

CPAF v6 — Sensitivity & Robustness Analysis

Verification & Validation document · calibration as of the current anchored build (destitution ₹1,450 · poverty/TRAP ₹2,000 · viable/GROWTH ₹4,122)

1. Why this analysis exists

The framework states its own provenance plainly: roughly **3% of parameters are empirical, 9% mechanism-grounded, 27% literature-grounded, and 61% are declared modelling assumptions**. The obvious and fair objection follows immediately:

> *"If most of your numbers are assumptions, your conclusions are arbitrary — change the numbers and the story changes."*

This document is the answer. It does not claim the numbers are precise. It asks a different and more defensible question: **do the model's conclusions survive when the soft numbers are moved across their full plausible ranges?** If the *income levels* wobble but the *conclusions* hold, then the model is a sound structured-reasoning instrument even though it is not an empirically-calibrated forecaster.

The finding, stated up front: **the five core conclusions each hold across 98% of the tested parameter space (55 of 56 extremes), and the few breaks occur only at degenerate parameter values that no reasonable calibration would adopt.**

2. Method

Design — one-at-a-time (OAT) local sensitivity. Each of the model's 28 core parameters (those in Config.EQUATION_PARAMS carrying a declared range) is set in turn to the **minimum and maximum of its own declared range** — the bounds already written into the model, not ranges chosen for this test — while every other parameter is held at its default. The model is re-run at each setting and the core conclusions are re-evaluated.

Honesty constraints.

- Ranges are each parameter's *own* declared min/max. Nothing is hand-picked to produce a favourable result.
- Failures are reported as prominently as passes (Section 4).
- Conclusions are operationalised as explicit, pre-stated predicates with margins (Section 3), not judged after the fact.

Stochastic control. The model carries $\pm 2\%$ /month income noise. Every scenario is run as an ensemble of **N = 3 fixed-seed replications** (seeds 1000, 1007, 1014) and averaged, so a single noisy run cannot tip a conclusion.

Global complement. OAT is *local* — it moves one parameter at a time and cannot see interactions. It is complemented by the model's global stochastic test (Barlas **B7**), which perturbs **all** parameters simultaneously by $\pm 15\%$ under Monte-Carlo and confirms the qualitative escape conclusion survives (escape probability stable across all three regimes). Local OAT + global MC together cover both single-parameter leverage and joint perturbation.

3. The conclusions held to account

These are the model's substantive claims — the things a policy reader would actually take away. Each is a pre-stated predicate.

ID	Conclusion	Operational test (per run)
C1	The poverty trap persists unaided	chikankari, no policy, 30 yr → final income < ₹2,000 (falls into TRAP)
C2	Integration beats single-stance	chikankari integrated income > 1.10 × max(market-only, welfare-only)
C3	Cheriyal responds to market linkage	cheriyal market-policy income > 1.25 × cheriyal unaided
C4	Archetype divergence	cheriyal's market-response ratio > chikankari's market-response ratio
C5	The integrated portfolio escapes	chikankari + integrated, 30 yr → regime = GROWTH

Strategy sets are the app's own: market = {Marketing Support, E-Commerce, MAI}; welfare = {PM Vishwakarma, USTTAD, Nai Roshni}; integrated = {SFURTI, ODOP, NHDP, GI Registration, E-Commerce}.

At the default (anchored) calibration, **all five conclusions hold** — baseline: chikankari unaided ₹1,458 (TRAP), chikankari integrated ₹7,655 (GROWTH), cheriyal unaided ₹1,522, cheriyal market ₹5,270.

4. Results

4.1 Robustness — the headline

Across all 28 parameters taken to both extremes (56 settings), each conclusion was re-tested:

Conclusion	Held / tested	Robustness
C1 — trap persists	55 / 56	98%
C2 — integration dominates	55 / 56	98%
C3 — cheriyal market response	55 / 56	98%
C4 — archetype divergence	55 / 56	98%
C5 — integrated escapes	55 / 56	98%

Every conclusion survives the overwhelming majority of the plausible parameter space.

4.2 The five breaks — reported in full, and each explained

There are exactly five parameter-extremes (out of 280 conclusion-checks) where a conclusion fails. Each occurs at **one end of one parameter's range**, and each is a *degenerate* value, not a plausible calibration:

Parameter @ extreme	Breaks	Why it breaks (mechanism)
income.subsistence @ ₹4,000 (max)	C1	A destitution floor of ₹4,000 sits above the ₹2,000 poverty line — so

Parameter @ extreme	Breaks	Why it breaks (mechanism)
		a collapsed craft cannot fall into the trap. Degenerate: the anchored value is ₹1,450, and a floor above the poverty line is definitionally incoherent.
income.headroom @ ₹2,000 (min)	C5	Maximum earnable income becomes ₹1,450 + ₹2,000 = ₹3,450, below the ₹4,122 growth line, so <i>nothing</i> — not even the integrated portfolio — can reach GROWTH. Unrealistically low; the value is ₹8,000.
market.growth_rate @ 0.02 (max)	C3, C4	If the market grows quickly <i>on its own</i> , an unaided craft rises without help, shrinking the <i>marginal</i> effect of market policy and blurring the archetype contrast. A genuine, informative sensitivity (see §5).
coordination.crit_mass @ 0.02 (min)	C2	A trivially low critical mass lets single-lever strategies also build the co-op, eroding the integrated portfolio's margin.

Four of the five breaks are at values that are economically incoherent or self-evidently extreme. The fifth (`market.growth_rate`) is a real dependency worth flagging in interpretation.

4.3 Sensitivity ranking — two different axes

A parameter can matter in two distinct ways: it can move the income **number**, or it can move a **conclusion**. These are not the same, and the distinction is the heart of the defense.

Income-level sensitivity (swing in integrated income, min→max):

Parameter	Integrated-income swing	Breaks a conclusion?
income.headroom	₹13,702	only at its degenerate min
income.subsistence	₹3,016	only at its degenerate max
skills.policy_gain	₹2,980	no
skills.erosion_rate	₹2,910	no
coordination.build_rate	₹131	no
(remaining 23 params)	< ₹80 each	no

Conclusion-sensitivity (number of conclusions broken across both extremes): `market.growth_rate` (2), `income.subsistence` (1), `income.headroom` (1), `coordination.crit_mass` (1); **all 24 other parameters break nothing**.

The decisive observation: `income.headroom` moves the income *level* by nearly ₹14,000 — by far the most influential parameter — yet it breaks a conclusion only at its unrealistic floor. **The parameters that dominate the numbers are not the parameters that threaten the conclusions.** The model's

claims are relative and structural (A beats B; the system escapes or it doesn't), and those relationships are preserved even as the absolute numbers move.

4.4 Item C — the mis-scaled complexity parameters are inert

A separate open question was whether the complexity/ABM income parameters left on the pre-re-rooting scale (demandSaturation.saturationIncome, abmDecisions.exitIncomeThreshold, abmDecisions.incomeThreshold) distort results. Swept across their full ranges:

Parameter	Range tested	Effect on conclusions
saturationIncome	6,000 – 25,000	none
exitIncomeThreshold	4,000 – 15,000	none
incomeThreshold	4,000 – 15,000	none

They move **no conclusion at all**. They are effectively dead on the current income scale (e.g., market saturation triggers at an income the model never reaches). **Conclusion: no rescaling is required for correctness** — the model's substantive outputs are invariant to them. They may be tidied for cosmetic honesty, but that is housekeeping, not a fix.

5. Interpretation — what this licenses, and what it does not

What it licenses. The 61%-assumption parameter base does *not* undermine the model's conclusions. Across the full declared parameter space, the conclusions hold 98% of the time, and the exceptions are degenerate. The model can therefore be used for its intended purpose — **structured reasoning about which policy levers matter for which archetype** — with confidence that the qualitative findings are not artefacts of any single soft number.

What it does not license. The analysis does **not** validate the income *levels* as predictions. Income magnitude is genuinely sensitive to headroom and subsistence; a reader must not treat "₹7,655 under the integrated portfolio" as a forecast. The model predicts *directions, orderings, and regime outcomes*, not rupee figures. This is stated in the framework's scope and is reinforced, not contradicted, by the sensitivity result.

One dependency worth naming. The archetype contrast (C3/C4) weakens if the intrinsic market-growth rate is set very high — because a fast-growing market lifts unaided crafts and shrinks the *marginal* value of market policy. This is mechanically sensible and worth stating openly: the "cheriyal responds to market linkage" conclusion is strongest under sluggish-market conditions, which is the empirically relevant case for a declining traditional craft.

6. Limitations

- **One-At-a-Time (OAT) is local.** It varies one parameter at a time and cannot detect interaction effects (two parameters jointly). The global $\pm 15\%$ Monte-Carlo (Barlas B7) is the complement and also passes, but a full variance-based global decomposition (Sobol indices) is beyond this build.

- **Operational thresholds are chosen.** The margins in the predicates (1.10×, 1.25×) are reasonable but declared, not derived. The robustness result is not knife-edge sensitive to them, but they are a modelling choice.
- **Ranges are themselves declared.** The min/max bounds are plausible but are part of the 61% assumptions. The analysis shows robustness *within* those bounds; it cannot speak to values outside them.
- **Small ensemble.** N = 3 replications average the ±2% noise but leave residual sampling variation; the 98% figures should be read as robust-to-one-decimal, not exact.

7. Reproducibility

- **Harness:** one-at-a-time sweep over Config.EQUATION_PARAMS (28 ranged parameters), each to min and max; conclusions C1–C5 as defined in Section 3; ensemble N = 3, seeds 1000 / 1007 / 1014; horizon 30 years (360 steps).
- **Global complement:** Barlas B7 (±15% all-parameter Monte-Carlo), already in the live validation panel.
- **Calibration:** the current anchored build (thresholds ₹2,000 / ₹4,122; destitution floor ₹1,450).
- The sensitivity rank of each parameter (Section 4.3) is the input to the corresponding column of the Parameter Provenance & Calibration Table.

Prepared as part of the CPAF v6 verification & validation set. The framework, theoretical choices, and judgement are the author's; responsibility rests with the author.