

When the AI and the MSC selection panel disagreed

Description

[Tweet](#)

Most Significant Change works by selection. People write stories about changes they think have been significant; a panel reads these and chooses the most significant of all. Although it is possible to have parallel selection panels, reflecting different stakeholder perspectives, this is not common, probably primarily because of the additional logical demands involved. But now with the availability LLM AI models they can be used to provide alternative perspectives, very quickly and at negligible cost

I have a set of MSC stories where selection panel choices have been documented, separately from the MSC accounts themselves. CCDB's Bangladesh programme collected MSC stories through 1994; I currently have digital copies of 29 of these, from November and December, of which seven fall in the 'people's participation' domain. A HQ panel in Dhaka selected one of these from each month, as most significant. So I asked the [Text Cluster Analysis Lab](#) a free browser tool that scores a set of texts against criteria it proposes, clusters them, and lets you question the result in plain language to predict which two the HQ panel had chosen, and to explain its reasoning.

What follows is condensed from that session. I have kept the order in which things actually happened, because the sequence matters: the interesting part is not the prediction but what came after it.

What the Lab chose

It picked **1994_11_Mohanpur_Participation**, in which a village Forum was handed responsibility for hosting a visit by foreign donors and organised the whole thing itself – hesitating at first, calling its own meeting, then confidently walking visitors round vaccine camps, tube-wells and cattle-rearing plots. And **1994_12_ChapaiParticipation**, in which Maunjera Begum took her husband, a police officer, to court over her dower and won a Tk. 20,000 judgment.

Its reasoning was coherent. The Mohanpur story showed a transfer of organisational ownership, with a deliberative process visible from start to finish. The Chapai story showed a woman using the formal justice system against a state actor – the most institutionally durable outcome in the whole participation set. Both, it argued, showed participants taking initiative beyond what the CCDB programme had asked of them.

The other five it set aside as 'reactive' – recovery from dysfunction rather than advancement.

What the HQ panel chose

Both Manda stories. Both about loan default.

In November, six members defaulted and openly declared they would not repay; the shomiti escalated through village matbors to the Union Parishad chairman to formal arbitration, and recovered the money. In December, a member's husband refused to let her repay and abused the group; the shomiti

resolved to mortgage her property if necessary, and he capitulated.

These were two of the five the Lab had dismissed.

Reconstructing the panel’s logic

I told Claude AI what the panel had actually chosen and asked what might explain the difference – making clear I was not looking for a correction, but for the reasoning.

Its account was more interesting than a simple concession. The Manda pair, it observed, share something its own picks did not: both are about *shomiti survival under threat*. That is a thematic coherence, not a coincidence of two strong stories.

It then rebuilt the case it had missed. The November story tests participation under adversarial conditions rather than favourable ones, and documents a precise escalation ladder through local governance structures – arguably a more sophisticated demonstration of capacity than a single well-organised donor visit. The December story shows a shomiti inventing a quasi-legal enforcement mechanism of its own, and resisting male interference in a women’s organisation. Both concern the most common operational threat the programme faced, and both were resolved without staff intervention – which makes them directly useful to managers elsewhere.

Then it named the difference:

	The Lab’s logic	The panel’s logic
Evidence of participation	Proactive, forward-looking action	Participation tested under stress
Valued outcome	Institutional confidence, legal rights	Institutional resilience, shomiti survival
View of conflict	Less compelling than advancement	The truest test of capacity
Replication focus	Inspiring models of what is possible	Practical models for common threats

In summary was that it had been choosing *showcase stories* – thinking like someone selecting evidence for an external audience – while the panel had been thinking like *programme managers*, asking not “what is most impressive?” but “what shows participation working when it really matters?” It added a further thought: by elevating loan-recovery stories, the Dhaka office may have been signalling back down to field offices which outcomes the programme cared about – the MSC process working as communication in both directions.

Why this is the interesting result

Had the Lab guessed correctly, I would have learned very little. The mismatch is what produced the insight, because it forced an articulation of criteria that had never been written down. The panel’s preference for stress-tested over exemplary participation is a real evaluative position, defensible and consequential – and it surfaced only because a second reader with different priors arrived at a different answer and could be asked to explain both.

This is not the tool validating human judgement, nor outperforming it. It is a second reading whose disagreements are diagnostic.

Whose change is it?

That is what the session showed. What follows is my own reading of it, drawing on knowledge of CCDB and the NGO sector that was not in the corpus and not in the dialogue â?? and worth marking as such, for the reason the next section gives.

Both of the panelâ??s choices concern the credit system. Loans recovered, defaulters brought back in, a husbandâ??s interference overcome. A functioning credit system is the precondition for everything else the programme does, and in a field crowded with organisations offering rural credit, it is also a condition of the organisationâ??s own continuation. These are changes CCDB produced, can evidence, and needs.

The court case is not like that. Maunjera Begumâ??s shomiti membership gave her, in the storyâ??s own account, the consciousness to act â?? but the judgment against her husband is not a programme output. It cannot be counted, replicated on request, or claimed in a report without heavy qualification. CCDBâ??s contribution is real and distal.

So the two selections differ in something beyond evaluative temperament. One kind of change *reproduces the institution that produced it*. The *other outruns it* â?? it is significant on the participantâ??s own terms, and would remain significant if CCDB vanished tomorrow.

I would not read this as the panel serving its own interest against the participantsâ??. A working credit system is what members need too; the interests are largely aligned, which is exactly what makes the distinction hard to see. The question the contrast raises is narrower and more awkward: *was the MSC selection systematically favouring changes a programme can account for?* If the criteria stay tacit, attributable change may win quietly and repeatedly â?? not because anyone decided it should, but because it is the kind of change that is visible as evidence.

And what would that cost? Maunjera Begumâ??s judgment was not only a change in one household. If women in that district learned from it that a shomiti member could take a police officer to court and win, the change runs into the local legal order itself â?? the kind of wider systemic shift no programme can so easily claim. Those are precisely the changes a selection process tuned to attributable evidence would risk letting through unnoticed.

I do not think this session answers the cost question. But it is the question I came away with, and I would not have quickly arrived at it by reading the seven stories again on my own.

Date

17/08/2026

Date Created

17/08/2026

Author

admin