

Authors

Meredith Startz
Dartmouth College

Flagging Failure

IPA research affiliate Chris Blattman occasionally takes a break from his [rigorous impact evaluations of clowns](#) to send a great idea out into the blogosphere. A couple of weeks ago, he suggested that [what development organizations really need to document is not best practices, but “worst practices”](#). The idea is that sharing worst practices could prevent organizations from repeating one another’s mistakes.

IPA and other development research groups should jump on this bandwagon by doing more to publicize negative and inconclusive results. We should make sure that we talk about what *doesn’t work* in poverty alleviation for several reasons:

First, not repeating a costly mistake is good, but never making it in the first place would be even better. If a many-million dollar program didn’t work in Indonesia or Malawi, by all means let’s write it up in our worst practices and prevent someone else from implementing the same program. But research groups can potentially avert big expenditures on an intervention that is doomed to failure in the first place by publishing negative results from pilot projects. Surveys and randomized trials aren’t costless, but they’re a steal if they save us from a major aid initiative that does zilch.

Second, speaking up more about negative results and non-results would be a proactive response to the current debates over [datamining](#) and the [value of randomized trials](#). Let’s make sure that the full set of results is available, so that everyone knows if an intervention improved some outcome measures but not others, or if one part of the program worked and the other didn’t. The reports that come out of randomized trials are often pretty up-front about this already – see this paper on [sex education in Kenya](#) and keep an eye out for the results of this [remittance savings program](#) in Mexico. But to be fair, this is something that we can always articulate more clearly and publicize better.

Independent research groups are ideally situated to take the lead on crucial but un-sexy aspects of evaluation that seem to be strongly disincentivized for academics, such as publishing non-results and replicating studies. As independent researchers, we are also free to flag failure without having anything personally at stake, unlike implementation organizations, which will have to explain to their funders why they just spent five years and ten million dollars on a dud. We can and should afford to be just as focused on providing evidence about what doesn’t work as what does.

May 27, 2009