and $\psi^\dagger = 1.326$, so the qualitative results are disciplined by inequalities rather than by one exact calibration; Appendix B records how those thresholds move with φ and δ . Because $\delta = 0.03$, the baseline peak near $\mu \approx 0.76\%$ corresponds to a stationary immigrant stock ratio $I/N \approx \mu/\delta \approx 0.25$, so seemingly small annual inflow rates should be read together with their implied stock counterparts.

Table I: Illustrative magnitudes for the stage-specific ranking

Scenario	Baseline peak	Peak after +25% $\bar{a}_1$	Peak after +25% $\bar{a}_2$	Late–early gap
Baseline ($\psi = 2.2$, $\bar{a}_2 = 0.30$)	1.117	1.126	1.128	0.002
Higher innovation weight ($\psi = 3.0$)	1.231	1.241	1.246	0.005
Tighter later stage ($\bar{a}_2 = 0.20$)	1.097	1.103	1.108	0.005
Both ($\psi = 3.0$, $\bar{a}_2 = 0.20$)	1.190	1.201	1.211	0.009

Notes: All rows keep $\delta = 0.03$, $\bar{a}_1 = 0.60$, $p = 5$, and $\varphi = 0.60$ fixed. Entries report the peak of $G^{sf}(\mu)$ under the common-congestion benchmark and after a proportional 25% increase in the indicated capacity. The table illustrates the benchmark ranking and should not be read as a quantitative forecast; under stage-specific congestion schedules, global magnitudes and even peak ordering can move even though Corollary 1 remains valid locally.

Table I moves from the local theorem to a benchmark global comparison. Corollary 1 does *not* prove peak ordering; it ranks local slopes at $\mu = 0$. The table therefore reports only what the common-congestion benchmark implies for global peaks. Under the baseline parameterization, the peak occurs at about 0.76% of the native stock per year and returns to baseline at roughly 1.65%. Raising $\bar{a}_1$ alone and $\bar{a}_2$ alone both move the peak to about 0.80%, but later-stage easing raises the peak slightly more. The baseline wedge is deliberately conservative — about 0.002 in peak height, or roughly 0.2 percent of the no-immigration benchmark — because the late bottleneck is only moderately tighter than the early one and ψ lies only modestly above $\psi^\dagger$. The point of the table is therefore not the exact baseline magnitude; it is the sign of the ranking and the way it widens when either selection into H becomes stronger or the later bottleneck becomes tighter. Consistent with that reading, raising ψ to 3.0 or tightening the later stage from $\bar{a}_2 = 0.30$ to $\bar{a}_2 = 0.20$ more than doubles the late-minus-early wedge. Appendix Figure A1(a) visualizes the near overlap in the baseline comparison, panel (b) contrasts innovation intensity with the level benchmark, panel (c) shows how the benchmark hump shifts with p , and panel (d) plots the late-minus-early gap against ψ under two values of $\bar{a}_2$.

5. Conclusion

The note delivers one central result and one auxiliary benchmark. The central result is Corollary 1: at the no-immigration baseline, proportional easing of the later stage raises innovation intensity by more than equal easing of the earlier stage when fully integrated workers carry the larger innovation weight and $\bar{a}_1 > \bar{a}_2$. The auxiliary result is Proposition 1: under the benchmark congestion schedule, the same two-stage pipeline can also generate a non-monotonic schedule for innovation intensity when the local-gain condition $\Lambda > 1/\delta$ holds. Appendix A clarifies why the ranking is absent from the one-stage benchmark, and Appendix B shows how the threshold conditions and derivative formulas move with φ and δ .

That ranking is a local comparison of long-run capacities, not a global welfare ranking. In settings where classroom support clears relatively quickly but licensing, recognition, or network-entry margins remain tight, the model points to those later margins as the more responsive local lever for innovation intensity. Table I suggests that, under the common-congestion benchmark, the same ordering can persist into benchmark peak comparisons, but that extension is illustrative rather than general and depends on the benchmark congestion technology.

Longitudinal administrative data — such as Canada’s IMDB/LAD linkages or Germany’s IAB-

SOEP-type panels — would support estimation of a_1 and a_2 through survival models or policy-event studies based on language-course expansions, recognition-office staffing, or processing reforms. A second step could combine employer frontier-ness, linked patent records, or inventor-network indicators with stage information to recover the innovation weights attached to M and H . In that sense, the theory maps most directly into patents per worker, inventor employment per worker, or related innovation-output-per-worker measures.

The analysis is deliberately stationary and the values of p and ψ are illustrative rather than structural. A natural extension is to embed the same pipeline in a transitional setting, where new inflows initially accumulate in U and M and generate classroom queues, recognition backlogs, temporary underemployment, and delayed entry into innovation-intensive teams before the new long-run allocation is reached. Another extension would let stage-specific reforms change the innovation weights themselves by opening access to frontier firms or inventor networks. Even with those extensions left for future work, one implication remains robust: for innovation intensity, what matters is not only how much capacity is expanded, but where in the integration pipeline it is expanded.

Appendix A. One-stage benchmark

If immigrants move only from U to H , then

$$\dot{U}_t = \mu N - (\delta + a)U_t, \quad \dot{H}_t = aU_t - \delta H_t,$$

with $a(\mu) = \bar{a}(\delta/(\mu + \delta))^p$. Stationarity implies

$$u_1(\mu) = \frac{U}{N} = \frac{\mu}{\delta + a(\mu)}, \quad h_1(\mu) = \frac{H}{N} = \frac{a(\mu)}{\delta} u_1(\mu) = \frac{\mu a(\mu)}{\delta(\delta + a(\mu))},$$

and therefore

$$G_1^{sf}(\mu) = \frac{1 + \psi h_1(\mu)}{1 + \mu/\delta}.$$

The local-gain term is $\Lambda_1 = \psi \bar{a} / [\delta(\delta + \bar{a})]$. A proportional increase in $\bar{a}$ shifts the single transition schedule uniformly. There is therefore no analogue of Corollary 1: the one-stage model has no separate late-stage margin whose easing can dominate locally.

Appendix B. Threshold dependence and derivative formulas

The two threshold functions summarize how parameter changes move the qualitative results. A useful intermediate form is

$$\Lambda = \frac{\bar{a}_1(\delta\varphi + \psi\bar{a}_2)}{\delta(\delta + \bar{a}_1)(\delta + \bar{a}_2)}. \quad (\text{B.1})$$

Differentiating (B.1) with respect to $\ln \bar{a}_1$ and $\ln \bar{a}_2$ gives

$$\frac{\partial \Lambda}{\partial \ln \bar{a}_1} = \bar{a}_1 \frac{\partial \Lambda}{\partial \bar{a}_1} = \frac{\bar{a}_1(\delta\varphi + \psi\bar{a}_2)}{(\delta + \bar{a}_1)^2(\delta + \bar{a}_2)}, \quad (\text{B.2})$$

and

$$\frac{\partial \Lambda}{\partial \ln \bar{a}_2} = \bar{a}_2 \frac{\partial \Lambda}{\partial \bar{a}_2} = \frac{\bar{a}_1 \bar{a}_2 (\psi - \varphi)}{(\delta + \bar{a}_1)(\delta + \bar{a}_2)^2}, \quad (\text{B.3})$$

which are the expressions used in Corollary 1. Solving (B.3)>(B.2) for ψ yields (13).

From (10),

$$\frac{\partial \psi_{\min}}{\partial \varphi} = -\frac{\delta}{\bar{a}_2} < 0, \quad \frac{\partial \psi_{\min}}{\partial \delta} = \frac{2\delta + \bar{a}_1 + \bar{a}_2 - \varphi \bar{a}_1}{\bar{a}_1 \bar{a}_2}.$$

Under the maintained range $0 \leq \varphi \leq 1$, this derivative is positive because

$$2\delta + \bar{a}_1 + \bar{a}_2 - \varphi \bar{a}_1 = 2\delta + \bar{a}_2 + (1 - \varphi)\bar{a}_1 > 0.$$

Thus larger intermediate innovation weight φ makes a hump-shaped schedule easier to obtain, while larger turnover δ makes it harder. From (13),

$$\frac{\partial \psi^\dagger}{\partial \varphi} = \frac{\bar{a}_1 \bar{a}_2 + 2\delta \bar{a}_2 + \delta^2}{\bar{a}_2(\bar{a}_1 - \bar{a}_2)} > 0, \quad \frac{\partial \psi^\dagger}{\partial \delta} = \frac{2\varphi(\bar{a}_2 + \delta)}{\bar{a}_2(\bar{a}_1 - \bar{a}_2)} > 0.$$

Hence later-stage dominance is easier to obtain when fully integrated workers carry much larger innovation weights than intermediate workers and when exit is slower. With $(\delta, \bar{a}_1, \bar{a}_2, \varphi) = (0.03, 0.60, 0.30, 0.60)$, later-stage dominance starts only once ψ exceeds about 1.326, so values only moderately above that cutoff generate a small but systematic wedge. These threshold relations are why the ranking result is disciplined by inequalities, not by one particular numerical calibration; Table I and Appendix Figure A1(d) simply illustrate how that disciplined ranking looks under the common-congestion benchmark.

References

- Chiswick, B. R., and P. W. Miller (2001) “A Model of Destination-Language Acquisition: Application to Male Immigrants in Canada” *Demography* 38, 391-409.
- Damelang, A., S. Ebersperger, and F. Stumpf (2020) “Foreign Credential Recognition and Immigrants’ Chances of Being Hired for Skilled Jobs—Evidence from a Survey Experiment among Employers” *Social Forces* 99, 648-671.
- Dustmann, C., and F. Fabbri (2003) “Language Proficiency and Labour Market Performance of Immigrants in the UK” *Economic Journal* 113, 695-717.
- Dustmann, C., and A. Glitz (2011) “Migration and Education” in *Handbook of the Economics of Education*, Vol. 4, E. A. Hanushek, S. Machin and L. Woessmann, Eds., Elsevier: Amsterdam, 327-439.
- Foged, M., and C. van der Werf (2023) “Access to Language Training and the Local Integration of Refugees” *Labour Economics* 84, 102366.
- Glennon, B. (2024) “Skilled Immigrants, Firms, and the Global Geography of Innovation” *Journal of Economic Perspectives* 38, 3-26.
- Jones, C. I. (1995) “R&D-Based Models of Economic Growth” *Journal of Political Economy* 103, 759-784.
- Lissoni, F., and E. Miguelez (2024) “Migration and Innovation: Learning from Patent and Inventor

Data” *Journal of Economic Perspectives* 38, 27-54.

OECD/European Commission (2023) *Indicators of Immigrant Integration 2023: Settling In*, OECD Publishing: Paris.

Romer, P. M. (1990) “Endogenous Technological Change” *Journal of Political Economy* 98, S71-S102.

Terry, S. J., T. Chaney, K. B. Burchardi, L. Tarquinio, and T. A. Hassan (2026) “Immigration, Innovation, and Growth” *American Economic Review* 116, 828-861. DOI: 10.1257/aer.20211601.

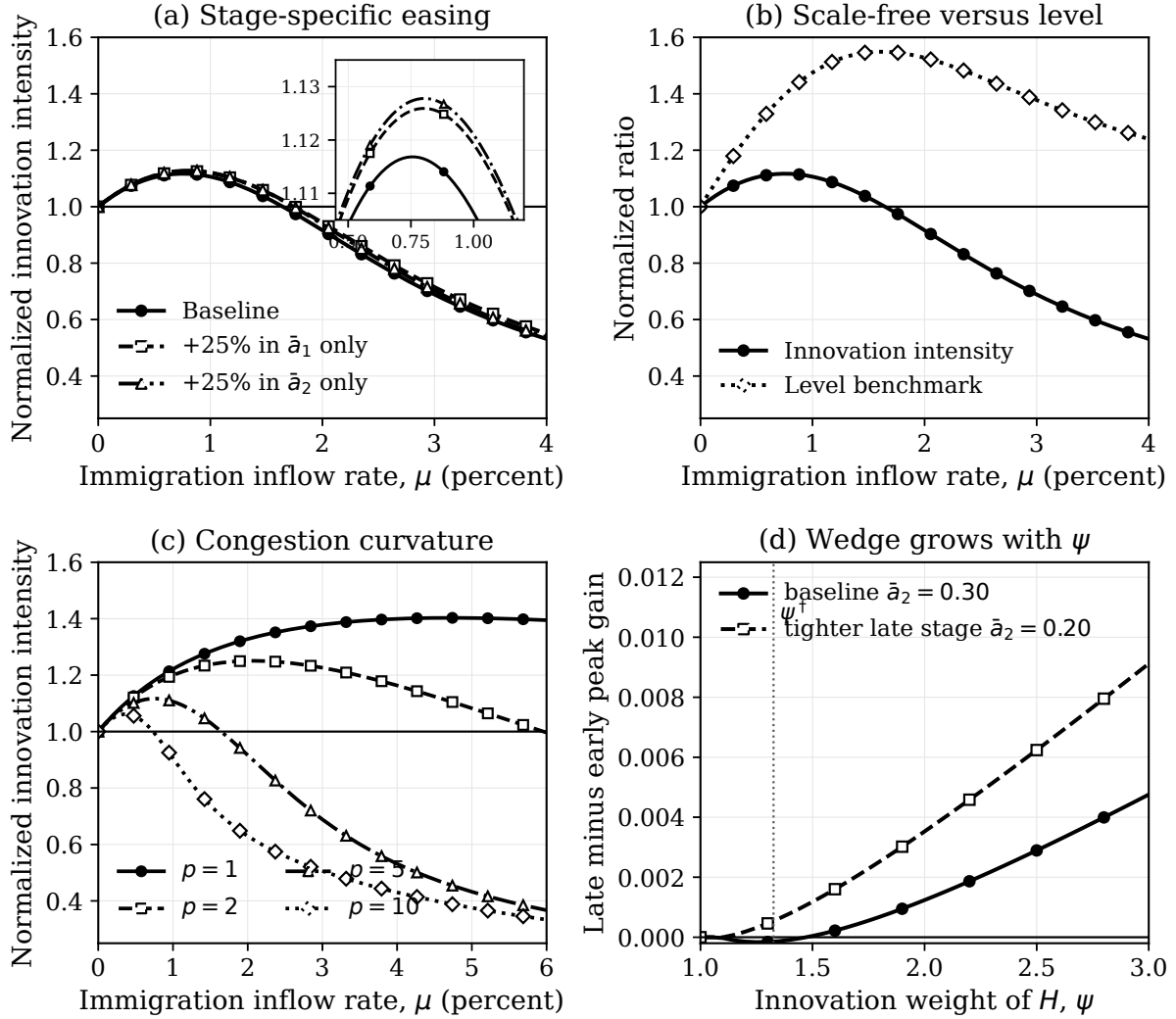


Figure A1: Comparative statics under the common-congestion benchmark. Baseline parameters are $(\delta, \bar{a}_1, \bar{a}_2, p, \varphi, \psi) = (0.03, 0.60, 0.30, 5, 0.60, 2.20)$ unless otherwise stated. Panel (a) shows stage-specific easing under the baseline and magnifies the peak region in an inset. Panel (b) contrasts innovation intensity with the level benchmark. Panel (c) varies the congestion-curvature parameter p . Panel (d) traces the late-minus-early peak wedge against the innovation weight ψ for $\bar{a}_2 = 0.30$ and $\bar{a}_2 = 0.20$; the vertical line marks the baseline threshold $\psi^\dagger = 1.326$, and Section 4 explains how that cutoff relates to Table I.