

Populace dynamics: an open, scored longitudinal layer for policy microsimulation

Max Ghenis

July 2026

Abstract

Dynamic microsimulation — the longitudinal aging of a person-level population through earnings, family structure, disability, mortality, and program participation — underpins retirement and social-insurance policy analysis, yet the benchmark United States models — DYNASIM, MINT, and CBOLT — are closed: internal to government, tied to restricted administrative records, or accessible only through institutional relationships. This design paper specifies an open alternative built as an extension of Populace, PolicyEngine’s open-source microdata stack, whose cross-sectional layer is the certified default United States microdata in PolicyEngine after outperforming its predecessor on held-out administrative targets. The design contributes four elements: a trajectory-weighted kernel in which multi-period calibration cannot silently destroy panel structure; a Dynamics operator that treats state transitions as conditional models, mixing deterministic demographic hazards with machine-learned earnings processes; an explicit domains-of-validity framework that refuses point forecasts where parameter uncertainty dominates — including the 75-year actuarial balance — and instead publishes sensitivity surfaces; and a scoring protocol under which every claim resolves against administrative publications, backtests with leakage control, or computes exactly from statute, with contributions merging only when they improve held-out scores. United States Social Security is the first validation domain; the layer itself is country-agnostic.

1 Introduction

Analysts who want to model taxes can run open, calibrated models directly: Tax-Calculator (Policy Simulation Library 2026) and PolicyEngine (PolicyEngine 2026) are openly callable on public data, alongside source-available models with restricted inputs (The Budget Lab at Yale 2026) and proprietary models used for outside-facing analysis (Tax Policy Center 2025; Institute on Taxation and Economic Policy 2025; Tax Foundation 2025; Penn Wharton Budget Model 2025). Analysts who want dynamic microsimulation — the longitudinal modeling that retirement and social-insurance policy requires — find the benchmark models closed: SSA projects retirement income and the distributional effects of policy proposals with MINT (Social Security Administration 2024), CBO produces its long-term Social Security projections with CBOLT (Congressional Budget Office 2018, 2024), the Urban Institute projects retirement income and long-term care with DYNASIM (Favreault et al. 2015; Urban Institute 2024), and Morningstar studies retirement adequacy with its Model of US Retirement Outcomes (Look and VanDerhei 2024). Outside users reach each only through an institutional relationship. The nearest open analogue, the Cato Social Security model (Chanwong 2026), simulates roughly 10,000 households from the 2007 CPS ASEC under the Social Security Administration’s assumptions and reports trust-fund metrics and reform scores, without published validation against administrative benchmarks.

This paper specifies the design of an open longitudinal layer — *Populace dynamics* — as an extension of Populace, PolicyEngine’s open-source microdata stack. It is a design paper: the cross-sectional

foundation is in production; this paper specifies the longitudinal layer, which the project builds in the open behind the scoring protocol of Section 5.

The design makes four contributions:

1. **A trajectory-weighted kernel with explicit alignment.** One weight per trajectory: multi-period calibration stacks constraint rows over the same weight vector, so that hitting cross-sectional totals in multiple periods cannot silently destroy panel structure — combined with event-selection alignment for period-by-period control, since weights alone cannot separate a correct life course from a correctly timed one (Section 4).
2. **Transitions as conditional models.** A single operator interface covers deterministic demographic hazards and machine-learned earnings processes, so candidate architectures compete under one evaluation standard (Section 4).
3. **Domains of validity as shipped metadata.** The model states which questions it will not answer with false precision — led by the 75-year actuarial balance, where input uncertainty dominates model fidelity — and publishes long-horizon results as sensitivity surfaces rather than point forecasts (Section 3).
4. **A scoring protocol in place of fidelity-only validation.** Claims resolve against administrative publications on an annual calendar, backtest against realized history with leakage control, or compute exactly from statute; contributions merge only when they improve held-out scores (Section 5).

United States Social Security is the first validation domain because benefit adequacy and reform incidence depend jointly on lifetime earnings, marriage and survivorship, disability, differential mortality, and claiming. The layer itself is country-agnostic and extends to other pension and benefit systems as PolicyEngine’s country coverage grows.

2 Terminology and scope

This paper uses “dynamic” in the microsimulation field’s standard sense, following Orcutt et al. (1961) and the tradition carried by DYNASIM, MINT, CBOLT, and SimPaths (Bronka et al. 2025): a longitudinal model that ages a person-level population through time (Li and O’Donoghue 2013). It does not mean “dynamic scoring” — the tax-policy usage denoting macroeconomic feedback in revenue estimation — and the design is not an overlapping-generations general-equilibrium model of the Auerbach–Kotlikoff kind, such as the Penn Wharton Budget Model operates (Penn Wharton Budget Model 2025). The model claims neither macroeconomic feedback nor equilibrium closure. Behavioral responses enter as labeled scenario inputs with documented ranges, not as point estimates carrying model authority.

3 Domains of validity

All models are wrong; a model earns its keep only where it improves predictions. The design therefore begins by bounding its own claims.

3.1 Parameter uncertainty dominates the long horizon

The 75-year actuarial balance of a pension system is a function of a small set of exogenous assumptions — fertility, mortality improvement, net immigration, real wage growth, interest rates — whose uncertainty dominates the microsimulation machinery that processes them. The public record makes this concrete for United States Social Security. The Trustees’ own low- and high-cost scenarios bracket a range of 75-year balances wider than the intermediate deficit itself (Board of Trustees, Federal Old-Age and Survivors Insurance and Federal Disability Insurance Trust Funds 2025). The Congressional Budget

Office’s long-term projections differ from the Trustees’ — chiefly for assumption reasons, though method differences such as CBO’s micro-founded projection of the taxable share of earnings also contribute to the divergence (Congressional Budget Office 2024). Assumption variance dominates the headline; model structure matters for distribution and the near term, which is the division of labor the output tiers encode. Successive Technical Panels convened by the Social Security Advisory Board have recommended revising fertility and mortality-improvement assumptions as realized values ran persistently outside the projected path (Technical Panel on Assumptions and Methods 2023). Two institutions with administrative data and decades of refinement disagree with each other and have both missed realized demographic trends; a better microsimulation does not repair that, because the variance lives in the inputs.

3.2 Three output tiers

The model sorts its outputs into tiers, each carrying the strongest claim it can support.

Tier 1: distributional analysis under fixed assumptions. Reform analysis is a difference — outcome under reform minus outcome under baseline, holding the population and assumption path fixed — and much of the unforecastable demographic uncertainty is common to both arms and cancels in the difference. The required ingredients — a calibrated joint distribution of lifetime earnings, family structure, and differential mortality, plus an exact rules engine — are the components the scoring protocol validates directly. The design labels rather than buries the slice of a reform delta that does not cancel: reform-induced claiming responses, and interactions between the reform and uncertain dynamics — the mortality gradient enters many deltas as a covariance with benefit position, not a level, so it does not difference out. Deltas evaluated past trust-fund depletion also require an explicit scheduled-versus-payable baseline convention, which the model states with every such output. For claiming, scenario ranges anchor to the quasi-experimental record on retirement-age responses (Mastrobuoni 2009; Behaghel and Blau 2012), and the model publishes them as a scenario library rather than embedding them as point estimates.

Tier 2: near-term components that resolve. Over roughly a ten-year horizon, mechanics rather than demographic extrapolation dominate outputs: beneficiary counts by type, average benefits, covered earnings and taxable payroll, claiming-age distributions, disability incidence. These resolve against administrative publications each year, and the protocol scores them (Section 5).

Tier 3: the long horizon as a sensitivity surface. The model publishes long-horizon outputs as surfaces over documented assumption ranges — how the balance, cohort replacement rates, or distributional outcomes move as fertility, mortality improvement, and immigration vary — never as point forecasts. The distinction is between computing and blessing: the model computes 75-year balances and depletion dates *conditional on named assumption paths*, including the Trustees’ intermediate path, so the numbers the policy debate runs on remain available as labeled conditional outputs; what the model declines is presenting any single path as its own forecast. Incumbent practice publishes the point estimate up front and the sensitivity analysis in an appendix; an open model can invert that and make the sensitivity the interface.

Every API response carries its tier, assumption path, and calibration history as metadata, so a downstream consumer — human or machine — weights the output by demonstrated reliability rather than by the producer’s reputation.

4 Architecture

4.1 The cross-sectional foundation

Populace builds a calibrated synthetic population entirely from primary-source government data — the Current Population Survey ASEC, the IRS Public Use File, the Survey of Consumer Finances, SIPP, CPS outgoing-rotation groups, MEPS, and the ACS — synthesizing missing variables with weight-aware conditional models and calibrating to administrative aggregates treated as uncertainty-weighted facts. In June 2026 it replaced PolicyEngine’s enhanced CPS as the certified default United States microdata in PolicyEngine, after a matched, symmetric-refit comparison on 41,314 households with a 739-target holdout (Table 1).

Table 1: Certification comparison from the Populace release manifest. The enhanced CPS wins more individual targets by small margins while its largest misses dominate the loss; the scoring protocol requires publishing the count that cuts against the headline.

Metric (lower is better)	Populace	enhanced CPS
Holdout loss (739 held-out targets)	0.038	0.317
Training loss	0.190	1.089
Full loss	0.228	1.405
Per-target wins	1,040	2,613 (51 ties)

4.2 Trajectory weights and population accounting

The longitudinal extension follows two kernel rules set in Populace’s charter. First, **one weight per trajectory**: multi-period calibration targets stack as (target, period) constraint rows over a single trajectory-level weight vector, so that calibrating cross-sections independently — which severs the trajectory-level consistency a panel exists to provide — is a kernel-level error rather than a modeling temptation. Second, **population is not closed**: trajectories carry entry and exit markers (birth, death, immigration, emigration), and a trajectory’s weight contributes to a period only while the person is present. Household and couple links are period-scoped, so family recomposition preserves per-period accounting identities; couples carry a shared unit weight derived from their trajectory weights so that spousal and survivor benefits have a well-defined representation.

Weights alone are one layer of alignment, not the whole answer: a weight cannot distinguish a correct life course from a correctly timed one, and reweighting trajectories to hit future cross-sectional cells risks selecting on entire correlated life courses. Period-by-period control therefore also operates through event selection — ranking individual transition probabilities and selecting the number of events an external control demands, the mechanism CBOLT and DYNASIM use — with trajectory weights reserved for base-year representation and slow-moving composition (Li and O’Donoghue 2013; Dekkers and Cumpston 2012). The kernel solves the stacked constraint system as uncertainty-weighted penalized least squares against target standard errors, so infeasible combinations resolve by SE-weighted compromise rather than silent failure, and projected demographic controls pin cohorts born after the base year rather than leaving them free.

4.3 The Dynamics operator

Dynamics is an operator from a population and a transition specification to a population with extended periods. A transition is a conditional distribution — the probability of next-period state given current state and covariates — which is the same interface Populace’s synthesis models already implement. The shipped baseline for earnings is a regime-gated, sequentially chained, weighted quantile-regression-forest imputer (Meinshausen 2006) whose zero-inflation gate doubles as a nonemployment model; richer

architectures (zero-inflated neural distribution models, normalizing flows) are candidates that must beat the baseline on held-out longitudinal moments to merge.

The design is hybrid. Where the evidence base is tabular — mortality from official life tables with published income gradients, fertility from vital statistics, marriage and divorce from ACS- and CPS-based rates (federal collection of detailed marriage and divorce statistics ended in the 1990s), disability incidence from program statistics — transitions are deterministic hazards, auditable row by row. Marriage requires a matching model, not only a hazard: spousal and survivor benefits depend on assortative matching over lifetime earnings, which the design treats as a first-class estimation target rather than an afterthought. The design reserves machine learning for processes with rich conditional structure, led by earnings dynamics — where the process is not stationary: volatility and mobility vary by cohort and period in administrative data (Sabelhaus and Song 2010; Kopczuk et al. 2010), so transition models carry cohort and period conditioning rather than pooling across decades. Chained one-period models also understate long-spell persistence unless spell structure enters the model explicitly, and backcasting is a distinct conditional object from forward simulation rather than the same operator reversed; both enter the evaluation as targets in their own right, disciplined by held-out panel moments such as higher-order earnings-change distributions (Guvenen et al. 2021).

The design states one measurement caveat rather than hiding it: the most demanding earnings-dynamics moments come from administrative records that public panels understate, and survey- and administrative-based estimates disagree on volatility levels and trends. Where the project cannot recompute a published administrative moment on held-out public data, matching it is calibration, not validation, and the scorecard labels it as such.

4.4 Rules and delivery

Statute is the deterministic slice of any policy forecast. Core retirement benefit formulas — average indexed monthly earnings, primary insurance amounts, actuarial adjustments — and benefit taxation compute exactly today through PolicyEngine’s rules engine via Populace’s rules-adapter protocol, vectorized over person-periods. Auxiliary, spousal, and survivor benefit formulas are explicit build items in the validation program: the current engine carries them as calibrated aggregates rather than person-level formulas, and the scorecard treats “computes exactly” as a per-rule status each formula earns, not a blanket claim.

The rules adapter is engine-agnostic. PolicyEngine-US implements it today; Axiom — an open project that encodes statute declaratively and compiles it to Rust — is the next adapter, and the performance headroom matters when benefit formulas run over person-periods across the full trajectory panel and many reform scenarios. In that architecture, PolicyEngine is a composition: Axiom supplies the rules, Populace supplies the population, and behavioral responses enter as the labeled scenario layer.

Data governance is a design requirement, not an afterthought. Estimating transition models on restricted-use panels and publishing the estimated parameters is settled practice — open models such as OG-USA, and DYNASIM itself, estimate on the PSID and publish what they learn. Releasing a synthetic *microdata* artifact informed by such panels is a stricter problem, because donor-based samplers can emit observed training values. The design answers it structurally: released records originate from Populace’s public-use cross-section, restricted panels train processes rather than donate records, samplers smooth or noise their draws so no verbatim donor value ships, and every release passes nearest-neighbor disclosure checks alongside its accuracy scorecard — and the project engages data producers directly where their terms of use warrant it. Uncertainty budgets therefore attach only to the components statute does not fix. The deliverable is a versioned artifact — a longitudinal population release with a manifest and scorecard, certified through the same path as the cross-sectional release — exposed through a Python library, a REST API, and a Model Context Protocol server so that AI agents can run baseline distributions and reform analyses with validity metadata attached.

5 Scoring and resolution

Validation by fidelity — does the model match published aggregates? — is necessary but weak: a model can reproduce the tables its authors fit it to. This project’s standard: a claim counts as validated when it improves prediction of something that later resolves. Five scoring surfaces implement it.

1. **Annually resolving components.** Beneficiary counts by type, average and aggregate benefits, covered earnings and taxable payroll, cost-of-living adjustments, disability incidence, and claiming-age distributions resolve against administrative publications each year. Every published forecast cell carries a resolution rule naming the exact table and vintage that settles it.
2. **Forecasting the forecasters.** Official projections revise every year; predicting the next revision of headline quantities resolves in months rather than decades and is decision-relevant to anyone who acts on the official number. This is partly forecasting an assumptions process — panels advise, committees adopt with a lag — so each cell pre-specifies the naive baseline it must beat (an assumption random walk plus mechanical data update).
3. **Retrodiction with leakage control.** Retrodiction builds the model from data vintages available at a historical date and scores it against realized outcomes. Populace’s versioned data registry pins vintages going forward; pre-registry history is a reconstruction problem — survey redesigns, revised administrative tables, re-released panels — so the protocol grades pre-registry backtests as pseudo-vintage, with a published log of deviations from true vintage, and no registry pins specification leakage: a “2005-vintage” model built today knows 2008 happened, and the protocol says so. Retrodictive calibration under the historical regime does not guarantee calibration under a new one, so backtests complement rather than substitute for live resolution.
4. **Statutory resolution.** Where an output is fixed by law, the rules engine computes it exactly, and enacted policy settles the corresponding conditional cells immediately.
5. **Held-out panel moments.** The protocol scores the population layer against moments it never fit: earnings-mobility matrices, autocorrelation and higher-order moments of earnings changes, cohort age-earnings profiles, and family-transition rates on held-out panel records.

Two governance rules complete the protocol. **Merge on score:** a contribution — a mortality module, a claiming model, an earnings architecture, from any contributor — merges if and only if it improves the population’s score on held-out facts, the rule Populace already applies to its cross-sectional layer. **Publication discipline:** misses publish with the same prominence as hits; superseded methods keep their historical scorecards; and stage gates in the development roadmap are pre-specified score thresholds, not narrative judgments.

Openness supplies most of the refereeing. Resolution rules are pre-registered, scores recompute from public data, and anyone who distrusts a published scorecard can rerun it — or fork the project and publish a rival scorecard. Two pieces sit beyond reproduction. First, restricted-data checks: the most demanding earnings-history moments live in linked administrative records that the public pipeline never touches, so the reader must trust whoever runs the comparison — author or outsider — and a validator without a stake in the result adds evidence where rerunning is impossible. Second, judgment: gate thresholds, disputed resolution rules, and the assumptions library carry discretion that pre-registration narrows but does not remove. The project defines its validation procedures and gate thresholds before the components they gate, and seats an advisory board to review them — and disputes under them — in public. Independent scoring across models is a decision only a third party can make; the design encourages it — published projections from the closed models give any such body a comparison set without model access — but does not depend on it.

6 First validation domain: U.S. Social Security

Social Security is the first domain because eligibility and benefits depend on the highest thirty-five years of indexed earnings, marital and survivorship histories, disability pathways, differential mortality,

and claiming timing — jointly. A layer that scores well here earns reuse in adjacent domains — Supplemental Security Income interactions, long-term care, and retirement adequacy, the last of which is *harder*, not easier: wealth projection challenged even administrative-data models (Favreault and Smith 2016), and a wealth and pension roadmap is future work, not an assumed extension — and in other countries’ pension systems.

The domain also makes the validity framework concrete. The canonical output of existing Social Security models — the 75-year balance and depletion date — is the quantity Section 3 declines to forecast. What the open layer offers instead is the combination the field lacks: tier-1 distributional incidence of reforms with reproducible assumptions; tier-2 near-term components with a public resolution record; and tier-3 sensitivity surfaces that make the assumption-dependence of long-horizon claims the interface rather than the appendix. No existing model, closed or open, publishes that combination. The closed benchmarks do publish validation and cohort tables; what outsiders cannot do is rerun the pipeline, vary its assumptions, or audit the intermediate states behind the published numbers.

7 Gate 1 in practice: the first pre-registered runs

The protocol of Section 5 stopped being a design in July 2026. This section reports what it produced: a locked gate, five registered and scored model runs, and the findings the failures purchased. Every number below recomputes from committed artifacts in the project repository; the run log lives in the repository’s pull requests and the candidate registry in [issue #42](#).

7.1 The lock

Gate 1 — earnings-history credibility — locked on 2026-07-05. Thresholds derive from committed noise-floor artifacts measured at the scale the protocol scores at (two disjoint 20%-of-persons samples of the PSID family earnings panel scored against each other), with every derivation stated as floor mean plus a named multiple of the floor’s seed standard deviation and enforced by a test: a floor rebuild that shifts any artifact fails continuous integration rather than silently orphaning the rationale. Before ratification the proposed thresholds went through three rounds of adversarial review, published in full on the ratification pull request. Round one found that the proposed numbers did not follow their own stated rule and were calibrated against a floor at five times deployment scale; round two found the same wrong-scale defect reintroduced on a second view and demonstrated that no window-2 statistic can catch a chained one-period model, forcing the persistence guard onto the battery’s long-horizon autocorrelation bands; round three verified the amendments and ratified. The merge of the ratification pull request is the lock event; the thresholds have not changed since, and any change requires a public amendment and a fresh review round.

7.2 Five runs, four failures, and what they isolated

Each candidate registers its complete specification — every modeling degree of freedom pinned — before its single scored run. The log autocorrelation ladder at 2, 4, and 10 years (locked bands 0.730 ± 0.05 , 0.657 ± 0.06 , 0.539 ± 0.07) tells most of the story:

candidate	2yr	4yr	10yr	verdict
chained weighted QRF (baseline)	0.726	0.573	0.333	fail
+ person effect as feature	0.726	0.688	0.649	fail

candidate	2yr	4yr	10yr	verdict
+	0.722	0.695	0.647	fail
persistence-aware decomposition				
structural	0.464	0.401	0.354	fail
three-component assembly				
donor splicing	0.779	0.704	0.616	fail
(single-donor)				
donor splicing	0.720	0.631	0.490	fail
(segments)				
Gaussian-copula	0.716	0.653	0.507	fail
rank dynamics				
empirical rank	0.692	0.548	0.381	fail
kernel				
k-NN rank	0.719	0.636	0.459	fail
bootstrap				
permanent-rank	0.791	0.733	0.670	fail
matching				
calibrated blend +	0.757	0.677	0.514	fail
Q0 regime				
inner-validated	0.773	0.695	0.510	fail
composition				
inner-validated	0.773	0.695	0.510	pass
composition				
(re-registered)				

The baseline failed exactly where the pre-lock review predicted a one-step chain must: window-2 geometry passed on every seed while the 10-year autocorrelation collapsed toward the Markov value. The second and third candidates added a latent person effect as a conditioning feature — first from a naive variance decomposition, then from a persistence-aware one that halved the drawn variance — and produced statistically identical ladders. That invariance is the program’s sharpest finding to date: a quantile forest rescales its response to a person-effect feature inversely to the feature’s scale, so conditional-draw generation transmits the raw within-panel person-mean variance share (0.647 on the training splits, inflated by transitory persistence over few biennial observations) no matter how the feature is constructed. No feature-side dial can set person-level variance. The fourth candidate set the variance structurally — assembling log earnings from an age profile, a person effect drawn at the decomposed permanent variance, a chained transitory component, and an observation-noise layer — and the variance landed as designed while the marginal distribution broke: parametric lognormal assembly inflated levels where the earlier candidates had inherited the data’s marginal from empirical conditional draws by construction.

The symmetry across the four failures defines the remaining design problem. Conditional-draw candidates preserve the marginal but cannot set person-level variance; the structural candidate set the variance and lost the marginal. The data demand both at once.

The generative track then ran the composition the taxonomy below suggests. A Gaussian-copula rank model — empirical quantile marginals, a calibrated permanent-plus-transitory latent, simulated-moment calibration — put the autocorrelation ladder inside its bands on every seed, the first generative candidate to do so, and its calibrated shares independently reproduced the variance decomposition estimated three candidates earlier; but Gaussian innovations churned earnings quintiles far too fast

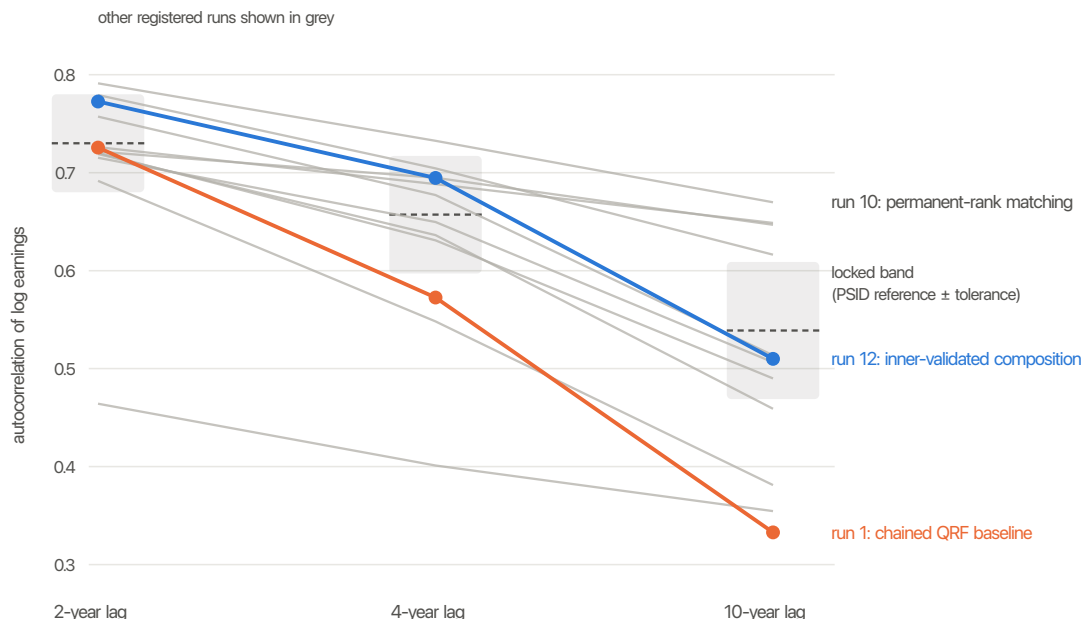


Figure 1: The autocorrelation ladder across the twelve distinct registered model runs. Shaded bands are the locked tolerances around the committed PSID reference at each horizon; run 13 re-registered the run-12 specification and reproduced it bit-exactly, so its curve coincides with run 12’s. Every value is computed from the committed run artifacts at figure-build time (`scripts/build_paper_figures.py`).

and the classifier read the synthetic joint easily. Replacing the Gaussian law with an empirical rank-transition kernel fixed the mobility matrix and halved the classifier gap while its one-step memory collapsed the ladder’s tail; adding two-step-plus-anchor memory through nearest-neighbor conditional draws brought the strongest scorecard of the sequence — the pairs-view classifier under its threshold on all five seeds, two seeds clearing the entire gate — with the window-3 classifier and the ten-year rung each missing on three seeds by thousandths. A final variant that matched donors on their estimated permanent rank rather than their observed anchor converted anchor noise into permanent signal and overshot every rung: the tenth registration, and the second time the program measured the same lesson from the opposite side.

7.3 How existing models solve the joint problem

Every serious earnings-history model confronts the same triple — the cross-sectional marginal, person-level persistence, and measurement noise — and the field has four working answers, each buying one thing at a stated price.

Reuse observed careers. MINT splices segments of donor workers’ administrative earnings records onto targets matched on demographics and recent earnings, adding a person-specific fixed effect in its regression projections (Social Security Administration 2024). Reuse dissolves the joint problem rather than solving it: the marginal is exact because the values are real, and person-level dependence is exact because whole careers come from one person. The price is donor support — no pattern appears that no donor exhibited, cohort drift needs ad hoc adjustment, and conditioning is only as rich as the match cells.

Generate parametrically, then align. DYNASIM carries an individual-specific error term in its

registered run	geometry seeds passed (of 5)	battery seeds passed (of 5)	pooled Q0	gate 1
1. chained weighted QRF (baseline)	○ ○ ○ ○ ○	○ ○ ○ ○ ○	not scored	× fail
2. + person effect as feature	○ ● ● ● ○	○ ○ ○ ○ ○	not scored	× fail
3. + persistence-aware decomposition	● ● ● ○ ○	○ ○ ○ ○ ○	not scored	× fail
4. structural three-component assembly	○ ○ ○ ○ ○	○ ○ ○ ○ ○	not scored	× fail
5. donor splicing (single-donor)	○ ○ ○ ○ ○	○ ○ ○ ○ ○	not scored	× fail
6. donor splicing (segments)	○ ○ ○ ○ ○	● ● ● ● ●	not scored	× fail
7. Gaussian-copula rank dynamics	○ ○ ○ ○ ○	○ ○ ○ ○ ○	not scored	× fail
8. empirical rank kernel	○ ○ ○ ○ ○	○ ○ ○ ○ ○	not scored	× fail
9. k-NN rank bootstrap	○ ● ○ ● ○	● ● ○ ○ ○	not scored	× fail
10. permanent-rank matching	○ ○ ○ ○ ○	○ ○ ○ ○ ○	not scored	× fail
11. calibrated blend + Q0 regime	○ ○ ○ ○ ○	○ ○ ● ● ○	×	× fail
12. inner-validated composition	○ ○ ● ● ●	○ ● ● ● ●	✓	× fail
13. inner-validated composition (re-registered)	● ● ● ● ●	○ ● ● ● ●	✓	✓ pass

● seed passed ○ seed failed
 gate 1 requires $\geq 4/5$ seeds on geometry AND $\geq 4/5$ on battery;
 runs 11–13 were also scored on the amended benefit-space block and pooled Q0; run 13's pairs-view classifier is gated by the amendment-2 rule (mean over 20 pre-registered seeds ≤ 0.53 with a per-seed cap of 0.554) rather than per-seed at 0.53.

Figure 2: Per-seed gate conjunction for the thirteen registered runs, from the committed run artifacts. Each cell shows the five locked holdout seeds; gate 1 requires at least four of five on geometry and on battery jointly. Runs 11–13 were additionally scored on the amended benefit-space block and the pooled zero-anchor gate, and run 13's pairs-view classifier is gated by the ratified amendment-2 rule — the mean over twenty pre-registered seeds at the unchanged 0.53 line with a 0.554 per-seed cap — rather than per-seed.

earnings equations (Favreault et al. 2015; Urban Institute 2024); CBOLT decomposes each worker’s earnings into a permanent shock — the long-run gap from the group mean — and a transitory shock, estimated on administrative panel records (Congressional Budget Office 2018). Both then align: simulated outputs are recalibrated, typically rank-preservingly, to external distributional and aggregate targets (Li and O’Donoghue 2013). Alignment is the field’s standing admission that generated marginals drift; it re-imposes them after the fact and guarantees the published tables. The documented cost is that alignment can distort the relationships between variables and hits its targets whether or not the underlying process is right (Li and O’Donoghue 2013) — the class of silent correction this project’s gate exists to expose rather than absorb: alignment would have masked the fourth candidate’s marginal failure entirely.

Fit the simulator, not the one-step step. The administrative-data literature estimates rich nonlinear processes by simulated method of moments — parameters chosen so that simulated trajectories reproduce the target moments jointly (Güvenen et al. 2021). This answers the composition problem directly: one-step conditionals estimated from noisy data need not compose into correct long-horizon dynamics, so the generator is scored as a whole, at the horizon that matters. The price is a parametric process and total dependence on the chosen moment set.

Put the dynamics in rank space. Nonlinear panel frameworks place persistence on a latent quantile position and read earnings off the empirical quantile function (Arellano et al. 2017). The marginal becomes untouchable by construction and all modeling effort concentrates on the rank process, where the permanent–transitory structure lives — at the price of heavier estimation machinery than any off-the-shelf learner supplies.

The fifth run tested the first strategy as a benchmark, in its simplest registered form: one donor per target, whole careers spliced by age with a level adjustment at the anchor. It failed — informatively. The marginal came through exactly as reuse promises (every distributional and tail metric passed comfortably), but the omnibus classifier still separated spliced from held-out trajectories on every seed, and the battery found whole-career reuse too persistent: mobility, exit rates, and zero-spell lengths all sit outside their bands, and the autocorrelation ladder overshoots at every horizon. The run’s own diagnostics name the artifacts — donor and target cohorts occupy offset points of the biennial age grid, so half the spliced values come from an adjacent age; scaling a whole career to a noisy anchor observation converts that noise into a permanent component; and one donor per career is strictly more persistent than the segment splicing MINT actually performs. The result bounds naive splicing rather than refuting the strategy. The registered segment variant — three-period segments from multiple donors, spliced by calendar period and level-adjusted at each segment boundary — then became the first candidate to clear the entire battery: all five seeds pass every dynamics tolerance, with the autocorrelation ladder, mobility, exit rates, and spell lengths inside their bands at once. What remains is a single metric family: the omnibus classifier still separates spliced from held-out trajectories by 0.017 and 0.027 above its two thresholds, with every other geometry metric passing on every seed. Six runs in, the gate has narrowed from “the dynamics are wrong” to “a classifier can still tell,” which is the precise question the generative track now has to answer. The generative track’s next design composes the third and fourth strategies with the project’s existing machinery: empirical quantile marginals, persistent dynamics in rank space, and innovation dispersion calibrated so iterated — not one-step — dynamics match the data’s ladder (Arellano et al. 2017; Güvenen et al. 2021).

Two auxiliary results harden the target. Excluding every PSID observation carrying an income-assignment flag (8.7% of positive-earnings person-periods overall, rising to 14–18% in the 2020 and 2022 waves) moves the 10-year reference by +0.012 — survey imputation does not explain the gap the candidates must close. And the statutory-resolution surface of Section 5 produced its first artifact: the project’s Python benefit oracle and the Axiom rules engine compute identical primary insurance amounts — 240 of 240 synthetic careers exact to the cent, both code paths pinned by revision — after the cross-engine comparison surfaced and fixed a statutory off-by-one in the elapsed-year count and a

rounding-direction error in a test fixture that no single-engine test had caught.

7.4 Recalibrating the gate toward decision relevance

Nine runs in, the maintainer asked the governance question the protocol exists to make askable: is the bar too high? The program answered with analyses rather than judgment. Classifier forensics showed the two best candidates' identical residuals were distinct defects, not a shared wall. A downstream-relevance analysis then pushed generated and real careers through the project's own statutory benefit calculator: the best candidate's benefit distributions were statistically indistinguishable from real ones — the distributional distance inside the real-versus-real noise floor, central gaps under two percent against a five-percent criterion — while the same analysis exposed a defect no locked metric had isolated: careers generated for people observed with zero earnings at their anchor overstated their benefits by nine percent, exactly the population a progressive benefit formula weights most. The locked gate was strict on an axis that does not matter for benefits and silent on one that does.

The response was the contract's own amendment mechanism, exercised in full for the first time: a proposal committed as an inert object changing nothing; a fresh adversarial referee round (which caught a misattributed criterion citation and forced disclosure that the proposed demotion flips four historical geometry verdicts — none of which changes an overall outcome); fixes; a verification pass; maintainer ratification by merge; and a follow-up flipping the ratified content into the locked block. The amended gate demotes the benefit-immaterial window-3 classifier to reported status and adds a gated benefit-space block — distributional and decile bands at the five-percent criterion, a distance bound derived from a committed real-versus-real anchor, and a pooled band on the zero-anchor subgroup where reality measures under three percent and the best candidate nine. Recalibration added strictness where the evidence says it matters: the best candidate still fails the amended gate, and reality still passes it.

7.5 Nested validation and the forecast discipline

Two practices matured in the runs that followed the amendment. The eleventh candidate composed the two best-understood mechanisms and failed through two couplings its registration had not foreseen — including a swing of the zero-anchor benefit gap from +9 to −18 percent through an interaction between its memory coordinate and the very subgroup it was fixing. The response was methodological rather than another guess: the one-shot rule protects the outer holdout, so model development may iterate freely on inner splits carved from each seed's training complement. An inner-validation harness now mirrors the amended gate at inner scale, and a design sweep raced the candidate mechanisms head to head on it — establishing with no outer-holdout contact that the zero-anchor participation refit alone closes the benefit gap, that the drawn persistent-state coordinate fixes the dynamics battery completely while wrecking the joint completely, and that one composition stood near-clearing everything at once. The twelfth candidate froze that composition, every constant selected by nested validation.

Each of the last two runs also carried a pre-registered forecast — component probabilities logged on the public registry before the run, graded against the outcome after. The twelfth run graded almost exactly: forecast 0.42 with modal failure named as the pairs-view classifier one seed short, and that is what happened — the battery passed (the ten-year rung in-band on every seed for the first time), the zero-anchor gap closed to +0.04 percent, the per-seed benefit metrics passed everywhere, and two seeds clipped the pairs-view classifier by 0.0015 and 0.0030. Twelve registered runs in, the program's entire remaining distance to its own gate is a classifier residual of thousandths on one view.

A twenty-seed extension of that measurement — reported, not gated — placed the residual against the matched real-vs-real floor (Figure 3). The twenty-seed mean sits 0.0066 below the locked 0.53 line; the five locked gate seeds alone average 0.0021 below it, and the two seeds that clip the line are the two largest candidate scores across all twenty. Whether that pattern is seed noise or a seed-set property

became the second amendment proposal, and the adversarial referee round earned its keep: the referee found the locked seeds non-exchangeable with the fresh ones in exactly the damaging direction (a random five-subset has a mean that high with probability 0.009), found the proposal's headline margin describing the twenty-seed mean where the gate scored the locked-five mean — one standard error below the line, not five — and named the self-rescue: unlike the first amendment, whose trigger still failed after ratification, this one's triggering candidate would have flipped to a pass.

7.6 The first pass

The reworked amendment survived a verification round and was ratified as an estimator change that refuses its own trigger. The pairs-view classifier now gates on the mean over twenty pre-registered seeds at the unchanged 0.53 line — an operating characteristic *stricter* above the line than the per-seed rule it replaces — with a per-seed catastrophe cap re-derived from a classifier-version-matched floor, and two standing rules: no candidate's committed verdict changes under a rule proposed after its own run, and floor derivation and candidate scoring must share a classifier version. Run 12's verdict stands as a fail permanently.

The thirteenth registered run was therefore a fresh registration of the identical specification, with the protocol's plainest disclosure yet: its pairs-classifier outcome was already public before the run, so the registered forecast was 0.97 with the residual entirely on execution error. The run reproduced every committed baseline bit-exactly — all twenty pairs scores, the locked-seed battery, benefit, and geometry blocks, to the last digit — and passed every block of the amended gate: geometry five of five, battery four of five, pooled zero-anchor gap +0.04 percent, twenty-seed classifier mean 0.5234 against 0.53 with a maximum seed at 0.5330 against the 0.554 cap (Figure 2). The first pass of the program arrived, in other words, not from a better model but from a better-measured gate — and the record distinguishes those two things explicitly, which is the point of keeping one.

A registered reform-delta diagnostic followed the pass (reported, not gated): two opposite-incidence mechanical reforms — the first PIA factor raised from 90 to 95 percent, and the taxable maximum removed from the AIME step — computed on real versus generated histories under the locked holdout protocol, each gap measured against a real-vs-real half-split floor. The aggregates a reform score leads with land inside the floor on both reforms: mean benefit change (gap \$0.29 against a \$0.48 floor bottom-loaded; \$4.02 against \$7.77 top-loaded), winners share, and the zero-anchor subgroup. The fine incidence curve is not fully reproduced: the bottom-loaded reform misses only the concave first-bend decile (1.9 times its floor, about one percent of the flat \$51 monthly delta), while cap removal shifts roughly a quarter of the top compared decile's gains downward — real +\$75.76 a month at the ninth decile against generated +\$55.33 — with six of seven compared deciles individually outside their floors. The named mechanism is that the generator under-concentrates cap-riding careers, diluting the extreme top of the earnings distribution, consistent with the run-12 microtexture forensics; it is the program's next target, found by the protocol's own diagnostics rather than by a downstream user.

7.7 What failure buys

Forty-five scored runs across six gates — thirteen at gate 1, sixteen at gate 2, nine at gate 2b, two at gate 2c, one at the disability gate, four at the transport gate — thirty-nine published failures against six passes under a rule that refused to rescue its own trigger, seven gates locked through their own adversarial rounds with seven ratified amendments among them — two at gate 1, at gate 2 an estimator aligned to its operating characteristic and then a tranche structure made explicit, and at the transport gate a mis-anchored family demoted on committed evidence after the lock and the gated surface then twice more narrowed on registered forensics, sixty-five cells at first lock and forty-four at the pass, every removal machine-reasoned and forensics-proven, no committed verdict changed — the marital-transition tranche passing at its sixteenth candidate, the household-composition tranche at its ninth, the marriage-by-earnings tranche at its second — the shortest of the family ladders — the

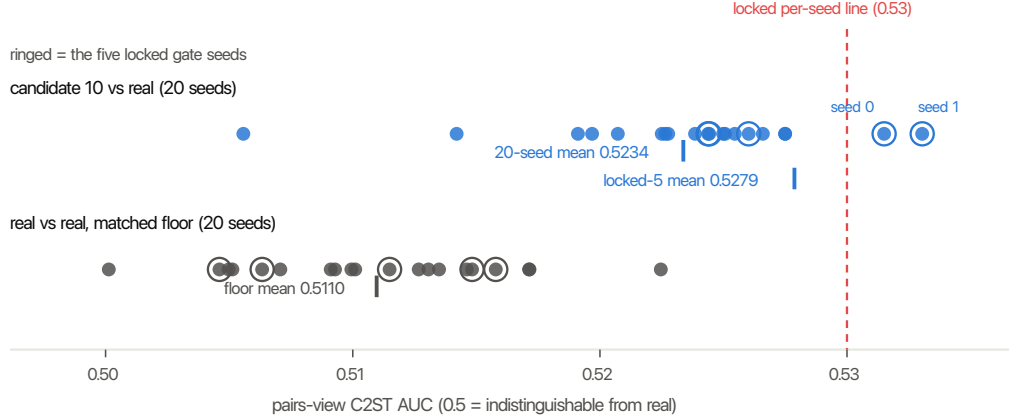


Figure 3: The twenty-seed extension of the run-12 pairs-view classifier measurement (reported, not gated), from the committed diagnostics artifact. Blue: candidate 10 scored against real holdouts; grey: the matched real-vs-real floor. Ringed points are the five locked gate seeds; both the twenty-seed and locked-five candidate means are marked.

disability gate at its first, the representative-frame transport gate at its fourth, verified bit-exactly on the forty-four-cell surface — the funded crux resolved — and the temporal-holdout projection-drift gate locked as the seventh, its candidate twice a designed stop before scoring rather than an improvised run — first on the missing harness, graded a registration error, then on the missing 2014-or-earlier external bindings — neither a pass nor a failure, the run lane’s discipline — its scored run still ahead — twelve registered forensics rounds and a reform-delta diagnostic feeding the next registration, a cross-anchor cost-ordering synthesis and a same-frame revenue pseudo-projection that between them isolated the compression the representative-frame transport was built to remove, a caregiver-credit benefit seam run end to end through the production tax calculator, six external-anchor replications reported at the same standing as any gated run, and fifty-six forecasts registered and fifty-six graded — the fiftieth, the transport gate’s second candidate, has since run, and the ladder closed at the fourth-candidate pass — is the protocol working as designed. Each failure published with the same prominence as a pass would, each narrowed the design space with a finding that transfers beyond this project, and none required trusting the authors: the registrations predate the runs, the artifacts recompute, and the thresholds cannot drift to accommodate a result — when a gate itself was recalibrated, it moved only through a public, refereed, ratified amendment whose every verdict change is disclosed in the contract’s own history. The contrast with validation by fidelity — where a model meets the tables it was fit to and the first hard test arrives after adoption — is the point.

8 The demographic layer and the replication anchors

The tally above compresses several developments the same protocol governed once gate 1 produced its first pass; this section details the first two and the sections that follow the rest: a second stage gate, opened and locked on the family-transition layer, which its sixteenth registered candidate passed on the marital-transition and fertility tranche, and the reform-scoring surface, tested against external anchors on real microdata. Neither program is closed — the gate’s marriage-by-earnings tranche and forward projection remain — and both are on the public record with the discipline the earnings-history gate established.

8.1 The second gate’s lock ceremony

Gate 2 — family and benefit outputs — governs the demographic transitions survivor, spousal, and caregiver reforms score on: first-marriage, divorce, widowhood, and remarriage hazards, cohort nuptiality, fertility, and the dissolved-state stock shares eligibility rides on. It locked on 2026-07-08, through the ceremony gate 1 established: draft floors, an adversarial round, fixes, verification, and ratification by merge.

The draft built the reference moments and a person-disjoint half-split noise floor on five split seeds, with each per-cell tolerance the floor mean plus four standard deviations across forty gate-eligible cells. The adversarial round returned *amend before lock*, and its objection was quantitative: on five seeds the floor’s standard deviation is a five-draw estimate of a half-normal, so the committed tolerances realized anywhere from 0.9 to 5.6 times their own measured noise, and a faithful candidate — one drawing from the reference process itself — cleared the four-of-five gate with probability 0.023. The referee also showed that the draft’s floor scale, pass statistic, and seed rule described three different experiments, and that a verbatim copy of a training half passed at the noise floor, which no moment gate can prevent — the memorization defense is procedural (registration, holdout exclusion, and the no-self-rescue rule), not a property of the cell set.

The fixes rebuilt the floor on one hundred split seeds, stabilizing the estimator to roughly 3.2 times each cell’s own noise; added a power cap that gates a cell only when its stabilized tolerance is at most $\ln(1.5)$ — a 1.5-times rate error — demoting under-powered per-age cells to report-only and recovering their coverage through pre-registered aggregates; and added the sequence and stock statistics a banded-marginal gate misses: origin-split remarriage, cohort ever-married-by-40, and dissolved-state stock shares by age and sex. Under the same rule on the rebuilt floor, the faithful-candidate operating characteristic is 0.969 over the resulting forty-six gated cells. Verification confirmed the amendments, and maintainer ratification by merge was the lock event. Every figure recomputes from the committed floor artifact (`runs/gate2_floors_v2.json`).

The gate’s external anchors are shape reports, not level gates, and the ceremony made the reason explicit. The raw recent PSID marriage and divorce rates run about 2.4 times the national vital-statistics rates — but the PSID counts persons transitioning where the vital-statistics series counts couples (a factor of two exactly) against a person-year denominator for the age-15-plus population (a further 1.22). Attributing those two concept factors leaves residuals of 1.036 for marriage and 0.989 for divorce: the anchor is a near-bullseye once the population concepts are aligned, and the period-matched fertility series sits at a median PSID/NCHS ratio of 0.93.

8.2 The gate-2 candidate ladder

Sixteen candidates have run against the locked family-transition gate, each registered on the campaign registry ([issue #42](#)) before its single scored run and graded against a pre-registered forecast after. The ladder narrows the failure from fifteen distinct cells to a single pass, and the record of what each step cost is the evidence the gate exists to produce. The base composition — stratified empirical hazards with a mortality-composed widowhood component — fails all five seeds across fifteen distinct cells. Adding an age-by-sex first-marriage interaction and a decade-period widowhood mortality makes the composed widowhood backfire, widening the failure to nineteen cells. Pooled-rate shrinkage un-explodes that cascade back to eight; a parametric per-sex mortality trend then fixes the 65–74 female widowed stock but lets the residual go diffuse. Replacing the trend’s source — PSID’s own male mortality slope was unstable — with an external NCHS life-table trend, and adding single-year-kernel fertility and an origin-split remarriage table, narrows the failure to six cells. The sharpest fix is the sixth: source-aligning the spouse-death level to the surviving spouse’s own marriage-history widowhood incidence removes a sex-asymmetric wedge in the earlier death-record level — which understated female widowhood by 1.9 to 2.8 times while overstating male — lifting the 75-plus female widow stock from one seed to four and clearing one full seed. Its residuals are the untargeted male lifetime-marriage

sequence cell, two seeds over its 0.047 tolerance; a three-cell miss on a third seed; and the fertility clip inherited from candidate 5 on a fourth — the first candidate to clear a seed.

Candidates 7 through 9, still under the gate’s original single-draw estimator, worked the two chronic marriage-count cells. Order-conditioned remarriage (candidate 7) fired the registered change-what-works risk — seed 0 regressed from forty-six passing cells to forty-four — and fixed the sign, since the generator under-produces lifetime marriages; an aggregate-preserving order split (candidate 8) isolated the deficit as compositional rather than a remarriage-level error, reconciled to a zero residual; and observed undatable-marriage initial states (candidate 9) cured the male count while the observed residual overshoot the female one. By candidate 9 every remaining failing cell’s twenty-draw-mean tilt measured sub-tolerance while single draws still decided verdicts: the binding constraint had become the gate’s estimator, not the model, and the ladder paused for a refereed amendment.

The seven candidates that followed ran under the amended mean-over-draws estimator against one named level target — the elderly-widow stock — that four forensics rounds resolved in turn. Candidate 10 confirmed the stock cell as an unambiguous level under-production rather than absorbed draw noise; candidate 11’s elderly-remarriage split exposed two compensating errors in the pooled band; candidate 12 cleared the cell on all five seeds with observed already-widowed initial states, while falsifying a co-registered spousal-gap-draw delta on the run’s own inertness test; candidates 13 and 14 re-banded the young and the oldest surviving-spouse widowhood, trading the stock against aging-in; and candidate 15 removed the deployment-time NCHS mortality trend the gate’s untrended reference does not carry, reaching three of five seeds — the ladder’s best short of a pass. Candidate 16 conditioned the widowhood hazard on one observed covariate and passed.

Table 2: The gate-2 candidate ladder, from the committed run artifacts (`runs/gate2_hazard_v1.json` through `v16.json`). “Distinct failing cells” counts gated cells missing on at least one of the five locked seeds; the gate requires at least four of five seeds with every one of the forty-six gated cells inside its locked tolerance, which candidate 16 is the first to meet. Candidates 10–16 are scored under the amendment-1 mean-over-draws estimator; candidates 1–9’s committed verdicts stand.

candidate	distinct failing cells	seeds passing	headline lesson
1 — stratified hazards, composed widowhood	15	0 / 5	the base composition; young first-marriage and female widowed-stock cells drift
2 — + age×sex first marriage, period widowhood mortality	19	0 / 5	the composed widowhood backfires — the widest failure
3 — + pooled-rate mortality shrinkage	8	0 / 5	shrinkage un-explodes the cascade; the widowed stock persists
4 — + parametric per-sex mortality trend	8	0 / 5	fixes the 65–74 female widowed stock; the residual goes diffuse
5 — + external NCHS mortality trend, kernel fertility	6	0 / 5	an external trend replaces PSID’s unstable male slope

candidate	distinct failing cells	seeds passing	headline lesson
6 — + source-aligned spouse-death level	5	1 / 5	removes the sex-asymmetric widowhood wedge (1.9–2.8×); first seed cleared
7 — + marriage-order remarriage, marital-status fertility	7	0 / 5	the change-what-works risk fires — seed 0 regresses 46→44; the generator under-produces lifetime marriages
8 — + order-split remarriage, aggregate counts preserved	6	0 / 5	the marriage-count deficit is compositional, not remarriage-level (reconciled to zero residual)
9 — + observed undatable-marriage initial state, low-parity fertility	7	0 / 5	the count fix cures males but overshoots females; sub-tolerance tilts expose the single-draw estimator as binding
10 — first run under the amended estimator; + age-band remarriage	4	1 / 5	under the mean estimator the 75+ widow stock is a level miss, not draw noise
11 — + 50–64 / 65–74 / 75+ remarriage split	4	1 / 5	the pooled elderly band hid two compensating errors; splitting trades outflow for aging-in observed
12 — + entry-widowed initial states, age-conditioned gap draws	3	2 / 5	already-widowed states clear the 75+ stock 5/5; the gap-draw delta is proven inert
13 — + young surviving-spouse widowhood bands (18–34, 35–44)	2	2 / 5	removes an order-of-magnitude young-widowhood rate error; both counts clear, but less aging-in lowers the stock
14 — + split the 75+ widowhood band (75–84, 85+)	2	2 / 5	re-banding recovers incidence (0.93→0.95) but a fixed aggregate cannot lift the stock

candidate	distinct failing cells	seeds passing	headline lesson
15 — — NCHS mortality trend inside the gate	2	3 / 5	the gate’s reference is the untrended panel; removal reaches three seeds but leaves a survival-to-75 stock leak
16 — + widowhood support-composition stratum	1	4 / 5	PASS — conditioning on the observed support window closes the yield leak; the sole miss is an RNG-isolated fertility split artifact

8.3 The estimator amendment

Candidate 9’s grading named a constraint the model could not move. Every failing cell’s mean over twenty pre-registered simulation draws sat inside its tolerance, yet single draws — one frozen replicate per seed — were deciding the verdicts. The gate had ratified its operating characteristic on a draw-noise-free basis: a faithful candidate — one drawing from the reference process itself — modeled per cell as a half-normal with the floor’s own standard deviation and no simulation-draw term, which gives a per-seed pass probability of 0.9404 and a four-of-five gate pass of 0.9685. But the single-draw estimator the gate shipped with injected a per-cell draw-noise term the tolerances never budgeted for, dropping that same faithful candidate to roughly 0.885 and 0.896 on the six measured cells: the pass probability the gate was ratified to deliver was not merely hard but unachievable under its own estimator.

The fix — proposed as an inert object, carried through an adversarial round and verification, and flipped live in a follow-up merge, the ceremony gate 1’s second amendment established — scores each cell on the mean cell rate over twenty pre-registered draws rather than one, leaving the tolerances and the forty-six-cell four-of-five conjunction untouched. Averaging over twenty draws shrinks the injected noise toward the draw-noise-free rate the tolerance was measured against, restoring the faithful candidate to roughly 0.939 and 0.967 — back to the numbers the gate was locked with, now achievable because the estimator finally shares their derivation basis. It is an estimator aligned to its own operating characteristic, not a loosened error budget: the four noise-dominated cells collapse toward zero clip probability while the one systematically mistuned marriage-count cell fails harder under the mean, not softer, and a candidate with gross level errors still fails by orders of magnitude at any number of draws. The amendment is prospective only — candidates 1 through 9 stand as committed failures under the no-self-rescue rule, and candidate 10 was registered as the first fresh run under it. The operating characteristic recomputes from the committed floor artifact (`runs/gate2_floors_v2.json`), and the record keeps the better-measured gate and the better model distinct exactly as gate 1 did.

8.4 Four forensics rounds

Between the scored runs the ladder registered four diagnostic rounds on the campaign registry, each reported, never gated, published whatever it found, and cited by the candidate it licensed. They are the method’s throughline: candidates 9, 11, 12, and 16 — the pass among them — each registered only after, and citing, the round that preceded it.

The first round decomposed the two chronic marriage-count cells on the training halves and found the male deficit was not a rate error but a measurement-concept residual — the reference carries

lifetime-marriage counts from episodes with undatable start or dissolution years that no hazard model can generate, so the fix is an observed initial state, not a dial. Its stability sub-question, twenty redraws of the same specification, measured every remaining clip’s mean tilt sub-tolerance, the evidence that motivated the estimator amendment. The second round decomposed the elderly-widow stock gap and the female count residual and found one mechanism behind both: a pooled fifty-plus current-age remarriage band applying a rate roughly nine times too high to 75-plus widows, depleting the stock from the outflow side while carrying the female count over-production. The third round audited how the reference constructs the 75-plus widowed stock and found that 12.1 percent of it is carried from spouse deaths predating the person’s panel support — structurally unreachable by any transition rate, and entirely fixable by an observed initial state — while the simulated young-widowed pool ran to 3.16 times reference at ages 15 to 49, fed by the youngest surviving-spouse band edge. The fourth round split its two remaining cells: seed 2’s fertility clip decomposed as a split artifact — a systematic deficit comfortably inside tolerance, wrapped in a maximum-of-five reference draw and a minimum-of-five simulation draw, isolated from the widowhood delta’s random stream and so failing regardless, which fixed the pass path at seeds 0, 1, 3, and 4 — and the elderly-stock leak decomposed as a survival-to-75 yield gap: real widowhoods correlate with long observed support, which a uniform within-band hazard cannot see, so simulated 50-to-64-onset widowhoods reached age-75 windows at 39 percent against the reference’s 57 percent. Candidate 16 conditioned on exactly that window.

8.5 The first gate-2 pass

Candidate 16 added one covariate to the surviving-spouse widowhood hazard: whether a person’s observed support window reaches age 75, a binary stratum known in advance from the panel attributes alone — the same observed-data class as the initial-state fixes — true of 3,147 of 41,409 persons, a 12.7 percent exposure-weighted share. Within each age band and sex the two strata are train-estimated and recombine to candidate 15’s band aggregate by an exposure-weighted identity, so aggregate widowhood incidence is preserved (recombination residual $1.7\text{e-}18$, reconciled to zero; 75-plus incidence moves from 1.060 to 1.061 of reference) while the event composition matches the reference’s window correlation, and all four gated widowhood-incidence cells hold on every seed.

The delta closed the yield leak the fourth forensics round had sized. The 50-to-64-onset survival-to-75 yield rose from 0.581 to 0.816 of reference as the share of those widowhoods whose window reaches 75 moved from 0.391 to 0.551 toward the reference’s 0.572, lifting the 75-plus female widowed stock from 0.841 to 0.914 of reference. The chronic stock cell cleared all five seeds — scores 0.057 to 0.111 against a 0.185 tolerance, with the two seeds candidate 15 had failed both flipping — and the marriage counts the recomposed exposure put at risk held on both sexes, minimum margins plus 0.008 and plus 0.011. The gate passed four of five.

The scope is on the record with the pass. It is a model delta under a locked gate — thresholds, protocol, and the ratified estimator amendment all fixed before the run was registered — not an amendment rescuing its own trigger; candidates 1 through 15 stand as committed failures. It certifies the marital-transition and fertility tranche the locked cells score, and no more: the household-composition tranche and the marriage-by-earnings joint the gate’s own scope note names remain unscored, and the sole failing cell anywhere in the passing run is seed 2’s completed-fertility clip — the split artifact the fourth forensics round diagnosed, byte-identical to candidate 15 and isolated by construction from the delta that carried the gate.

8.6 The replication anchors

While the generative track worked the demographic gate, the reform-scoring surface — the statutory AIME/PIA chain and the survivor, spousal, and caregiver benefit plumbing, run on real PSID careers — was tested against six external DYNASIM anchors. Each is reported, not gated: registered on the campaign registry before the run and published regardless of outcome. Four score real careers

directly; the price-indexing anchor additionally routes generated careers through the same calculator. A sixth, added once the gate-2a pass validated the marital histories it needs, recomputes Mermin’s exact shared-earnings quintile concept on real couples and closes the own-versus-shared concept delta the price-indexing anchor had carried (Section 9).

Table 3: The six external-anchor replications, reported not gated, from the committed diagnostic artifacts; each registered on [issue #42](#) before its run. The shared-earnings row (Section 9) was run after the gate-2a pass validated the marital histories it requires.

provision (anchor)	achieved replication	artifact / registration
progressive price indexing — Mermin (2005)	the price-indexing scalar lands at 66.73% of scheduled against DYNASIM’s 67.8%; the progressive-price-indexing quintile pattern reproduces — real and generated careers decline together from 100% to 84% of scheduled across quintiles, matching DYNASIM’s monotone gradient — with real-versus-generated gaps inside the floor at four of five quintiles	<code>replication_ppi_mermin_v1.json</code> ; #42 c4907444903
normal-retirement-age increase to 70 — Mermin (2005)	79.75% of scheduled overall, cross-quintile spread 0.045pp, each quintile within 0.03–0.33pp of Mermin’s 79.4–79.9% row	<code>replication_mermin_rows_v1.json</code> ; #42 c4911609804
COLA reduced 0.4pp — Mermin (2005)	99.2% of scheduled at ages 62–67 (anchor 98.9%) and 93.0% at 80–85 (anchor 92.4%)	<code>replication_mermin_rows_v1.json</code> ; #42 c4911609804
earnings sharing — Favreault and Steuerle (2007)	the best-populated cell — married women, the large gainers — lands on DYNASIM: gain 20% 49 vs 44, gain 5% 63 vs 60, lose 5% 24 vs 22; all four registered directional calls hold	<code>replication_r7_sharing_v1.json</code> ; #42 c4911171806
caregiver credit — Smith et al. (2020)	all four candidate plans place 58–71% of aggregate gains in the bottom lifetime-earnings quintile, in or above the anchor’s 52–62% band; the concentration and reach rankings match (Spearman 0.8)	<code>replication_caregiver_v1.json</code> ; #42 c4911453454

provision (anchor)	achieved replication	artifact / registration
shared-earnings progressive price indexing — Mermin (2005)	Mermin’s exact shared-lifetime-earnings quintile concept on 3,131 real couples (own-record benefits, shared ranking): shared Q1 lands at 99.49% of scheduled against the anchor’s 98.7 where the own-record analogue read 100.00, Q1–Q3 each move closer to the anchor, the distribution stays monotone (99.49 → 86.62), and the price-indexing wedge is flat at 66.73% against 67.8; 45.8% of the persons change quintile between the two rankings; all four registered calls hold	<code>replication_ppi_shared_v1.json</code> ; #42 c4931009783

The discipline that makes those comparisons legible is that every level difference from the DYNASIM projections is named, not absorbed. The populations differ — observed PSID retirees eligible 2005–2019, born 1943–57, against DYNASIM’s projected 1960–80 cohorts evaluated in 2049 — as do the windows (a biennial PIA-proxy convention, the national average wage projected forward to each indexing year) and the earnings concept (the earnings-sharing exercise scores real couples with both spouses computable, a shared-versus-individual distinction the incidence turns on). Where the older PSID cohorts hold more single-earner couples, married men lose harder than DYNASIM’s 2049 projection, and the direction, not the level, is what the test certifies. The same discipline caught an error in an anchor: Favreault and Steuerle’s Table 3 carries a package-1c married-women column summing to 109.1, a 9.1 in a “no change” cell where every comparable earnings-sharing cell reads 0.0 — a printed typo, flagged and carried verbatim, feeding no scored result.

Two things are true at once, and the record keeps them distinct. Distributional scoring on validated real microdata is anchor-grade across the cataloged provision classes — price indexing, retirement-age and cost-of-living reductions, earnings sharing, and caregiver credits — each reproducing its DYNASIM incidence pattern within the noise floor or a named delta. The generated-population demographic track has now passed every tranche of its own gate — the marital-transition and fertility surface at the sixteenth candidate, household composition at the ninth, the marriage-by-earnings joint at the second — leaving no family tranche unscored; the open frontier is the representative-frame transport and forward projection (Section 10), not a demographic tranche. The paper reports all of it because the protocol scores all of it, and a reader who wants to know which claims are load-bearing today can read it off the record rather than the prose.

9 The reform-scoring surface, tested against itself

The gate-2a pass validated the marital histories the anchor replications need, and three registered diagnostics followed — each reported, not gated, each graded against a forecast filed before the run. One closed the concept delta the price-indexing anchor had carried; the other two scored the committed encodings against each other and added the revenue side the projection roadmap requires, and both surfaced the same distortion a representative-frame transport exists to remove.

9.1 The shared-earnings concept, closed

Section 8’s progressive-price-indexing anchor reproduced Mermin’s quintile gradient in shape but not in grouping: Mermin (2005) ranks retired workers by *shared* lifetime earnings — own earnings when single, half the couple’s when married — and the Phase-A run could only rank by own record. With the gate-2a marital histories in hand the exact concept becomes computable on real couples, and the sixth anchor in Table 3 runs it, holding the benefit reform at the Phase-A encoding verbatim and changing only the ranking variable. On the 3,131 couples with both spouses computable, 45.8 percent of persons change quintile between the own and shared rankings, and the regrouping moves the bottom onto the anchor: shared Q1 lands at 99.49 percent of scheduled against Mermin’s 98.7, where the own-record analogue read a flat 100.00, and Q1 through Q3 each sit closer to the anchor than the own-record version. The distribution stays monotone and the price-indexing wedge is flat at 66.73 percent against Mermin’s flat 67.8. The one carried delta is stated rather than chased: the upper quintiles move slightly *away* from the anchor under shared ranking — the arithmetic complement of pulling the low-earning spouses of high earners up out of the top groups — and the top stays compressed against Mermin’s 71.7 by the truncated observation window. All four registered calls held; the own-versus-shared concept delta the earlier anchor named is closed on real data, and the generated-couples version of the same concept waits on the marriage-by-earnings tranche.

9.2 Cost ordering and a revenue projection: two compression fingerprints

The anchor replications each scored one provision family in isolation. A cost-ordering synthesis then scored every committed encoding once on a single common frame — 1,549 sex-resolvable careers under the Phase-A 2050 transport — and tested the aggregate cost deltas *ordinally* against the anchors’ published cost columns, since observed completed careers are not a Trustees projection and the levels differ by construction. The signs all agreed, eight of eight — four Mermin provisions score as savings, four caregiver plans as costs — and the caregiver ordering matched the anchor at rank correlation 0.913, its one inversion a pair the anchor itself ties. But the Mermin quartet’s cost ordering broke on a single adjacent swap: the frame ranks the retirement-age increase above progressive price indexing, where DYNASIM ranks them the other way. The mechanism is the compressed support the earlier anchors already named — on these observed careers most indexed earnings sit below the thirtieth-percentile bend that progressive price indexing protects, so it bites lightly (−11.0 percent) while the uniform retirement-age reduction cuts across the board (−20.2 percent). The forecast of a perfect ordering failed, and the miss is a diagnosis: progressive price indexing is exactly the provision whose relative magnitude depends on getting the earnings distribution right.

The same distortion reappeared on the revenue side. The first roadmap milestone added a taxable-payroll aggregation, a present-value balance analogue, and an endogenous trust-fund-exhaustion ledger to that common frame, calibrated so the baseline exhausts in the 2034 year Smith (2015) reads from its own DYNASIM run, and scored the five solvency provisions Smith projects — levels frame-relative and ungraded, signs and orderings the content. Every sign held, fourteen of fourteen, and a registered *disagreement* held in its stated direction: raising the full retirement age to 72 exhausts Smith’s fund by under a year because his increase phases in against a fixed horizon, but on completed careers with no phase-in the cut never lets the calibrated fund exhaust at all, so it ranks first where Smith ranks it last. The revenue ordering then broke on one adjacent swap, as the benefit ordering had — the frame puts a two-point payroll-rate rise above full removal of the taxable maximum, where Smith has removal first, because only about 12.7 percent of the frame’s taxable payroll sits above the wage base, under the 16.1 percent break-even where removal’s revenue would overtake the rate rise and under the roughly 17–18 percent the administrative data carries. It is the taxable-maximum analogue of the bend-point swap: the same truncated-career compression, once at the benefit bends and once at the contribution ceiling. The record commits both as before-and-after test cases for the representative-frame transport, the milestone whose whole job is to make the frame’s distribution representative enough that both orderings come out right.

10 Four gates passed, the funded crux resolved, and the projection gate locked

Alongside the surface work, an external architecture review reset the engineering standard for deployment; a third and a fourth stage gate locked — on household composition and on the marriage-by-earnings joint — and both then passed, at the ninth and the second registered candidate; a disability gate locked and passed at its first; the representative-frame transport gate locked, failed its first candidate, was three times amended on what its own registered forensics proved, and passed at its fourth, verified bit-exactly, on a surface those amendments narrowed from sixty-five gated cells to forty-four — the funded crux resolved; an interim benefit seam carried a reform end to end through the production tax calculator; a temporal-holdout gate locked on projection drift, its first candidate stopped before scoring on a registration error and its scored run still ahead; and the program published the capability path from here to a full projection. No candidate moved a locked tolerance, and the one contract that changed after locking — the transport gate, first because its own referee-established contradiction had made it structurally unpassable, then twice more on registered forensics — changed only through the full amendment ceremony, with every demoted anchor retained verbatim, every removed cell machine-reasoned, and every failed candidate left failed.

10.1 Hardening the ledger for deployment

An external architecture review of the candidate chain through its fourteenth member returned a one-line thesis — preserve the research ledger, stop extending its execution architecture, which had grown too implicit for deployment on the production population — and its ten findings dispatched into the contract’s own machinery. A blocker, that gate 2 overclaimed which tranche it certified, became the ratified tranche amendment that split the family gate into its marital-transition, household-composition, and marriage-by-earnings tranches and promoted the rule binding each to its scored surface. Three priority-zero fixes landed as a legacy digest manifest with a no-overwrite guard, per-marker test tiers, and a seam-ownership record. The priority-one item became a flattened component registry that ports the passing sixteenth candidate into immutable, injected components with no runner imports, carrying a compatibility certificate that reproduces the candidate’s committed 20-by-46-by-5 rate cube bit-for-bit — 4,600 of 4,600 values equal to the IEEE-754 bit, signed zeros included — and its gate-2a pass verdict. A referee certified the certificate real rather than vacuous by re-deriving the cube independently and confirming that every mutation it injected was caught. The same port has since carried the household-composition tranche’s passing ninth candidate into the registry under its own compatibility certificate, the same bit-for-bit discipline.

Provenance itself then hardened. Every artifact can now carry a typed contract identity — the Git state of the locked contract plus pinned Python, NumPy, pandas, scikit-learn, SciPy, and platform metadata — written as an opt-in, append-only `.env.json` sidecar beside the artifact, with every existing call site unchanged; the transport gate’s forensics round below was the first artifact to adopt it. And a typed generic twenty-draw evaluator now re-derives what each gate’s own bespoke bindings assert — validating the exact decimal floor derivations, the registered draw stream, the locked floor references, the scoring semantics, and the four-of-five conjunction — and reproduces the three committed gate verdicts from the committed artifacts and the locked contract alone: the gate-2a pass at four of five with its seed-2 completed-fertility miss (0.1816 against 0.171), the gate-2b pass at four of five with its seed-3 `hh_size.5+` miss (0.1042 against 0.094), and, against the gate-2c floor’s own rates in candidate shape, a five-of-five compatibility result — a shape-compatibility check, not a candidate verdict. No committed run artifact, frozen file, or locked block changed: one evaluator, three gates.

10.2 The household-composition gate

The gate-2 pass certified the marital-transition tranche and named household composition as unscored; it now carries its own locked gate. Gate 2b — coresidence of spouses, children, parents, grandchildren, multigenerational households, and household size, read from the PSID relationship matrix — locked on 2026-07-10 through the ceremony gate 2 established, and the ceremony is the point, because the draft floor did not survive its adversarial round. The referee returned *amend before lock* on nine findings, three of them blockers, and earned the standing by reproducing the floor twice: a full rerun bit-identical to the committed artifact, and a fully independent recompute from its own parsers that matched all 81 reference moments, all 405 per-seed cell values, and all five holdout hashes to nine decimals. The three blockers were the discipline working, not failing. The draft protocol had reintroduced the single-draw estimator gate 2’s own amendment had already retired, making its stated operating characteristic unachievable. The multigenerational family carried a code-frame bug that mapped first-year cohabitators to a spurious extra generation, which the referee’s independent recompute showed flipped 13,400 of 91,261 multigenerational person-waves — 14.7 percent, concentrated in exactly the young cells the family gates. And the headline estimand, “waves 1969–2023,” was not the scored surface: a weight rescaling leaves the first 28 waves carrying 0.19 percent of the estimand’s weight, so the gate effectively certifies 1997–2023, which the standing rule that a tranche describe exactly its scored surface — promoted a day earlier by the tranche amendment above — forbids it to hide.

The eight required fixes adopted the ratified mean-over-draws estimator, corrected the cohabitor code and pinned the multigenerational concept to three distinct generations, restated the estimand as effectively 1997–2023, and — the review’s own precondition for any level anchor — bundled a concept-decomposed Census anchor before the flip rather than after. That anchor carries the discipline the marriage-and-divorce anchor established one tranche earlier: the PSID family unit is not the Census household, so each ratio carries a named concept factor, and after the partner-inclusion bridge all fourteen coresident-spouse cells land between 0.75 and 1.14 of their Census counterparts, the widowhood-asymmetric 0.84 at the oldest women stated honestly below one. Verification returned *lock as-is*, the maintainer ratified by merge, and the flip that inserted the thresholds — 46 gated cells against 47 report-only, the faithful-candidate operating characteristic recomputing to 0.9397 per seed and 0.9678 at four of five — passed its own flip-fidelity referee, who recomputed every number, re-fetched the ceremony comments, and mutated the new bindings to confirm they bite, clearing all eight verification sections before the merge. What the lock claimed was the threshold, not a pass: at ratification no candidate had run against gate 2b. The gate was built first; the model has since cleared it — nine registered candidates and five forensics rounds later — and the ladder that follows is the record of what the clearing cost.

10.3 The gate-2b candidate ladder

Nine candidates have run against the locked household-composition gate — its 46 gated cells, its five locked person-split seeds, and the twenty-draw mean-over-draws estimator gate 2a ratified and gate 2b adopted at lock — each registered on the campaign registry ([issue #42](#)) with a pre-registered forecast before its single scored run and graded after. Interleaved with them ran five registered forensics rounds — diagnostic, never gated, published whatever they found, each cited by the candidate it licensed — the same throughline that carried the gate-2a ladder through its four. The ladder narrows the distinct failing cells $21 \rightarrow 16 \rightarrow 14 \rightarrow 9 \rightarrow 8 \rightarrow 7 \rightarrow 6 \rightarrow 4 \rightarrow 1$ — seven candidates at zero seeds, the eighth clearing one, the ninth clearing four — and the record of what each step cost is, again, the evidence the gate exists to produce.

The base composition — candidate 1, a structural generator built from the certified gate-2a marital core, with spouse presence from the candidate-16 machinery, a logistic parental-home exit hazard, multigenerational entry and exit hazards, coresident children from the certified maternal fertility kernel, household size composed from the simulated states, and the grandchild family composed-only — fails

all five seeds across twenty-one distinct cells. What cleared from the start is as instructive as what failed: the multigenerational stocks held because observed initial states were applied from the first candidate — the gate-2a lesson applied, not rediscovered. The two dominant gaps were both concepts rather than rates. The certified machinery generates legal marriage, but the gate-2b reference’s spouse concept counts cohabiting partners — PSID MX8 codes 20 and 22 — the exact partner-inclusion concept the gate’s Census anchor bridge had quantified at lock, now expressing as a candidate failure: a measurement-concept residual of the same species as gate 2a’s undatable marriages. And the maternal fertility kernel under-attributes paternal coresident children.

Candidate 2 supplied both — a cohabiting-partner overlay of code-22 entry and exit hazards unioned onto the legal-marriage state, and paternal child attribution from birth-history father links — and narrowed the failure to sixteen cells. The overlay delivered the missing partner concept and brought the young spouse cells near clear; the father links overshot, attributing non-coresident and non-custodial children, which flipped the child residual to over-production and named custody conditioning as the next fix. Candidate 3 added the conditioning — a father-linked child counts only when the training data places the child in the father’s family unit that wave — alongside a household bridge of train-fitted non-family household members, added to the household-size composition only and feeding no coresidence cell, and skipped-generation coresidence: fourteen cells. Its structural reveal was that candidate 2’s near-clear had not held — the byte-carried spouse family fails three to four cells on every seed and is the binding constraint — and that the bridge’s minimal-count reading protected the tail of the size distribution while leaving its middle, size 3, off.

The first forensics round decomposed the four residual mechanisms. The spouse problem is pure allocation — 0.00 percent of reference spouse mass sits outside codes 20 and 22, so there is no unsupplied concept mass, unlike gate 2a’s undatable marriages — the generator over-produces young cohabitation and under-produces legal-spouse stock at 25–34 and 65-plus. The male child overshoot is roughly 88 percent linked-custodial: the gate, not the shadow kernel. The size-3 over-fill and the $2+ \rightarrow 2$ tail cap are two distinct mechanisms where the registration had guessed one. And the grandchild remainder is a skipped-generation level shortfall. Candidate 4 shipped that quartet — age-refined cohabitation, an additive legal-spouse residual top-up, a custodial-gate refinement with a tail spread, and a skipped-generation level rebuild — and reached nine cells. The spouse family cleared in both directions, the ladder’s flagship fix, its family score moving 0.74 to 0.97; and the run’s most valuable finding was an honest negative: the remaining child overshoot lives in observable-subset selection — linked children observed with their fathers are disproportionately coresident — which no refit of that same subset can remove.

The second round took the three concept-class residuals and earned the diagnostic its keep by returning two of its three registered guesses wrong. The custodial gate is faithful exactly where the overshoot had been hypothesized — for young married fathers the child-record rate reverses, higher rather than lower. The unreachable grandchild mass is a three-generation decoupling: the reference joint of multigenerational membership and a coresident own child is five times the independence product the simulation’s separate components collapse to, so the fix class is state coupling, not hazard levels. And the over-produced size-3 route is three adults, its 0.088 core gap splitting exactly into bridge-reach mass, 0.051, plus composition, 0.037 — fertility and the marital joint exonerated. Candidate 5 answered with a multigenerational-adult-child coupling — a train-fitted joint replacing the independence product — a not-married custodial correction, and bridge reach for size-3 cores: eight cells. The coupling cleared grandchild 55-plus female outright, zero seeds to five with the family score moving 0.50 to 1.00. But the not-married correction missed, the overshoot being married-dominated; the child-record basis, higher at ages 0–4, tipped the 15–24 male cell from five seeds to none; and the two household-size mechanisms traded off under the current bridge.

The third round worked the five-cell endgame. The 0–4 discrepancy is a join-denominator artifact — the observable basis is right at ages 0–4, the child-record basis right at school ages — and the three

failing child cells carry three different mechanisms. The 25–34 female miss is a cohabitation-overlay shortfall, minus 0.045, not a legal-core one — a registration miscitation owned on the record. And question 10 proved the decisive positive: a train count-conditional bridge clears all three household-size cells simultaneously on the simulation’s own core distribution. Candidate 6 pulled the four measured levers — the 0–4 basis revert, adult-child exit timing on both parent sides, a female cohabitation lift at 25–34, and the count-conditional bridge — and reached seven cells: `hh_size.3` and `hh_size.4` clear on all five seeds, `hh_size.5+` on four, and the failure surface narrows to one mechanism wide. One lever proved inert — the female-cohabitation refit coincided with the existing estimator — and though the cell moved, the attribution was deferred to forensics rather than banked.

The fourth round sized linked-father supply and ran two attribution checks, and its central number reframed the child overshoot: the dominant linked driver is unenumerated non-joinable supply, plus 0.035 and plus 0.036 — the committed mechanism had been applying custody probabilities to biological children with no enumerated household record, coresidence the reference roster can never observe, 25.8 percent of linked exposure. The spell channel is real but secondary and opposite in sign — the simulation fragments episodes, mean 3.57 waves against the reference’s 5.93, it does not lengthen them — and the lane caught mid-run that a naive enumerated-only reference would have falsely confirmed the registration. The round also priced what it did not fix: the fragile spouse cell is inherited machinery riding the tolerance line at two of five split seeds, carried rather than banked, and the `hh_size.5+` seed miss is the three-plus-child fertility deficit — structure, not noise. Candidate 7 answered with enumeration conditioning — the paternal linked-coresidence draw restricted to enumerated children — and episode persistence, a correlated entry–persist–exit process with sibling-synchronized frailty — per-child independent chains provably remove zero occupancy — fitted to the training episode-length distribution while preserving the per-wave marginal. Both child male cells cleared on all five seeds and the failure narrowed to six. But the conditioning exposed a joinable adult-child supply deficit at 55–64 and 65–74 male that the unenumerated inflation had been masking — 65–74 male now fails all five seeds — a compensating-errors reveal of the same species as gate 2a’s pooled-band findings: fixing a concept artifact uncovers the real deficit beneath it. The priced risks — the fragile spouse cell, the structural `hh_size.5+` seed — materialized as priced.

The fifth round decomposed the revealed deficit and ran the convergence test. Fertility dominates at 55–64 male; at 65–74 male the near-zero adult-child kernel attenuates the fertility endowment, so exit and link coverage co-dominate; and 45–54 female is over rather than under, exit over-retention at plus 0.079 — the same hazard family needing band-specific corrections in both directions. Question 15 proved a completed-fertility lift clears `hh_size.5+` and 55–64 male but not 65–74 male; question 16 proved the cohabitation lift clears the fragility without collateral. Candidate 8 shipped the two proven levers — the fertility-core lift, the cohabitation lift at 25–34 female — and a band-signed adult-child retention refit at parent ages 45-plus, lifting 65–74 male retention up while bringing 45–54 female over-retention down with the same train-fitted structure. It cleared the ladder’s first seed, at four distinct failing cells: 65–74 male moved from zero seeds to five and both band signs held. But the global fertility lift overshot four previously-cleared middle-cohort child cells — a registration scope error, priced below, that only the holdout could catch. Candidate 9 registered one delta — the same lift, confined to exactly the deficit cohorts the fifth round had measured — and passed, four seeds of five with a single failing cell.

Table 4: The gate-2b candidate ladder, from the committed run artifacts (`runs/gate2b_hazard_v1.json` through `v9.json`). “Distinct failing cells” counts gated cells missing on at least one of the five locked seeds, the same convention as Table 2; the gate requires at least four of five seeds with every one of the 46 gated cells inside its locked tolerance, which candidate 9 is the first to meet. All nine candidates are scored under the mean-over-draws estimator gate 2b adopted at lock.

candidate	distinct failing cells	seeds passing	headline lesson
1 — structural composition from the certified 2a core	21	0 / 5	legal-marriage machinery is not the 2b partner concept; the maternal kernel under-attributes paternal children
2 — + cohabitation overlay, paternal father-link attribution	16	0 / 5	the overlay supplies the partner concept; father links overshoot — custody needs conditioning
3 — + custodial conditioning, household bridge, skipped-generation	14	0 / 5	the byte-carried spouse family is the binding constraint; the bridge leaves the size-3 middle off
4 — + the forensics-1 quartet	9	0 / 5	spouse clears both directions; the child overshoot is observable-subset selection, unfixable by refit
5 — + multigen-adult-child coupling, not-married correction, bridge reach	8	0 / 5	the joint clears grandchild 55+ female 0/5→5/5; the not-married correction misses the married-dominated overshoot
6 — + the four measured levers	7	0 / 5	hh_size clears; the failure surface is one mechanism wide; the female-cohab refit proves inert
7 — + enumeration conditioning, episode persistence	6	0 / 5	both child cells clear, but conditioning reveals the masked 65-74 male supply deficit — a compensating-errors reveal

candidate	distinct failing cells	seeds passing	headline lesson
8 — + two proven levers, band-signed retention refit	4	1 / 5	first seed cleared; a global fertility lift overshoots on a train-side-proven scope that does not transport
9 — + cohort-scoped fertility lift	1	4 / 5	PASS — scoping the lift to the measured deficit cohorts clears the ladder; the sole miss is the registered hh_size.5+ residual

What the ladder certifies must be stated as carefully as what it found, because a ladder is an audited iterative search, not nine independent experiments. The candidates are one-shot — each specification frozen and its forecast registered before the single scored run — but the forensics questions between them are selected by the holdout failure patterns the previous run exposed, and that selection is a leak channel no per-candidate operating characteristic prices. The effective search is nine one-shot candidates and five forensics rounds against the same five holdout halves, roughly 2,070 scored cell-by-seed exposures. The locked operating characteristic — 0.9678 at four of five, 0.9397 per seed, the figures the lock ceremony recomputed — prices a single pre-registered attempt’s false-negative risk on the draw-noise-free floor basis; it is not a null false-pass rate, and no locked quantity supplies a family-wise number for an adaptively constructed ladder, so none is quoted. The certified claim is therefore the precise one: the generator matches the held-out household and relationship-composition moments after an audited nine-candidate, five-forensics search, under the program’s domains-of-validity doctrine — not on a first attempt out of sample. The registrations and gradings on issue #42 are the audit trail.

Two standing rules were adopted mid-ladder, at candidate 8, out of a process audit rather than a failure: pass-run verification — a pass does not enter the record tally until an independent adversarial round has reproduced it bit-exactly, mutated its bindings, and scrutinized its deltas, the treatment gate 1’s run 13 received — and the ladder-search disclosure above. The timing is the point: the program tightened its own governance before a pass existed to flatter, and that is the discipline working.

One registration error is worth naming rather than absorbing. Candidate 8 over-read a train-side proof as holdout-transportable: the fifth forensics round had verified that the completed-fertility lift held every already-cleared cell, but that verification was computed on the training half, and on the holdout the reference rate sits below the train rate for the already-adequate middle cohorts — so the same faithful lift overshoot four previously-cleared child cells. The error lived in the registration’s scope claim, not in the levers or the diagnostic; it was graded against the registrant and published as-run under the no-holdout-tuning rule, and candidate 9 then confined the lift to exactly the deficit cohorts forensics 5 had measured. That is the audited search visible in one candidate: a scope claim only the holdout could falsify, falsified in public, fixed by the next registration.

10.4 The first gate-2b pass

Candidate 9’s delta is one clause of scope. The completed-fertility lift candidate 8 had applied globally is confined to exactly the four deficit cohorts the fifth forensics round measured — 55–64 and 65–74 for men, 45–54 and 65–74 for women — and the non-deficit middle cohorts revert byte-identically to candidate 7. The candidate fitted zero new quantities; its one new degree of freedom is a discrete four-cohort scope taken from a train-side forensics measurement.

The write-gate implementation made the reversion checkable at bit precision. The eight reverted child cells are bit-identical to candidate 7 — a difference of 0.0 across all twenty draws and five seeds — and the thirty-three carried deficit-and-cleared cells bit-identical to candidate 8, so forty-one of the forty-six gated cells reproduce bits already published in the candidate-7 and candidate-8 artifacts, and the only genuinely new holdout numbers in the run are the five `hh_size` aggregates. The delta is auditable at the bit, not the narrative, level.

The gate passed four of five, per-seed passing-cell counts 46, 46, 46, 45, 46. The sole miss is `hh_size.5+` on seed 3, score 0.1042 against a 0.094 tolerance — the registered modal residual, and the one cell the pre-run analytic check had flagged in advance, its train-side prediction of 0.1291 landing beside the realized holdout seed-mean of 0.1286.

Under the standing pass-run rule the pass was independently verified before entering the record: a referee re-executed all five seeds from the committed specification and reproduced the run’s 20-by-46-by-5 rate cube bit-for-bit — 4,600 of 4,600 values — independently rescored all 230 cell-by-seed verdicts from the locked contract, ran a clean leak hunt confirming every fitted quantity is train-half only, and caught all three mutations injected to prove the checks bite. It is a model delta under a locked gate — thresholds, protocol, and the ratified estimator all fixed before the run was registered — not an amendment rescuing its own trigger; candidates 1 through 8 stand as committed failures. It certifies the household and relationship-composition tranche the locked cells score — coresidence of spouses (legal and cohabiting), children, parents, and grandchildren, multigenerational households, and household size — on the observed frame, and no more: the marriage-by-earnings joint — gate 2c — remains the locked-but-unpassed tranche, and the sole failing cell anywhere in the passing run is the registered `hh_size.5+` deep-tail residual.

10.5 The marriage-by-earnings gate

The gate-2 pass certified the marital-transition tranche and named the marriage-by-earnings joint — the surface spousal and survivor benefit *levels* key on — as a separate later tranche; it too now carries its own locked gate. Gate 2c — the own-by-spouse assortative-mating contingency, the earnings-conditional first-marriage and remarriage hazards, the around-event earnings dynamics, and the couple’s shared-earnings distribution, all on a per-year indexed-earnings axis over the selected universe of PSID couples with a computable earnings history for both partners (7,994 directed couples over 14,952 earnings-supply persons) — locked on 2026-07-10 through the same ceremony, and once again the ceremony is the point, because the draft floor did not survive its adversarial round. The referee returned *amend before lock* on eight findings, three of them blockers, and earned the standing the household referee had, by reproducing the floor twice: a full rerun bit-identical to the committed artifact, and a fully independent recompute from its own fixed-width PSID parser, AIME chain, and splitter that matched all forty-nine reference moments, all two hundred forty-five per-seed cell values, and all five holdout hashes to nine decimals.

The three blockers were the same three classes the earlier tranches had surfaced, each caught before the lock. The first was the floor: the draft measured its noise between two person-disjoint halves, but a couple surface is not person-disjoint — under the person split half the couples straddled the two sides (50.4 percent of mirrored pairs on one seed), so the floor priced person-level noise where the tranche gates couples, and the disclosed faithful-candidate operating characteristic of 0.95/0.977 was 0.83/0.80 at the couple-level noise a faithful model actually faces. The second was the estimand: the artifact described its earnings window as 1993–2022 while the committed panel spans 1968–2022 — a twenty-five-year misstatement of its own observation surface, with 69.0 percent of the earnings-supply persons carrying a pre-1993 positive year, and a candidate faithfully implementing the *described* window failing five of thirteen gated cells on the description alone; the standing rule that a tranche describe exactly its scored surface forbids it. The third was construct validity, and it is the assortative-mating sibling of the compression fingerprints the reform-scoring surface had already turned up: the career-summed

earnings proxy the draft gated correlated only 0.12 across spouses, but that number is the observation process, not earnings sorting — cross-sex tercile pooling and the anti-correlated observed-year counts of single-earner-era couples together halve it, and the proxy’s 0.977 split-half reliability proves the weak signal is what the career sum measures, not sampling noise. The within-couple rank on per-year indexed earnings is 0.49, literature-scale; gated on the proxy the joint would have certified the observation mechanics and inverted the tranche’s purpose, failing a candidate with true earnings sorting and passing one that merely reproduced PSID’s observation cadence.

The eight required fixes rebuilt the floor couple-disjoint by connected component of the couple graph (13,275 components, none larger than four persons, zero couples straddling), restated the true 1968–2022 estimand everywhere, re-specified the earnings axis onto per-year indexed earnings, pinned the candidate’s directed both-orientation couple emission (a single-orientation emission was measured non-conformant, off by up to 4.9 times), and detrended the around-event windows against a placebo drift deflator of 1.2019 — roughly three-quarters of each raw window’s magnitude was generic nominal life-cycle drift any window on this panel shows, so a real-terms candidate is no longer failed by construction. Re-priced on the rebuilt floor, twenty-seven cells gate against twenty-two report-only, and the faithful-candidate operating characteristic recomputes to 0.9641 per seed and 0.988 at four of five. The review’s precondition for the flip, as at the household gate, was an external anchor bundled before it rather than after: a concept-bridged assortative-mating anchor built from the CPS spouses’-earnings correlation of Schwartz (2010) and the educational assortative-mating series of Greenwood et al. (2014), reported and never gated. Its per-year rank of 0.49 sits above Schwartz’s 0.23 dual-earner annual-earnings correlation and its earnings-contingency diagonal delta of 1.37 below Greenwood’s 1.6-to-2.0 educational delta — both directions the named concept deltas predict — and it moves no floor value. Verification returned *lock as-is*, the maintainer ratified by merge — the ceremony branch re-homed across three pull requests as its merge reference settled, the earlier two closed unmerged on the same floor — and the flip that inserted the thresholds passed its own flip-fidelity referee, who recomputed every tolerance from the frozen floor, re-fetched the ceremony comments, re-derived the anchor against the archived Schwartz and Greenwood sources, and confirmed the new bindings bite across all eight verification sections. What the lock claimed, again, was the threshold, not a pass: at ratification no candidate had run against gate 2c. Two registered candidates have since settled it, and the ladder that follows — the shortest of the three family ladders — is the record of what the settling cost.

10.6 The gate-2c candidate ladder

Two candidates have run against the locked marriage-by-earnings gate — its twenty-seven gated cells, its five couple-disjoint split seeds, and the twenty-draw mean-over-draws estimator — each registered on the campaign registry ([issue #42](#)) with a pre-registered forecast before its single scored run and graded after.

Candidate 1 composed couple formation from the certified tranche-2a marital core and train-fitted joints, and failed, one seed of five passing, per-seed gated-cell counts 25, 26, 27, 26, 26 of 27. The registration had named its modal failure classes in order — the assortative-mating contingency diagonal first, the earnings-conditional first-marriage hazard cells second, the around-event cells third — and the run resolved the order. The diagonal held completely, thirty-five of thirty-five cell-seeds: the correlation structure the tranche exists to certify survived composition on the first attempt. Every other family cleared every cell on every seed, and the gate bound entirely on the second named class: the five misses, all marginal at 1.01 to 1.34 times tolerance, are three first-marriage-by-earnings cells — the bottom-tercile 18–24 female hazard on seeds 0 and 4, the top-tercile 25–34 male on seeds 0 and 3, the top-tercile 18–24 female on seed 1. The mechanism is structural rather than a mistuning: the certified marital core conditions first-marriage timing on age, sex, and cohort — not on the earnings axis — so a tercile-conditional first-marriage hazard is just the age-and-cohort hazard averaged over that tercile’s age and cohort mix. The registered forecast priced a pass at 0.10 to 0.25; the fail graded

as the forecast’s modal outcome.

Candidate 2 registered one delta: an earnings-conditioned first-marriage-timing modifier — train-fitted, multiplicative, by tercile within age band and sex, composed onto the certified 2a first-marriage hazard, with the earnings-blind core read and never re-simulated so the fit consumes no randomness — normalized so the certified age-by-sex timing marginal does not move (the per-band sum of modifier times certified timing is one, hand-verified to a maximum deviation of $2.2\text{e-}16$), the same marginal-preservation constraint class the household tranche’s coupling delta used. The gate passed, four of five, per-seed gated-cell counts 27, 27, 25, 27, 27. All five candidate-1 misses resolved, and the honest asymmetry publishes with the pass: the sole failing seed is seed 2 — the one seed candidate 1 passed — where the delta moved two 18–24 female first-marriage cells the earnings-blind core already matched from roughly 0.19 and 0.07 of tolerance to 1.036 and 1.065 times it (0.238 against 0.230, 0.287 against 0.269). Byte-carry makes the delta auditable at the bit level, as the household pass was: the twenty-five carried-family cells are byte-identical to candidate 1 across 2,500 per-draw comparisons (maximum deviation 0.0), the first-marriage family moved on all 120 of its cell-seeds by exactly the modifier (the candidate-2 rate equal to the modifier times the candidate-1 rate to $5.6\text{e-}17$), and nineteen of the twenty-seven gated cells reproduce bits already published in the candidate-1 artifact — the genuinely new holdout exposure was the eight gated first-marriage cells. The registered forecast priced the pass at 0.45 to 0.65 and the run realized it; the forecast’s modal-residual call — the top-tercile 25–34 male cell — did not describe the realized residuals, and that miss is graded and disclosed with the rest. An -invariance check executed rather than assumed: at of 0, 8, and 20 the verdict is identical, four of five with the same failing set.

Table 5: The gate-2c candidate ladder, from the committed run artifacts (`runs/gate2c_hazard_v1.json` and `v2.json`). “Distinct failing cells” counts gated cells missing on at least one of the five locked seeds, the same convention as Table 2; the gate requires at least four of five seeds with every one of the 27 gated cells inside its locked tolerance, which candidate 2 is the first to meet. Both candidates are scored under the twenty-draw mean-over-draws estimator and the couple-disjoint split seeds gate 2c locked with.

candidate	distinct failing cells	seeds passing	headline lesson
1 — couple formation composed from the certified 2a marital core	3	1 / 5	the assortative diagonal holds 35/35 on the first attempt; the gate binds entirely on earnings-conditional first-marriage timing, the registration’s second-named class
2 — + earnings-conditioned first-marriage-timing modifier, timing marginal preserved	2	4 / 5	PASS — one train-fitted delta resolves all five candidate-1 misses; the sole failing seed is the one candidate 1 passed

10.7 The first gate-2c pass

Under the standing pass-run rule the pass entered the record only after an independent adversarial round, and because the round was the tranche’s first, the referee verified the carried candidate-1 machinery too rather than assuming it. The re-execution was bit-exact — the 2,700-rate twenty-by-

twenty-seven-by-five cube reproduced bit-for-bit, the only deviations six elapsed-seconds fields — all 135 cell-by-seed verdicts rescored from the committed cube with no repository code, the byte-carry recomputed at 0.0 across all 2,500 comparisons, the χ^2 -invariance check executed, a leak trace confirming that no candidate-2 addition is holdout-fitted, the couple-component split verified with zero couples straddling the two halves on any gate seed, and every injected mutation caught.

The search accounting states what two candidates mean. This is the shortest ladder of the three tranches — sixteen candidates and four forensics rounds at the marital-transition tranche, nine and five at the household tranche, two and zero here, roughly 270 gated cell-seed scores — and it is short because it inherited everything: the certified 2a marital core, the per-year indexed-earnings axis and its cut provenance from the gate’s own fixes round, the marginal-preservation constraint class from the household coupling delta, and the ladder method itself. Candidate 2 was designed from candidate 1’s published holdout decomposition — one adaptive step, along a delta axis candidate 1’s registration had named as a modal failure class before any 2c holdout was seen, with a registered forecast pricing the outcome. The certified claim is scoped accordingly: the generator matches the held-out marriage-by-earnings joint moments after an audited two-candidate iterative search, under the domains-of-validity doctrine — not on a first attempt out of sample — and no locked quantity supplies a family-wise error number for a two-candidate adaptive family, so none is quoted.

The pass is also a milestone: every gate-2 tranche is now passed. All four demographic gates — earnings; marital transitions and fertility; household and relationship composition; the marriage-by-earnings joint — are locked and passed, each behind its own pre-registered contract, adversarial lock ceremony, one-shot candidate ladder, and independently verified pass run. On the observed frame the certified surface now covers the full demographic input surface the cataloged provision classes score on: career earnings histories; marriage, divorce, remarriage, widowhood, and fertility; household and relationship composition, from cohabitation and custody to skip-generation households; and who marries whom, by earnings. What remained at that milestone — the balance of this section, and the program’s open frontier — is the disability gate, the representative-frame transport, the projection engine with its temporal-holdout gate, and full revenue and trust-fund accounting.

With that, all three tranches of the family gate are locked and passed: the marital-transition surface at the sixteenth candidate, household composition at the ninth, the marriage-by-earnings joint at the second. Across the three the adversarial round kept returning the same three failure classes — a floor or estimator that priced the wrong noise, an estimand that misdescribed its own surface, and a construct whose signal was the measurement rather than the thing measured — and the ceremony quantified and fixed each before the threshold locked, none after. That is the design working rather than failing, and it is why a locked gate here is worth what it claims: the hard test comes before adoption, in public, against a referee who recomputes. One class the pre-lock ceremony did not catch — a gate that locks wrong anyway — arrived at the transport gate below, and the same machinery caught that too, after the lock.

10.8 The disability gate

The disability module added the last major transition process, first in reported-not-calibrated form and now behind its own locked and passed gate: DI incidence and recovery hazards from PSID self-reported work limitation, and the statutory conversion at full retirement age, where a disabled worker’s benefit becomes a retirement benefit at the same primary amount — a factor of exactly one — so that the claiming sampler’s excluded conversion mass reconstitutes the full administrative entrant mix. The gate’s first design decision is what it refuses to score. The PSID self-reported work limitation is not the SSA disability program: seven named concept deltas separate a self-reported labor-force status from an adjudicated award — definition, insured population, severity threshold, recovery churn, conversion denominator, biennial timing, and secular period — so no SSA DI level is gated anywhere. The estimand is the self-report panel’s own disability processes, with the administrative record entering as

concept-bridged anchors gated on shape and direction, never level. The module’s founding validation kept that pattern: the PSID disabled-to-retired transition, read against the archived conversion column of the 2023 SSA supplement, runs at 0.267 of the administrative share for women and 0.322 for men, reported as a ratio far below one and never calibrated toward it.

The gate is anchor-based in gate 1’s external-check style, not holdout-only, and its floors ceremony restored teeth rather than filing them down: the round-one referee’s first finding — a blocker — required restoring the pre-registered internal families the draft had weakened and adopting a mixed-multiple remedy — flow-hazard tolerances at the floor mean plus three seed standard deviations and the 50-to-59 prevalence stock at plus four — eight gated internal cells with the occupancy-stock teeth restored. The same round adopted the ratified draw stream and fresh-run schema and pinned the candidate-side anchor statistic — per gate seed, on the twenty-draw mean, with a pinned minimum margin of three standard deviations. The lock then added gate_m4 as a new top-level gate: twelve gated cells — the eight internal mixed-multiple cells plus four concept-bridged anchor margins — against twenty-four report-only, with the in-repo, hash-pinned 2023 DI Annual Statistical Report tables as the anchor, reported and never gated as levels; the faithful-candidate operating characteristic is 0.9408 per seed and 0.9689 at four of five. The flip referee’s first finding fixed the gate’s weight definition to the ratified start-wave weight and verified the fix with teeth: the 50-to-59 male incidence rate recomputed under start-wave weights equals the committed reference exactly, 0.02151029, where end-wave weights give 0.02015166 — the check discriminates.

The first candidate was the module’s own hazard machinery — refit per split seed on the training half, simulated forward on the holdout support with the ratified twenty-draw estimator and the pinned start-wave weights, no free hyperparameter — and it passed five of five, the program’s first five-of-five pass. All eight internal cells hold on every seed, the worst at 0.88 of tolerance (the 50-to-59 female incidence hazard on seed 4); all four anchor margins hold on every seed, minimum margin 4.39 standard deviations against the pinned 3.0, the prevalence age-shape margins running 6.5 to 7.5 and conversion-exit dominance 4.4 to 13.7; zero undefined draws. The registered forecast priced a pass at 0.55 to 0.75, and five of five is the faithful model’s own modal outcome, probability 0.737.

Verification applied the standing first-pass round — bit-exact re-execution (the 800-rate cube at 800 of 800 equal), all sixty cell verdicts rescored with no repository code, the candidate’s train-side copy equal to the frozen floor to 1.39e-16, the weight fix verified as above — and its sharpest work was a construction adjudication executed rather than argued. The fit re-initializes post-gap episode state from the train-fitted stock, and 20.2 percent of holdout person-years are episode starts — 14.6 first-wave, 5.65 post-gap — so the referee re-scored all five seeds and twenty draws under both alternative gap constructions, an ordinary hazard step and carry-forward of the previous state. Both still pass five of five with the same worst-cell identities, and the committed re-initialization scores worse on the gated prevalence stock than either alternative (seed-4 male 0.2646 against 0.2426 and 0.2374): the construction carries no score advantage and no stock-information leak. The leak hunt closed the loop from the other side — the simulator reads only person identity, period, age, weight, and sex from the holdout, a strictly smaller holdout footprint than the verified 2b and 2c candidates, which kept observed initial states; here even initial states draw from the train-fitted stock. One candidate, zero forensics rounds — the shortest ladder in the program — and the certified claim is scoped to match: the disability module’s hazard machinery reproduces the held-out PSID work-limitation moments and holds the concept-bridged anchors after a one-candidate registered search, under the domains-of-validity doctrine, and never an SSA DI level — the seven concept deltas stay report-only.

10.9 The transport gate

The transport of the certified generators from their PSID estimation basis onto the representative population is the program’s scientific crux and its funded build, and it now carries a locked gate — gate_w1 — a failed first candidate, a registered forensics round, and the program’s fifth ratified

amendment. The gate certifies deployment on the certified populace CPS frame — release us-4.18.8, sha-verified at deployment, 166,302 persons representing a weighted 340.0 million — a frame none of the generators was fitted to. The locked surface is 65 gated against 67 report-only cells in three target families. Family A, the CPS-observable joints — 53 gated, 52 report-only — floors on a hundred-seed household-disjoint half-split of the certified frame’s own weighted earnings, age, sex, marital, and household moments, each tolerance the floor mean plus four standard deviations rounded to three decimals and capped at $\ln(1.5)$, with a heavy-tail guard so a faithful candidate passes and the identity candidate — copying the scored columns, score zero at zero dispersion — declared non-conformant rather than victorious; the faithful-candidate operating characteristic is 0.922 per seed and 0.9481 at four of five. Family B, the SSA administrative anchors — 10 gated, 15 report-only — carries two simulated disability-conversion margins and eight DI age-composition prevalence bands. And family C gates the two compression fingerprints Section 9 committed as before-and-after test cases, as binary reversal checks whose required orderings derive from the committed Mermin and Smith anchors: progressive price indexing must rise past the retirement-age increase, and taxable-maximum elimination past the two-point payroll-rate rise. The floors went through the full ceremony — a round-one referee returned amend on ten findings, eight fixes followed, verification returned amend on the body text alone, the body was amended, and ratification came by merge — with the twenty-draw estimator adopted from the start.

One process violation is owned on the record rather than absorbed: the floors merge carried a test-tier manifest recounted in a gate virtual environment missing an optional dependency, under-collecting ten conditionally-collected unit tests — 189 against the continuous-integration count of 199 — and turning the mainline tier-policy check red. The repair restored the rule and named it: CI collection is authoritative, and the manifest is recounted CI-equivalent or not at all.

Candidate 1 was the first full transport deployment — the certified generators together, the gate-1 earnings process, the marital candidate 16, the household candidate 9, the disability module, and the reform ledgers, deployed on the certified frame — and it failed, zero seeds of five. Family A passes 28, 28, 32, 26, and 24 of its 53 cells per seed; family B passes zero of ten; neither family-C fingerprint reverses in full. The registered forecast priced a pass at 0.05 to 0.20; the zero-of-five fail graded as its modal outcome. The reads are the run’s value. The fingerprints, re-run changing only the frame: progressive price indexing does not lift past the retirement-age increase — the transported AIME mass above the second bend point is not enough — while the elimination-versus-rate-rise swap is realised, exactly the predicted mechanism, though the same heavier tail lifts a \$150,000-cap variant above both payroll-rate provisions, so the full committed ordering does not reproduce cleanly. Both readings are robust: the orderings are identical across three independent transport constructions — level-anchored, rank-flat, and rank with transitory mobility — and the transported frame carries 37.7 percent of taxable payroll above the wage base, an upper read (prime-age positive earners, no zero or low years) well above the administrative 17-to-18-percent share and the PSID frame’s 12.7; a heavier tail favours reversal, so the price-indexing non-reversal is conservative. Family A’s modal failures decompose by view: the marital share accounts for 45 cell-failures over the five seeds — the synthetic-panel deployment starts everyone never-married and the hazards do not equilibrate over the window, so younger cohorts under-produce married — and earnings participation for 25, because the gate-1 process carries no sex covariate and fits ages 25 to 59, so 18-to-24 and 62-to-69 extrapolate to the nearest fitted bin. The earnings regeneration itself is not degenerate: nineteen earnings cells pass, including the p90/p50 and p50/p10 dispersion ratios, with a maximum across-draw standard deviation of 0.475 — a real regeneration, not an identity map. And family B’s zero of ten has a shape: the simulated work-disability prevalence gradient transports, but it is far flatter than the SSA disabled-worker beneficiary stock’s older-age concentration — 18.6 percentage points deployed against the anchor’s 45.4 at the 60-to-full-retirement-age band — so the level and the steepness do not.

The registered forensics round — reported, never gated, the first artifact to adopt the environment sidecar — measured the five initialization, support, and scope mechanisms the candidate isolated, with

its instrumentation bit-identical to the committed transport machinery at 0.0. Extending marital exposure barely helps (a mean absolute contribution near 0.028) while the hazard-level residual dominates (near 0.095), and initializing from the frame’s own marital column is the prohibited regenerated-surface identity — so entry-state seeding, not observed initialization, is the necessary lever. Nearest-bin boundary handling clears zero of six boundary cells where a train-fitted boundary extension clears four — the 62-to-69 bands and the 18-to-24 profile — while 18-to-24 participation resists: PSID heads and spouses overstate it against a CPS all-person frame. Scope and composition telescope, and a large composition residual — over-generation of size-one lone-adult households — remains, so scope is not the whole miss. The corrected, lighter earnings tail moves the price-indexing savings down, 0.0169 to 0.0137, widening the gap to the retirement-age cut’s 0.202 — the non-reversal is robust in the conservative direction, and the elimination swap holds under both tails. And the fourth question returned a finding about the gate rather than the candidate.

That finding established, against the locked contract, that the eight family-B DI age-composition bands are unclearable by any contract-consistent candidate. The bands score the disability module’s work-disability point-prevalence against the SSA disabled-worker beneficiary stock — a duration-accumulated stock, read from the archived DI statistical tables — across a concept bridge the gate never defined; the concept delta dominates the miss (aggregate share 0.595, per-band 0.023 to 0.891; at the 60-to-full-retirement-age band, duration accumulation contributes +21.3 percentage points against +2.6 of simulated shape), the insured denominator is not even archived, and all eight bands miss by 2.9 to 21.9 times tolerance. Family B is a conjunction, so `gate_w1` as locked had a faithful-candidate operating characteristic of structurally zero: a gate no faithful candidate could pass. The response was the amendment ceremony, in full — a proposal committed inert with the locked block untouched, a round-one referee who returned amend-the-amendment, fixes, a verification round that ratified as-is, and the flip. The round-one demand is the ceremony’s teeth showing: the proposal had demoted the eight DI bands, and the referee required extending the demotion to all ten family-B cells, because the two disability-conversion margins fail the same candidate-independence test on committed evidence — the module’s conversion validation ratio, 0.267 and 0.322, was never a level match. The ratified flip demotes all ten to report-only with machine-readable reasons committed (`concept_bridge_undefined_di_stock`, `conversion_level_match_never_certified`), retains all ten anchor values and tolerances verbatim under `family_b.retained_anchors` — zero threshold movement, nothing deleted — and moves the roll-ups from 65 gated and 67 report-only to 55 and 77. The gate’s operating characteristic returns from a structural zero to the family-A characteristic, 0.9481, on the residual surface, still conjoined with the two family-C fingerprints. And the no-self-rescue rule holds on committed grounds: candidate 1 stands failed independently of the amendment — it fails family A at zero of five, both retained conversion cells, and family C — so the demotion rescues nothing; the amendment is prospective, and the forensics evidence behind it is train- and frame-side.

This is the fourth referee class, and the reason the amendment mechanism exists. The three family-tranche ceremonies caught a mispriced floor, a misdescribed estimand, and a measurement standing in for its construct — each before its lock. A gate can also simply be wrong: locked carrying an internal contradiction that makes it structurally unpassable, which no pre-lock round caught. The same ceremony catches that too, after the lock, on the gate’s own committed evidence, and the record keeps the two events distinct — an amended gate is not a better candidate, and the candidate that exposed the contradiction stays failed.

The second candidate deployed the same certified generators with the three changes the first forensics round had proved, and failed, zero seeds of five — but inside its registered failure band, priced at 0.25 to 0.45, with the registered modal shape almost exactly realised: the price-indexing binary plus a coherent band of family-A stragglers, though about 15 of them rather than the registered one to three. Family A rose to 35 to 39 of its 53 cells per seed against the first candidate’s 24 to 32, and the run’s value is again in what the failure isolated. The young-cohort married deficit did not close; it overshot — the 25-to-34 married share deployed at 0.555 against the frame’s 0.387 — while a 65-plus undershoot

emerged, 0.719 against 0.844: the first candidate had run under, the second ran over, the frame’s value between the two deployments. The price-indexing non-reversal was confirmed a second time, under the strongest deployment yet built — the committed evidence the gate’s second amendment would cite. And five coherent residual mechanisms fell out of the reads: the initial coresident-parent rosters and the fertility window (household size), the marital calibration frame at 25-to-34 and 65-plus, the 18-to-24 participation concept delta, and a missing interior sex covariate.

The second registered forensics round asked four questions — reported, never gated, its instrumentation bit-identical to the committed machinery at 0.0, every measurement train- and frame-side. On the marital calibration frame, the entry-state level dominates the 25-to-34 overshoot — contributing +0.124 and +0.093 of the miss against near nothing from the hazards — and the contract adjudication is exact: a CPS-anchored entry model is permitted, while back-solving entries to reproduce terminal states is the prohibited identity in disguise. The 65-plus channel is not widowhood — deployed widowed sits below the frame — but divorce over-accumulation, +0.123 and +0.124 of excess with deployed 65-plus divorced near 0.20 against the frame’s 0.05 to 0.10: a cohort-vintage mismatch between the pooled certified hazards and the frame’s older cohorts, which the permitted entry lever cannot fix. The other three questions returned a refutation, an exceedance, and a consolidation. The joint household-size feasibility was refuted: the two entry-state levers move every cell toward the frame but clear only size-2, and they partially self-offset — a fuller fertility window reverts the marital seeding — so the pre-registration was wrong, and the artifact says exactly how. The interior sex covariate exceeded its registration: four of four cells clear with no collateral damage to any passing cell — a missing covariate, not a hazard defect. And the concept cells hardened: the 18-to-24 participation gap measured 22.1 percentage points — a PSID head-and-spouse 0.89 against a CPS all-person 0.64 — reproduced directly on the train side by universe restriction, and the price-indexing non-reversal consolidated three ways, analytic plus two empirical runs. Against the locked contract the consequence was structural: with two 18-to-24 cells unclearable by any candidate (a population-concept delta), the price-indexing fingerprint unclearable, and the 65-plus marital pair unfixable in the deployed hazard class, the gate remained unpassable as amended — the second amendment had to precede the third candidate.

That amendment ran the full five-round ceremony — proposal, adversarial referee, fixes, verification, second fixes — and its teeth showed twice, both times on a false universal. The referee confirmed the demotion boundary was principled — its own five additional construction attempts moved the 65-plus cells away from the frame — but caught the blanket per-cell impossibility language: the first candidate had in fact passed one 65-plus coresident-spouse cell, the female one, at four of five. The fixes rebuilt the section on per-cell grounds; verification then caught that the fixes commit had reintroduced one re-blanketed impossibility sentence — false for the same cell — and inherited a second, and the second fixes round bound both, the mutation analogues now guarded. The ratified flip demoted six family-A cells to report-only with machine reasons committed — two 18-to-24 population-concept deltas, three 65-plus cohort-vintage hazard mismatches, and one scored duplicate of a demoted married quantity — demoted the price-indexing fingerprint on the twice-confirmed non-reversal, and retained every tolerance verbatim. The roll-up moved from 55 gated and 77 report-only to 48 and 84; the family-A operating characteristic recomputed from 0.922 per seed and 0.9481 at four of five on 53 cells to 0.9344 and 0.9623 on 47. And the amendment met the whittling question — whether successive demotions carve a gate toward whatever the next candidate can pass — in its first strong form: it named the shrinking surface itself, and it filed a series forecast that pre-named the household-size quad and the \$150,000-cap adjacency as the next candidate’s residual risks, before that candidate ran.

The third candidate bound all three forensics-proven levers — a CPS-anchored entry-level marital model, the interior sex covariate, and a co-designed coresident-parent roster and fertility window — and failed, zero seeds of five; the registered forecast had priced a pass near 0.15, and the fail landed in its roughly 85 percent modal mass. All three levers landed. Family A passed 43 of its 47 gated cells in-band on all five seeds, every 25-to-34 marital cell among them — those passes carrying the partial-self-reference caveat the registration had disclosed in advance — along with the interior sex cells

and size-2. The residual was the pre-named household-size quad, and here the forecast matched exactly: the failure union across all seeds was precisely sizes 1, 3, 4, and 5-plus, with size-2 clearing everywhere and size-4 clearing on seed 1 — the modal outcome the registration had named. The elimination fingerprint’s core mechanism realised a third consecutive time — the elimination-versus-two-point swap — but its full four-element ordering broke in the other pre-named direction: the \$150,000-cap variant landed at rank 2, the deployed ordering running elimination, cap, two-point, one-point against the required elimination, two-point, one-point, cap. That is the roughly 10 percent tail the forecast had assigned to the fingerprint’s non-reversal, materialising exactly where the second amendment’s series forecast had pre-named it.

The third registered forensics round asked two questions, both reported, both at bit-identity 0.0. On the cap adjacency, the entailment holds: the deployed frame’s above-cap payroll share owns 0.9925 of the gap between Smith’s one-year exhaustion delay and the deployed 16.7 years, configuration and vintage the residual 0.0075. The ledger arithmetic makes the inconsistency exact — the elimination-over-two-point swap requires a compression ratio above 0.1613, the cap-over-one-point leg above 0.0806, and Smith’s full ordering survives only inside a narrow window that any frame compressed enough to realise the certified swap overshoots, the deployed frame by roughly tenfold. The four-element ordering is therefore internally inconsistent as a transport target — the compression the fingerprint certifies mechanically busts the cap rank — and contract-permitted lever restoration enumerated empty. A pair-scoped fingerprint, the elimination-versus-two-point adjacent swap alone, is anchor-supported — Smith orders elimination’s savings above the two-point rise — and had realised three of three across the candidates. On the household-size residual, the round refuted its own registration and owned the miss. The coresidence channel is a near-exact size-1-to-size-3-plus mirror — -0.047 against +0.046, a ratio of 0.98 — and owns size-3 at 0.75, but fertility owns size-4 at 0.67 and size-5-plus at 0.78, and size-1 is mixed; the registered expectation that coresidence would dominate more than half of every failing cell was refuted, and the mixed-but-structured branch priced near 0.2 was the one that held. The levers are exhausted: a roster-ceiling probe pushed the seed past its permitted maximum and reached a terminal coresidence of 0.111, capped by the certified parental-exit hazard rather than the seed, and no untried permitted entry-state lever exists — the residual requires a household model outside the entry state. The round’s own draft summary, that coresidence dominates the quad, was refuted by its own decomposition before commit; the finding was corrected to the data, a tolerance-unit bug fixed, the measurement re-run, and the committed artifact carries the attribution as measured.

The third amendment ran the same five-round ceremony, and the ceremony again caught a false universal on its first round: the proposal had asserted the impossibility universally — no frame can both realise the swap and reproduce Smith’s cap rank — which the forensics round’s own compression-share-conditional analysis refutes, a frame at a share near 0.22 realising the swap and holding the full four-element order. The referee required the claim scoped to the pinned certified frame, where the deployed compression measures the cap at rank 2 bit-identically across all three candidates, and required the section to lead with its two dispositive grounds — the measured deployed-frame rank, and the absence of any published representative-frame anchor for the four-element order; verification extended the mutation guard to the whole earnings-distribution synonym class. The ratified amendment demoted the four household-size cells to report-only under one machine reason — no permitted entry-state lever reaches the cell — kept size-2 gated on the third candidate’s five of five, and re-scoped the gated fingerprint to the pair alone, its cap and one-point legs publishing report-only. The roll-up moved to 44 gated and 88 report-only, the family-A operating characteristic to 0.9403 per seed and 0.9684 at four of five on 43 cells. This put the whittling question in its strongest form — the 44-cell surface was now known passable by the already-committed third-candidate model — and the amendment carried the burden head-on: every demoted element was forensics-proven unreachable by any permitted lever, every one had been pre-named in the second amendment’s ratified series forecast before the third candidate ran, and all three prior candidates stood failed under the no-self-rescue rule. The demotion boundary is the permitted-lever line applied prospectively, not a fit to the model that would pass.

The fourth candidate was the third-candidate model with zero code changes — byte-identical model and runner — re-registered on the 44-cell surface: a deliberate reproduction-class run on the same pinned frame, the same registered streams, the same seeds. The single pre-registered exception was disclosed and continuous-integration-verified before execution: the family-C verdict adapter had to read the pair-scope rule from the live contract, its statistics untouched, the full deployed ordering still computed and published. It passed — five of five seeds across all 43 family-A cells in-band, the elimination-versus-two-point swap realised — at a registered price near 0.90. The standing addendum’s independent adversarial verification then reproduced the entire run bit-exactly in a fresh worktree: 4,300 of 4,300 cube values and 1,505 of 1,505 per-cell fields equal, only timing and path metadata differing, the frame re-export reproducing its checksum exactly. The pre-registered verdict adapter was needed and decisive — the old four-element rule returned false, the ratified pair rule true, both statistics published. Every registered forecast and forensic call in the four-candidate search is graded on the public record.

The pass certifies the live certification scope and nothing more: the transport of the already-fit PSID generators onto the certified populace frame — that stochastic regeneration through the deployed generators reproduces the CPS-observable cross-section within the frame’s own sampling floor on the 44 gated cells, the 43 family-A joints and the pair-scoped compression swap. It does not certify the four household-size composition shares or the cap and one-point ordering legs, report-only pending a household-composition model outside the entry state and the anchor-ordering re-specification; the 18-to-24 participation pair, the 65-plus marital and coresident quad, or the price-indexing fingerprint, report-only pending concept and cohort-vintage bridges; any SSA administrative margin, report-only since the first amendment, benefit levels published but not certified; the re-estimation of the dynamics themselves, which stay certified on their own PSID holdouts — earnings, marital, household, marriage-by-earnings, disability; or the projection.

The surface that passed is smaller than the surface first locked — 44 gated cells from 65 — and every removal carries its machine reason, its retained tolerances, and its public evidence chain: an audited search of four scored candidates, three registered forensics rounds, and three ratified amendments, no committed verdict ever changed. The representative-frame transport — the program’s funded scientific crux, roadmap milestone M5 — is resolved. The unified person-period panel now exists as a certified object: PSID-estimated lifetime dynamics deployed on the certified populace frame, faithful to the CPS-observable cross-section within its own sampling floors.

10.10 The projection engine and its temporal-holdout gate

The milestone after the transport is the projection engine: a wave loop over the single unified person-period panel, each period a fixed order of operations — mortality, aging, couple formation and dissolution, fertility and roster, disability, earnings, claiming, household reconciliation — forming a directed acyclic graph with one lagged edge, and composing the already-certified transition objects — the marital core, the household-composition candidate injected as state, the disability module — rather than re-estimating them; the transport gate is its period zero. The design went through an adversarial design-referee round before any floor was built, and the round returned a major revision on fourteen findings. It held the temporal-holdout architecture sound and cleared the design’s two hardest premises — the strict prohibition on any post-cutoff person-year entering any fitter, and confirmation from code that the embedded marital core is literally the certified object — while requiring the conceptual gaps closed before the floors ceremony: an undisclosed second-order leak (the specifications were structurally selected on the full sample, so the holdout is out-of-sample for fitted parameters but in-sample for structure), an internal contradiction over whether the household-composition core is injected or re-certified, a truth-side weighting-and-attrition friction, and the reference-year shock-window pin. All fourteen were addressed in the next revision and independently verified before the ceremony proceeded.

The gate fits on every PSID observation dated 2014 or earlier and projects and scores the held-out

2015–2022 window, PSID-projected against PSID-realized — so, unlike the transport gate’s SSA family, it needs no concept bridge: the projected and realized quantities are the same concept. The truth-side floors ran their own ceremony. The first floor priced the design’s age-band-by-sex surface and fired the gate’s operating-characteristic-before-lock pause in both directions: the certifiable flow surface was near-vacuous — one gated flow cell — and the combined faithful pass probability sat at 0.8449 against the 0.90 bar, so that first floor is frozen as the pause evidence. The coordinator adjudicated a surface redesign — two pinned ladders, candidate-blind by construction — and the second floor gated eleven cells, four flow and seven earnings, clearing at 0.9067.

The adversarial referee then rebuilt both floor artifacts bit-identically from the PSID, swept every contiguous pooling, and returned amend-four-then-proceed. The one substantive finding was an enumeration gap: the marital ladder had never enumerated the asymmetric age-two rung — lower bands merged wide, elderly tail isolated, the shape the mortality ladder already used — and that is the single rung where remarriage clears: remarriage at 18-to-64 passes at tolerance 0.403, over 143 weaker-half events, on every one of a hundred household-disjoint seeds. The third floor added exactly that rung; its fifth gated flow cell forced one further pinned earnings prune, leaving eleven gated cells — five flow and six earnings.

The gate locked as a new top-level gate on that third floor, the surface frozen and only ever read: eleven gated cells, a per-seed faithful pass probability of 0.8934 and a combined 0.9087 — flows 0.9822, earnings 0.9626. Its headline names what it does not certify, mortality drift first and at the same prominence as what it covers. No admissible pooling of the 25-to-84 mortality surface clears the $\ln(1.5)$ tolerance cap even fully pooled — the best achievable tolerance sits near 0.472 against the 0.4055 cap — the 85-plus stratum stays attrition-confounded, and pools that include it clear on power only by diluting the confound; so zero mortality and widowhood cells are gated, and any downstream trust-fund seeding rides an SSA/NCHS report-only mortality anchor, never certified drift. The 2020–2022 shock window is held outside the model class and partitioned out of every gated set; end-window stocks are report-only by the trivially-passable-stock problem; and the gated earnings cells ride a forward earnings law first certified here — the gate-1 certificate does not transfer to it.

The engine build reached three blockers, and each became a pinned design amendment rather than an improvised choice. First, the backward-chain discovery: the certified gate-1 earnings generator is a backward biennial imputation chain — it anchors each person’s chronologically latest real earnings and draws earlier ranks conditional on later ones — while the projection starts from realized 2014 and needs earnings forward, and inverting the conditional, extrapolating a future marginal, or splitting the biennial step into an annual kernel would each be a different stochastic law from the one certified. The build stopped, and the amendment specifies a forward conditional-rank chain fit from scratch on the pre-2014 data, mirroring the certified conditioning structure reversed in time. Second, the rank-to-level law: a pre-2014 fit has no 2016 age-bin cell to map a drawn rank onto a positive earnings level, and inventing one is prohibited; the amendment pins a wage-normalized, calendar-invariant cell-marginal quantile re-indexed by a wage index projected from the 2005–2014 trend — realized post-cutoff wage growth fenced off the scored path as leakage — plus the exact inverse re-ranking law, so a rank-to-level-to-rank round trip is closed. Third, the earnings-domain law: the closed-panel universe includes people the forward generator has no 2014 state for, and its initializer requires that state for every person; the amendment fixes the domain to the 2014-anchored realized-earnings closed panel and marks everyone outside it an earnings open-addition, report-only. The domain check found that of the third floor’s 13,163 gated-earnings support persons, 2,722 — about 21 percent — were later entrants outside that domain: the frozen floor had over-included open-additions, a conservative bias rather than a leak.

The engine’s first registered candidate stopped before scoring, by design, and the record grades it as such. The run was registered as a one-shot temporal-holdout projection, but the engine was build-only — the per-slice native-panel builders existed only as test stubs, the drift-scoring composition was

unbuilt, and the real-data pre-flight had run only on synthetic ensembles — so the run lane stopped before any scored phase rather than reconstruct the projected panel’s history schema inside a scoring lane, which would be unregistered modeling. No scored phase ran, no artifact was written, and the registration is graded a registration error on the public record. The unblock is a reviewed deliverable, not runner glue: a design amendment pins the scored-run harness with every modeling decision fixed and candidate-blind — its builders and scorer exercised on synthetic frames only until the registered run — and the harness build, comprising the year-0 realized slice, the projected-slice-to-native-panel builders, a drift-scoring layer that reuses the frozen floor’s cell functions verbatim under a byte-identity self-test, and the real-data pre-flight, merged after the adversarial referee returned merge-ready across its sixteen probes — no build-versus-spec deviation, the self-test proven load-bearing by a mutation run that failed loudly. A pre-flight sign-path correction followed — a design amendment and its alignment patch, both merged. The fresh registration then stopped before scoring a second time, by design, one layer deeper than the first: the harness and its guards all passed — the registration accepted, the stale identifier rejected, the frozen floor byte-verified, the eleven-cell registry and the four-of-five protocol resolved — but the runner is deliberately not self-starting, requiring a certified external-reference binding dated 2014 or earlier that does not yet exist, and the guard correctly refused the only committed claiming reference as a post-cutoff 2023 vintage rather than bind it; so again nothing scored, no artifact written, and a further design amendment is in flight to supply those bindings.

10.11 The candidate ladder against the projection gate

The bindings amendment landed, the seventh registration cleared, and the first candidate scored: FAIL, six of eleven cells, published under the gate’s publishes-regardless rule with every per-seed, per-cell margin in the committed artifact. The failures were informative in the way the gate was designed to make them: the earnings process carried too much conditional-rank memory — lag-2 autocorrelation and rank mobility both far outside tolerance — and the marital transport gaps concentrated in first marriage and remarriage. The candidate-2 program answered with two train-only deltas, each selected and frozen before any scored evidence existed: a stable-coordinate earnings rank refresh drawn per transition at a share q chosen by a twenty-one-rung train-only ladder under a one-standard-error smallest- q rule ($q^* = 0.55$), and a support-aware first-marriage estimator with its regularization chosen by the same discipline. The program also had to say aloud what it expected: a mandated candid forecast, built on a boundary-transport measurement, predicted the run would fail with remarriage the failing cell, and the registration proceeded anyway because the run’s registered purpose was scored evidence, not a predicted pass.

Getting candidate 2 to a verdict consumed four registrations and six invocations, and the record grades every one. Two aborted on environment guards that exist to refuse an unregistered parameter surface; one was rejected by a preflight whose certified premise the registered candidate itself falsified — the check assumed the two comparison arms differ only in marital provenance while the registered first-marriage swap made that false by design, so a ratified amendment rebound the reference arm to the registered law; one crashed at a serialization boundary and was fixed under pin; and the sixth ran eight and a half hours before the host machine slept and a network-dependent supervisor process died and took the runner with it — an external kill with zero information produced, adjudicated on the public record before the registration’s single disclosed re-execution was spent on a hardened topology with no supervisor in the process ancestry. The re-execution completed: FAIL, three of five seeds against a four-of-five bar, with the five-cell must-not-regress block passing on every seed — the two deltas surrendered nothing candidate 1 had won.

The verdict’s most consequential content was the failure map, not the failure. The forecast inverted: remarriage — the predicted failing cell, measured in transport at up to 193 percent of tolerance — came in at a median 73 percent across seeds, failing once, marginally; the verdict-pivotal cells were instead two earnings margins, prime-age mean growth (undershot in all five seeds) and lag-2 autocorrelation

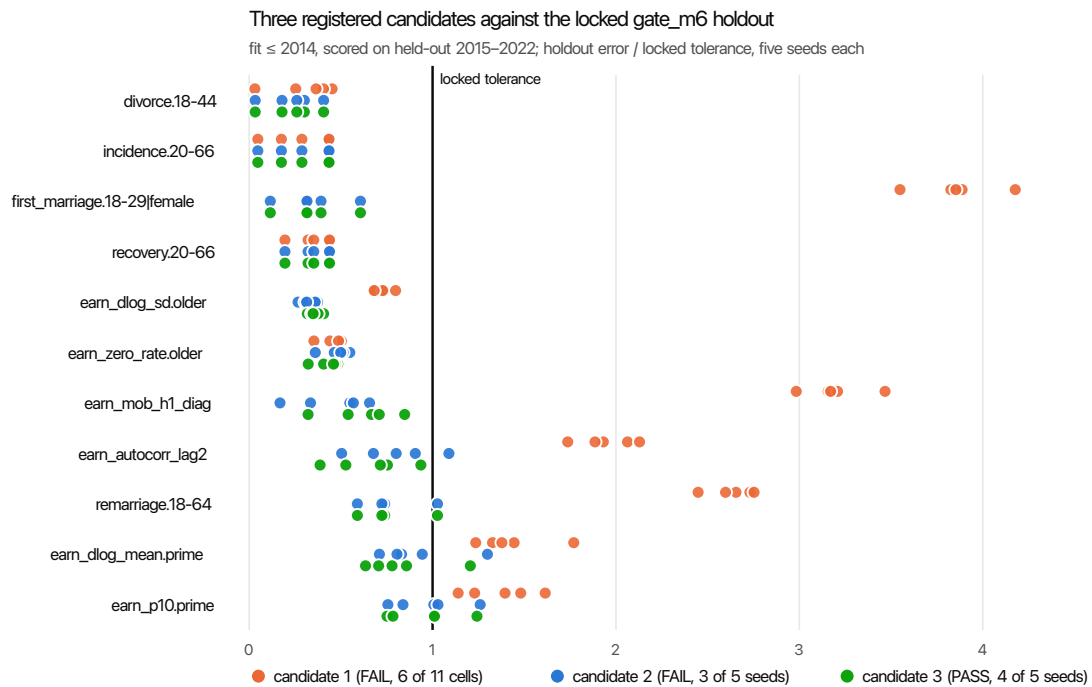


Figure 4: All three scored candidates against the locked gate: every seed of every gated cell, as holdout error over the locked tolerance. Candidate 1 fails six cells, several at two to four times tolerance; candidate 2 pulls ten of eleven cell-seed clusters inside but fails three of five seeds; candidate 3 passes — four of five seeds, with the lag-2 autocorrelation target inside tolerance in all five and only seed 2’s mean-growth and remarriage residuals over the line. All three artifacts are committed and referee-verified; the figure is rebuilt from them at build time (`scripts/build_paper_figures.py`).

(overshot in all five), each breaching in exactly one seed at thin margins. An artifact-only forensic program — every number quoted from the two committed artifacts, no fresh holdout read — established the discrepancies at two to two-and-a-half anchor-noise standard deviations, adjudicated the design directions a scored failure licenses (surface identity, never numbers, under an explicit leakage boundary), and survived four adversarial referee rounds in which each round caught a genuine record error, two of them the coordinator’s own. A train-only diagnostic then split the growth undershoot mechanically: at the deepest train boundary the same statistic shows no gap at all — the discrepancy is a transfer phenomenon across boundary depth and calendar regime, not a fitting failure — while the persistence excess is real in train, three-point-nine floor standard deviations, and completely insensitive to the wage-index path.

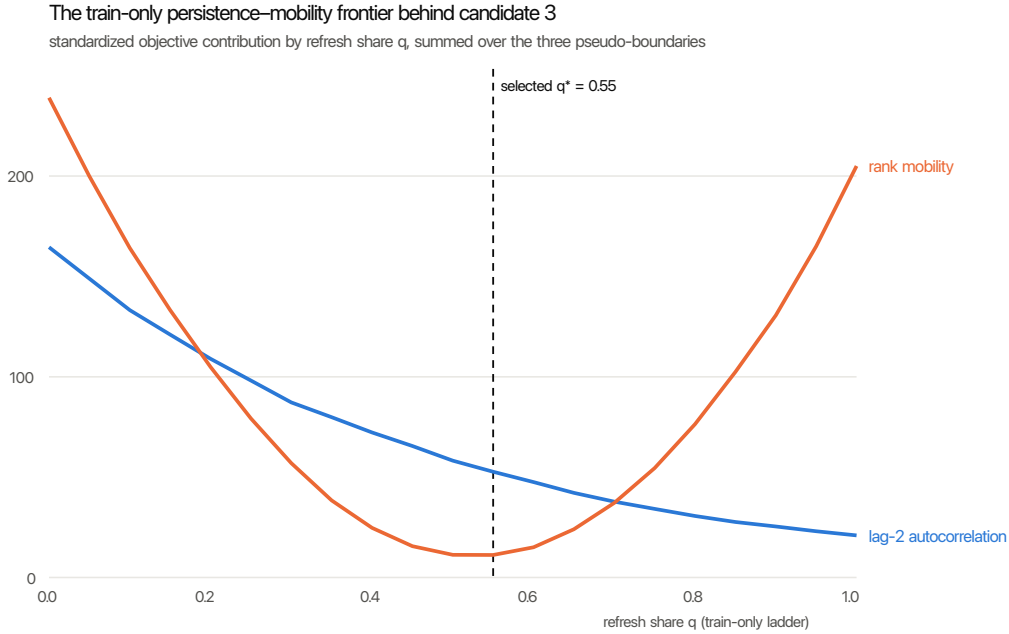


Figure 5: The train-only selection ladder behind the earnings law, summed over its three pseudo-boundaries. More refresh monotonically reduces excess lag-2 persistence and destroys rank mobility past its interior minimum; a single parameter controls both, which is the frontier the selected q^* sits on. Candidate 3’s amendment adds the missing degree of freedom — within-person correlation of consecutive refresh events — decoupling the joint no-refresh run law from the marginal rate.

That routing produced design amendment 6: the refresh events, independent per transition under candidate 2, gain a within-person first-order correlation at the frozen q^* , because under independence one parameter controls both the marginal refresh rate that rank mobility prices and the joint no-refresh run length that lag-2 persistence rides — the visible frontier of Figure 5. Negative correlation shortens unrefreshed runs at an unchanged marginal: the consecutive-no-refresh probability falls by sixty-one percent at the deepest admissible grid rung. The amendment fixes the selection protocol blind — a seventeen-rung non-positive correlation ladder re-using the frozen machinery, a closest-to-zero one-standard-error tie-break, and an explicit designed-pause outcome if zero survives.

The ladder selected a correlation of -0.60 , the tie-break rule engaging exactly as written; a second referee recomputed the selection bit-for-bit before the lock. Candidate 3 was registered under a restatement of the one-run terms, a strengthened post-2014 attestation, the diagnostic’s conditioning caveat carried verbatim, and — a house rule the candidate-2 experience earned — a transport-calibration datum pinning how far the previous registration’s boundary-transport forecast overshoot the scored window.

The registration’s candid forecast called FAIL the modal outcome while conceding, for the first time in the campaign, a materially nonzero chance of passing. The one-shot ran under a disclosed three-invocation chain — an environment-guard refusal, a host power-management incident that consumed the registration’s single re-execution allowance, then nine uninterrupted hours — and returned the campaign’s first PASS: four of five seeds against the eleven-cell contract with every must-not-regress constraint held. The routed target delivered outright — lag-2 autocorrelation came inside tolerance in all five seeds, at thirty-nine to ninety-four percent of it, where candidate 2 had breached — while the mean-growth cell cleared four of five with its residual breach shrunk, and remarriage reproduced candidate 2’s scored band almost exactly, vindicating the transport-calibration datum’s factor-of-two correction. The verdict artifact, its adversarial referee round, and the ratifying merge complete the arc: the correlated-refresh forward earnings law is first-certified on the registered 2016/2018 holdout surface, transferring no backward-law certificate and certifying nothing about mortality drift or the shock window.

The campaign also keeps a second kind of score. An append-only forecast ledger records the project’s own schedule estimates — median and eightieth-percentile dates against decidable resolution criteria, revised only by supersession with the evidence stated, graded against reality when they resolve. Its first resolved entries landed on their median dates; its first material revision, one day of median slip, carries the incident chain that caused it; and the certified-engine entry itself resolved two days inside its median when the candidate-3 verdict ratified. The ledger binds the project to the same discipline as the model: forecasts published before outcomes, misses kept on the page.

10.12 The interim benefit seam

The transport pass certifies the CPS-observable cross-section, not benefit levels — those await the transported AIME — so reform analysis still delivers benefit deltas through an interim seam, and a registered diagnostic — reported, not gated — ran it end to end. The dynamics layer computed caregiver-credit benefit deltas from a committed encoding (the Biden plan); an interim cell-mean mapping transported them onto the certified populace default file; and the managed PolicyEngine runner consumed them as a social-security perturbation, returning full current-year (2026) tax-and-benefit incidence with program interactions intact. What the run certifies is the seam, not the levels. The seam moved exactly the dollars it was handed — the engine’s change in social-security outlays matches the transport’s weighted sum to a relative gap of 6e-8 — and the interactions are enumerated rather than asserted: a gross benefit cost of \$41.35 billion nets to \$36.02 billion (0.871 of gross) after \$2.48 billion of income-tax clawback from benefit taxation and \$2.86 billion of net means-tested savings — SSI —\$1.91 billion, SNAP —\$1.13 billion, ACA premium tax credits +\$0.19 billion, TANF —\$0.01 billion, Medicaid inert at zero. Poverty moves where a caregiver credit should: the overall SPM rate falls 0.128 points, 13.86 to 13.73 percent, and the 65-plus rate 0.371, 15.67 to 15.30 — 2.9 times the overall reduction — while 56.2 percent of the aggregate net-income gain reaches the bottom half of the household income distribution. The bottom-quintile share, 20.9 percent, sits below the caregiver anchor’s 52-to-62-percent lifetime-earnings band — the dilution the registration anticipated for a recipient-only transport read on household income deciles rather than lifetime earnings. All four pre-registered expectations held: net cost inside the 75-to-92-percent-of-gross band, SNAP spending falls, the poverty reduction concentrates at 65-plus, and the decile gains concentrate in the bottom half.

The seam also states the program’s division of labor. PolicyEngine executes: modeled social-security deltas enter the current-year Populace file so the existing machinery calculates downstream incidence — taxation of benefits, means-tested interactions, marginal tax rates — unchanged. Axiom cross-validates the statute: the AIME-and-primary-insurance-amount chain is already cross-validated against the Axiom engine at 240 of 240 careers exact to the cent, and the auxiliary spousal and survivor rules are next to encode and cross-validate against this project’s frozen oracle.

10.13 The capability roadmap

These gates sit on a published roadmap of eight capability milestones that carries the program from replicating DYNASIM’s incidence patterns to a projection in its class. The layer certified or anchor-validated today — earnings, marital and fertility dynamics, household and relationship composition, the marriage-by-earnings joint, mortality, claiming, and the statutory formula with its auxiliary benefits — is the first milestone; the same-frame pseudo-projection and the disability module — now behind its own passed gate — are two more; the representative-frame transport — the funded crux — is resolved: its gate passed at the fourth registered candidate, verified bit-exactly, on a surface that three registered forensics rounds and three ratified amendments narrowed from sixty-five gated cells to forty-four, every removal machine-reasoned and no committed verdict changed; and the interim benefit seam above continues to deliver reform incidence through the production calculator. Ahead lie the year-by-year projection engine — its temporal-holdout drift gate already locked, eleven gated cells with the headline not-certified naming mortality drift first, its scored-run harness merged after its referee round, and its candidate twice a designed stop before scoring (the second on the missing 2014-or-earlier external-reference bindings, a further amendment in flight), no scored projection yet run — full revenue and trust-fund accounting, and per-year execution of the statutory rules on the projected panel. The architectural bet under all of it is a single object: one person-period panel — the cross-sectional population extended through time, with future births and immigrants entering as new persons — of which every product, from a one-year tax-benefit microsimulation to a local-area cut to the seventy-five-year projection, is a row, column, or time slice, rather than the separately funded and mutually unreconcilable models a projection of this kind has historically required.

11 Related work

Dynamic microsimulation originates with Orcutt et al. (1961); Li and O’Donoghue (2013) survey the field’s alignment and calibration practice. This design borrows the central structural lesson of DYNASIM (Favreault et al. 2015; Urban Institute 2024), MINT (Social Security Administration 2024), and CBOLT (Congressional Budget Office 2018) — an annual state engine with family links and external alignment — while departing on openness, trajectory-level weighting, and scoring; Dekkers and Cumpston (2012) treat the weighting question directly. The Cato model (Chanwong 2026) is the nearest open system.

The pattern repeats internationally — Pensim2 at the United Kingdom’s Department for Work and Pensions, MOSART at Statistics Norway, MIDAS at Belgium’s Federal Planning Bureau (Li and O’Donoghue 2013) — with exceptions. INSEE has published the source of Destinie 2 (Blanchet et al. 2010), the pension model behind France’s official projection exercises (github.com/InseeFr/Destinie-2). SimPaths, from the University of Essex’s Centre for Microsimulation and Policy Analysis, is an open-source life-course model estimated for the United Kingdom and being adapted to several other European countries (Bronka et al. 2025). And open frameworks exist without open populations: LIAM2 — built at the same Federal Planning Bureau that runs MIDAS (Menten et al. 2014) — and OpenM++, an open reimplement of Statistics Canada’s Modgen platform (OpenM++ development team 2026), supply generic simulation engines but ship no calibrated population and no rules stack. On OpenM++, WIFO’s microWELT models welfare transfers comparatively across Austria, Spain, Finland, and the United Kingdom, with a United States variant for labor-force projection (Spielauer et al. 2020) — one open engine carrying a multi-country dynamic model. None of these combines an open codebase with a certified calibrated microdata baseline and a public scoring protocol under which claims resolve. SimPaths is closest on openness; the difference here is inheritance — a certified cross-sectional population, one rules platform across countries, and credibility staked on resolved scores rather than publication.

On the measurement side, administrative-data studies of earnings dynamics discipline the earnings

process: lifecycle moments (Guvenen et al. 2021), long-run mobility (Kopczuk et al. 2010), non-stationary volatility (Sabelhaus and Song 2010), and the attenuating link between current and lifetime earnings (Haider and Solon 2006) — with nonlinear panel frameworks (Arellano et al. 2017) the academic cousin of the quantile machinery used here. The quasi-experimental record on claiming responses (Mastrobuoni 2009; Behaghel and Blau 2012) anchors the behavioral scenario library, and the assumption-failure record compiled by successive Technical Panels (Technical Panel on Assumptions and Methods 2023) motivates the domains-of-validity framework.

12 Status and roadmap

The cross-sectional foundation is in production. Populace’s charter specifies the longitudinal kernel rules. This paper is the project’s front door; the supplementary design appendices — operational chapters on earnings-history construction, family and auxiliary benefits, disability and claiming, mortality and projection drift, calibration targets, the Social Security validation program, and a source-based DYNASIM dossier — live in the project repository at github.com/PolicyEngine/populace-dynamics. Development proceeds through stage gates defined as score thresholds — earnings-history credibility first, family and benefit outputs second, forward projection third, productization last — and every gate’s evidence publishes whether it passes or fails; the first three stages have now passed their gates. The third — the forward-projection temporal holdout — certified at this revision on its third registered candidate, with two failures on the record first: six of eleven cells, then three of five seeds with every prior win held, then a pass at four of five seeds through a ratified correlated-refresh amendment whose routed target cleared in all five. Certification unlocks the next stage: the first end-to-end Social Security benefit and revenue aggregates on projected histories, published with their gaps disclosed.

The layer is open source under the MIT license. The project invites corrections, particularly to the benchmark characterizations, and contributions merge on the same standard as the authors’: improve the held-out score.

References

- Arellano, Manuel, Richard Blundell, and Stéphane Bonhomme. 2017. “Earnings and Consumption Dynamics: A Nonlinear Panel Data Framework.” *Econometrica* 85 (3): 693–734.
- Behaghel, Luc, and David M Blau. 2012. “Framing Social Security Reform: Behavioral Responses to Changes in the Full Retirement Age.” *American Economic Journal: Economic Policy* 4 (4): 41–67.
- Blanchet, Didier, Sophie Buffeteau, Emmanuelle Crenner, and Sylvie Le Minez. 2010. *The New Destinée 2 Microsimulation Model: Main Characteristics and Illustrative Results*. Document de Travail Nos. G2010-13. INSEE.
- Board of Trustees, Federal Old-Age and Survivors Insurance and Federal Disability Insurance Trust Funds. 2025. *The 2025 Annual Report of the Board of Trustees of the Federal Old-Age and Survivors Insurance and Federal Disability Insurance Trust Funds*. Social Security Administration. <https://www.ssa.gov/oact/TR/2025/tr2025.pdf>.
- Bronka, Patryk, Justin van de Ven, Daniel Kopasker, S. Vittal Katikireddi, and Matteo Richiardi. 2025. “SimPaths: An Open-Source Microsimulation Model for Life Course Analysis.” *International Journal of Microsimulation* 18 (1): 95–133. <https://doi.org/10.34196/ijm.00318>.
- Chanwong, Krit. 2026. *Social Security Cato Model*. https://github.com/kchanwong/social_security_cato_model.

- Congressional Budget Office. 2018. *An Overview of CBOLT: The Congressional Budget Office Long-Term Model*. Congressional Budget Office. <https://www.cbo.gov/publication/53667>.
- Congressional Budget Office. 2024. *CBO’s 2024 Long-Term Projections for Social Security*. Congressional Budget Office. <https://www.cbo.gov/publication/60392>.
- Dekkers, Gijs, and Richard Cumpston. 2012. “On Weights in Dynamic-Ageing Microsimulation Models.” *International Journal of Microsimulation* 5 (2): 59–65.
- Favreault, Melissa M, Karen E Smith, and Richard W Johnson. 2015. *The Dynamic Simulation of Income Model (DYNASIM): An Overview*. The Urban Institute. <https://www.urban.org/research/publication/dynamic-simulation-income-model-dynasim>.
- Favreault, Melissa M., and C. Eugene Steuerle. 2007. *Social Security Spouse and Survivor Benefits for the Modern Family*. Retirement Project Discussion Paper Nos. 07-01. The Urban Institute.
- Favreault, Melissa, and Karen E Smith. 2016. *The Accuracy of MINT Wealth Projections*. Urban Institute.
- Greenwood, Jeremy, Nezih Guner, Georgi Kocharkov, and Cezar Santos. 2014. “Marry Your Like: Assortative Mating and Income Inequality.” *American Economic Review: Papers & Proceedings* 104 (5): 348–53.
- Guenen, Fatih, Fatih Karahan, Serdar Ozkan, and Jae Song. 2021. “What Do Data on Millions of U.S. Workers Reveal about Lifecycle Earnings Dynamics?” *Econometrica* 89 (5): 2303–39.
- Haider, Steven, and Gary Solon. 2006. “Life Cycle Variation in the Association Between Current and Lifetime Earnings.” *American Economic Review* 96 (4): 1308–20.
- Institute on Taxation and Economic Policy. 2025. *ITEP Tax Microsimulation Model Overview*. <https://itep.org/itep-tax-model>.
- Kopczuk, Wojciech, Emmanuel Saez, and Jae Song. 2010. “Earnings Inequality and Mobility in the United States: Evidence from Social Security Data Since 1937.” *The Quarterly Journal of Economics* 125 (1): 91–128.
- Li, Jinjing, and Cathal O’Donoghue. 2013. “Alignment and Calibration of a Dynamic Microsimulation Model.” *Journal of Artificial Societies and Social Simulation* 16 (3): 1–15.
- Look, Spencer U., and Jack VanDerhei. 2024. *Beyond the Retirement Crisis Headlines: Why Employer-Sponsored Plans Are the Key to Retirement Adequacy for Today’s Workers*. Morningstar Center for Retirement & Policy Studies. https://www.morningstar.com/content/cs-assets/v3/assets/blt9415ea4cc4157833/bltd4bb26598046aed4/66a1535de91a178e5c15872a/Introducing_the_Morningstar_Model_of_US_Retirement_Outcomes_-_July_2024_-_final.pdf.
- Mastrobuoni, Giovanni. 2009. “Labor Supply Effects of the Recent Social Security Benefit Cuts: Empirical Estimates Using Cohort Discontinuities.” *Journal of Public Economics* 93 (11-12): 1224–33.
- Meinshausen, Nicolai. 2006. “Quantile Regression Forests.” *Journal of Machine Learning Research* 7: 983–99. <https://jmlr.org/papers/v7/meinshausen06a.html>.

- Menten, Gaëtan de, Gijs Dekkers, Geert Bryon, Philippe Liégeois, and Cathal O'Donoghue. 2014. "LIAM2: A New Open Source Development Tool for Discrete-Time Dynamic Microsimulation Models." *Journal of Artificial Societies and Social Simulation* 17 (3): 9. <https://doi.org/10.18564/jasss.2574>.
- Mermin, Gordon B. T. 2005. *The Effect of Benefit Reductions on the Distribution of Social Security Benefits*. Report No. 411260. The Urban Institute.
- OpenM++ development team. 2026. *OpenM++: Open Source Microsimulation Platform*. <https://openmpp.org>.
- Orcutt, Guy H, Martin Greenberger, John Korbel, and Alice M Rivlin. 1961. "Simulation of Economic Systems." *The American Economic Review* 51 (5): 893–907.
- Penn Wharton Budget Model. 2025. *Penn Wharton Budget Model: Microsimulation*. <https://budgetmodel.wharton.upenn.edu/model/microsimulation/>.
- Policy Simulation Library. 2026. *Tax-Calculator*. <https://taxcalc.pslmodels.org/>.
- PolicyEngine. 2026. *PolicyEngine: Open-Source Tax-Benefit Microsimulation*. <https://policyengine.org>.
- Sabelhaus, John, and Jae Song. 2010. "The Great Moderation in Micro Labor Earnings." *Journal of Monetary Economics* 57 (4): 391–403.
- Schwartz, Christine R. 2010. "Earnings Inequality and the Changing Association Between Spouses' Earnings." *American Journal of Sociology* 115 (5): 1524–57. <https://doi.org/10.1086/651373>.
- Smith, Karen E. 2015. *Can Social Security Be Solvent?* Program on Retirement Policy Report No. 72196. The Urban Institute.
- Smith, Karen E., Richard W. Johnson, and Melissa M. Favreault. 2020. *Five Democratic Approaches to Social Security Reform: Estimated Impact of Plans by 2020 Presidential Candidates*. Report No. 103050. The Urban Institute.
- Social Security Administration. 2024. *Projection Methodology: Modeling Income in the Near Term, Version 8 (MINT8)*. <https://www.ssa.gov/policy/docs/projections/methodology.html>.
- Spielauer, Martin, Thomas Horvath, and Marian Fink. 2020. *microWELT: A Dynamic Microsimulation Model for the Study of Welfare Transfer Flows in Ageing Societies from a Comparative Welfare State Perspective*. WIFO Working Papers 609/2020. Austrian Institute of Economic Research (WIFO).
- Tax Foundation. 2025. *The Tax Foundation's Taxes and Growth Model*. <https://taxfoundation.org/research/all/federal/overview-tax-foundations-taxes-growth-model/>.
- Tax Policy Center. 2025. *Microsimulation Model FAQ*. <https://taxpolicycenter.org/resources/tax-model-resources/tpcs-microsimulation-model-faq>.
- Technical Panel on Assumptions and Methods. 2023. *2023 Technical Panel on Assumptions and Methods: Report to the Social Security Advisory Board*. Social Security Advisory Board. <https://www.ssa.gov/policy/docs/assumptions/>.

[//www.ssab.gov](http://www.ssab.gov).

The Budget Lab at Yale. 2026. *Tax-Simulator: Microsimulation Model of US Federal Tax System*.
<https://github.com/Budget-Lab-Yale/Tax-Simulator>.

Urban Institute. 2024. *Urban's Dynamic Simulation of Income Model 4 (DYNASIM4)*. <https://www.urban.org/research/publication/urbans-dynamic-simulation-income-model-4>.