

Bayesian calibration for conference reviews

Harrison Ritz Centre for Neuroscience Studies, Queen's University

Contact: hjr@queensu.ca || Supported by C.V. Starr Fellowship (Princeton) and Connected Minds Fellowship (Queen's)

Abstract

Selecting among high-quality submissions is a challenge for universities, funders, and conferences.

A major difficulty in these evaluations is accounting for idiosyncrasies in reviewers' assessments.

Classical psychometric methods from Item Response Theory (IRT) can help mitigate these reviewer effects; however, these methods are rarely used in peer review.

Here, we leverage Bayesian IRT to reanalyze openly available reviews from CCN 2023. Our results suggest that IRT can support fairer assessment of conference papers and other high-stakes evaluations.

Methods

Dataset

Reviewers were asked to review a relatively large number of papers (7 or 9) on a 2D continuous scale. One axis reflected the 'clarity' of the paper and the other axis reflected the 'impact'. Reviewers were encouraged to use the entire scale. [527 papers, 589 reviewers, 4808 ratings]

Model

In item-response theory (IRT), attributes are assessed using latent variable inference. These models estimate (1) the assessed value, (2) the assessor bias to give higher or lower scores ("1P" models) and (3) the sensitivity of the assessor to the latent value of the assessed ("2P" models).

We fit these models using brms in R (5000 retained draws with 1000 warmup, parallelized over 16 chains). We captured clarity and impact ratings using a logit-normal model:

Each 'read' of a paper was Normally distributed with a mean ('quality') and covariance ('controversy' — disagreement between perfectly calibrated reviewers).

Reviewer generated {clarity, impact} scores by scaling the 'read' of each paper ('discrimination'), and then introducing bias ('bias') and observation noise ('noise').

'Quality' and 'controversy' were estimated for each score, and were correlated across papers (within and between scores). 'Discrimination', 'bias' and 'noise' were correlated across reviewers (within and between scores).

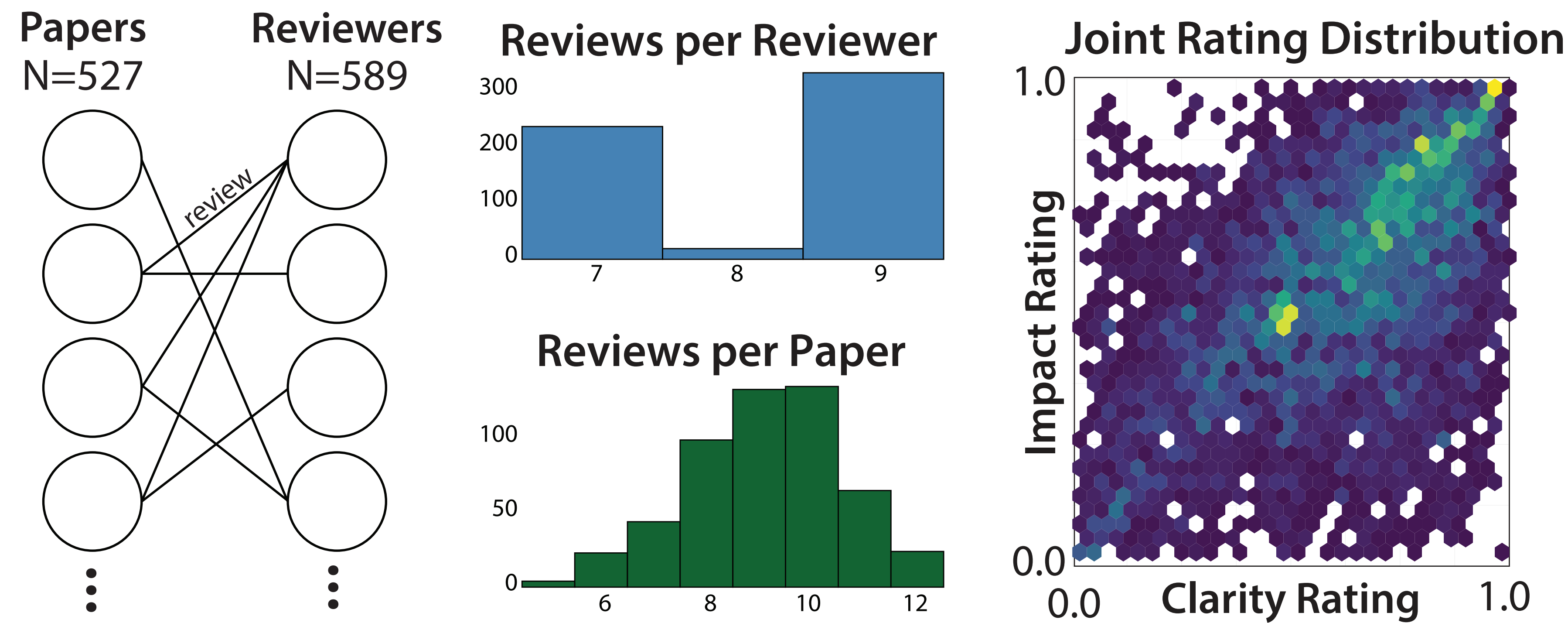
Discussion

The best-fitting model shared discrimination parameters between mean and covariance, consistent with the underlying process model.

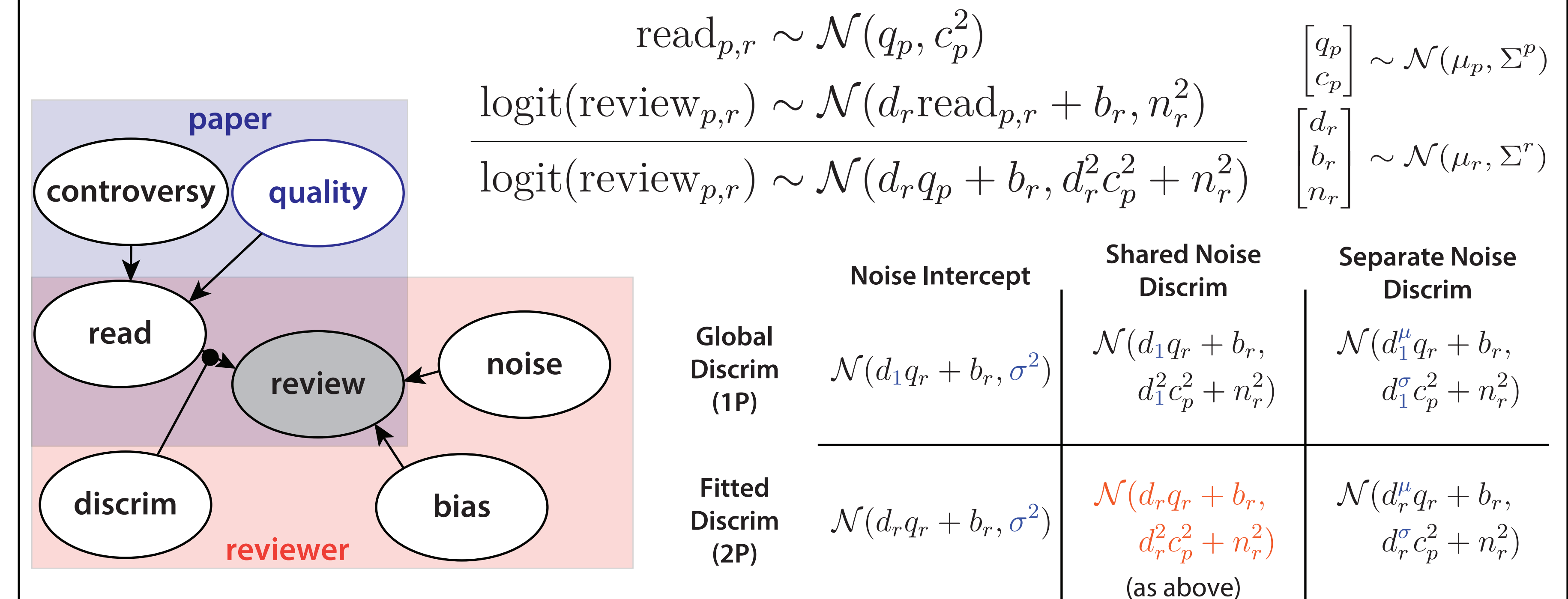
The best-fitting model had good parameter and model recovery. We could effectively capture the distribution of scores, and distribution of empirical reviewer discrimination. This model could generalize to the new OpenReview system, though potentially at lower sensitivity due to ordinal response and lower reviewer load.

The posterior estimate from this model could help guide conference and grant decisions, such as through our proposed lottery. However, practical deployment may also want to further calibrate paper selection based on factors like topic diversity and presenter demographics.

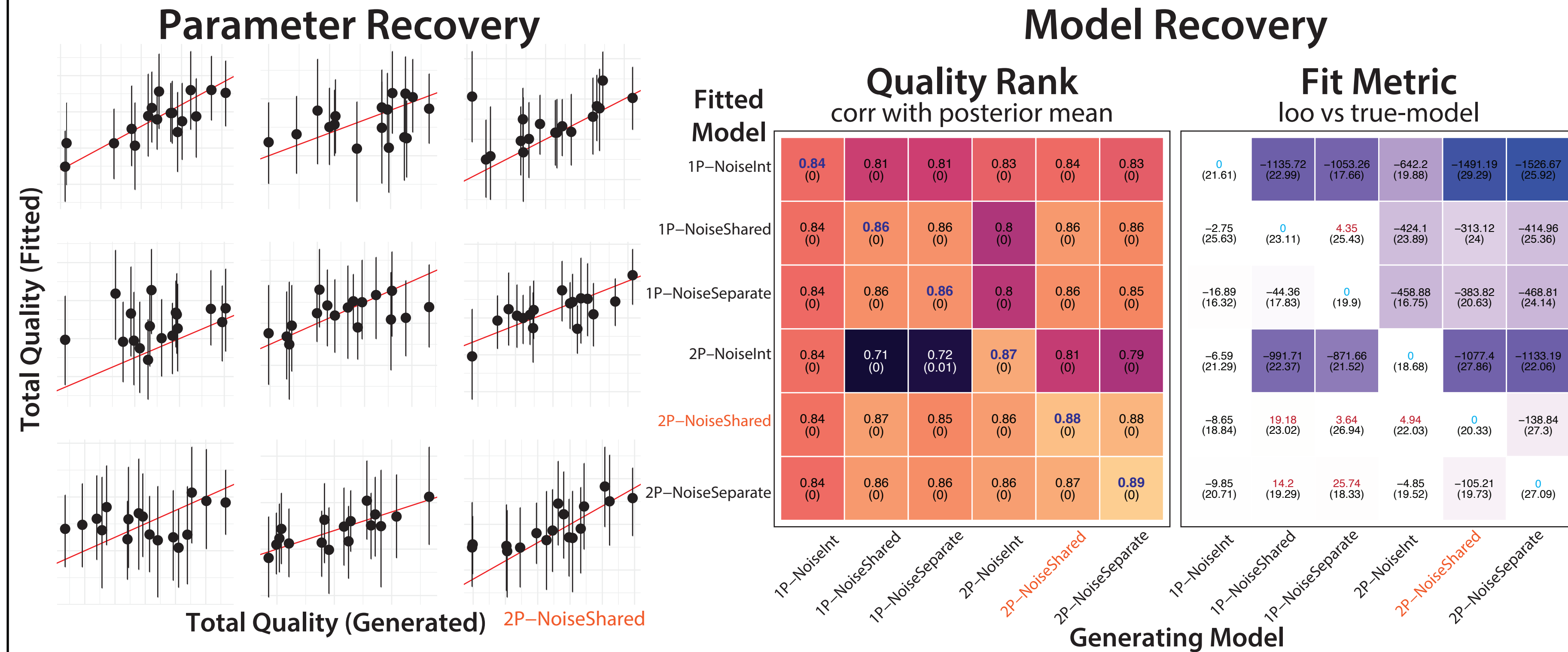
Review Database from CCN 2023



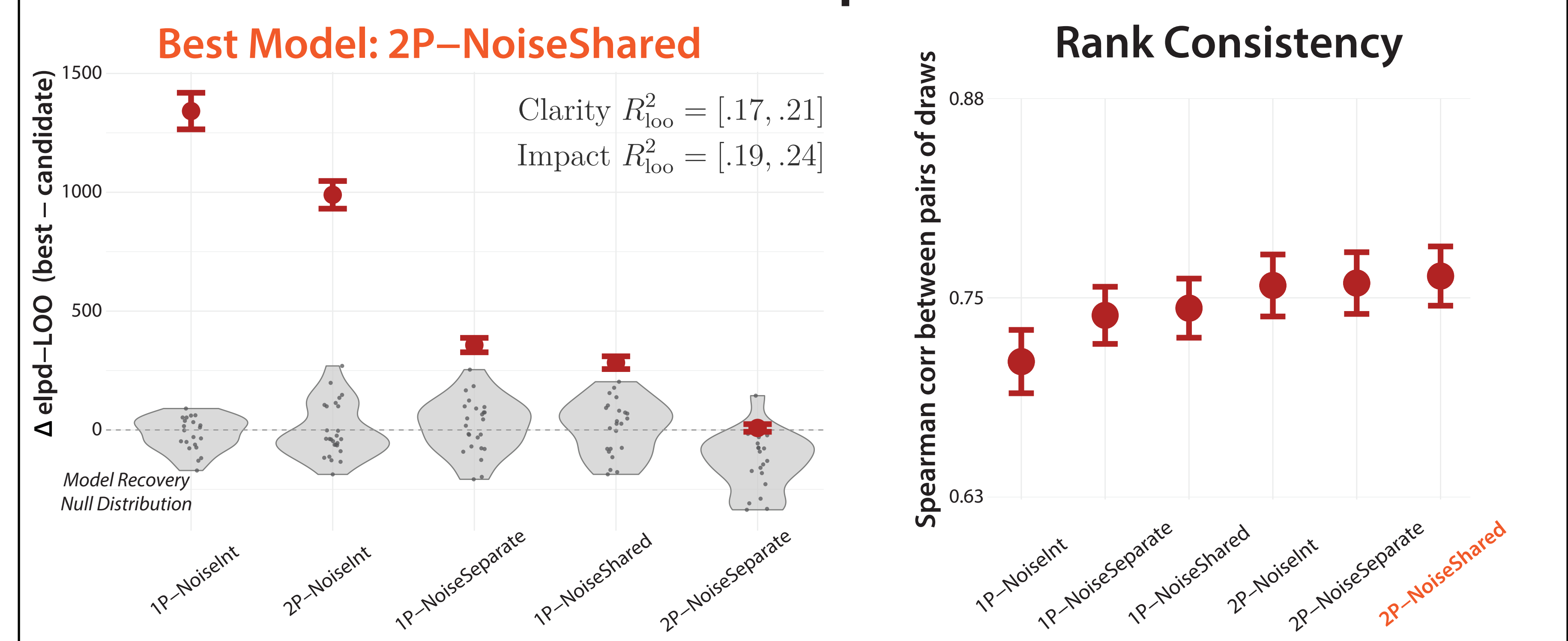
Item Response Theory Model



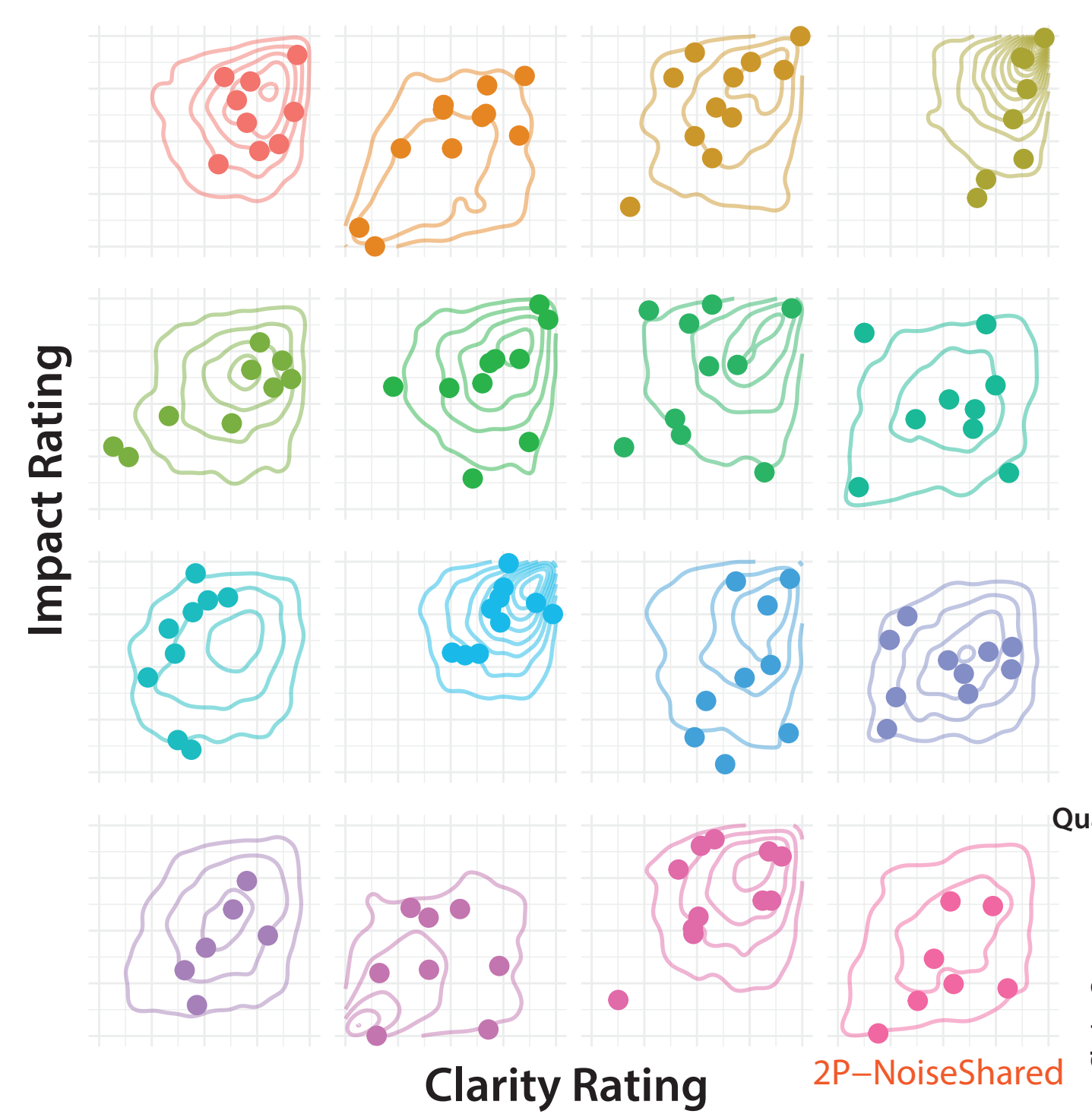
Model Validation



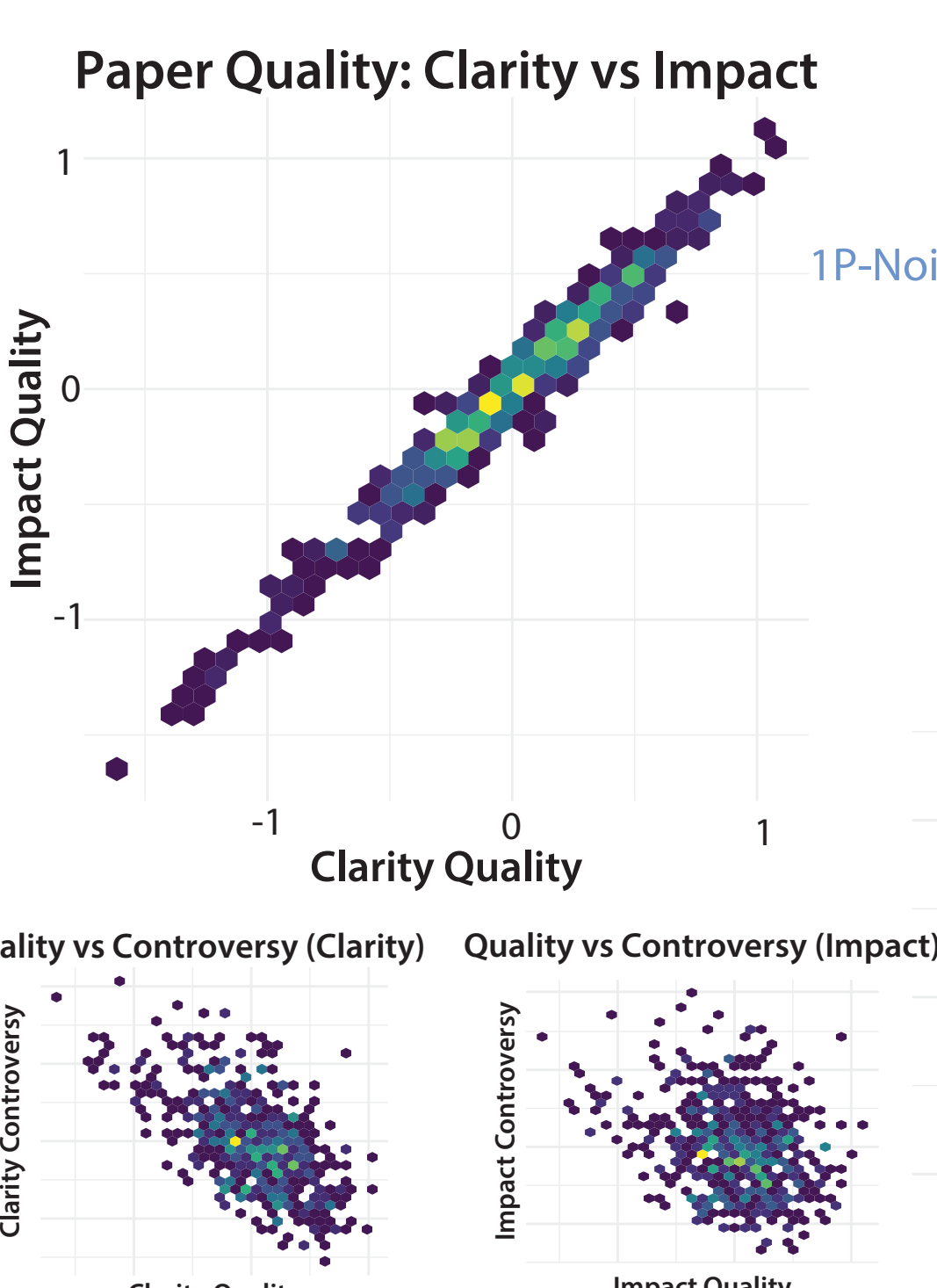
Model Comparison



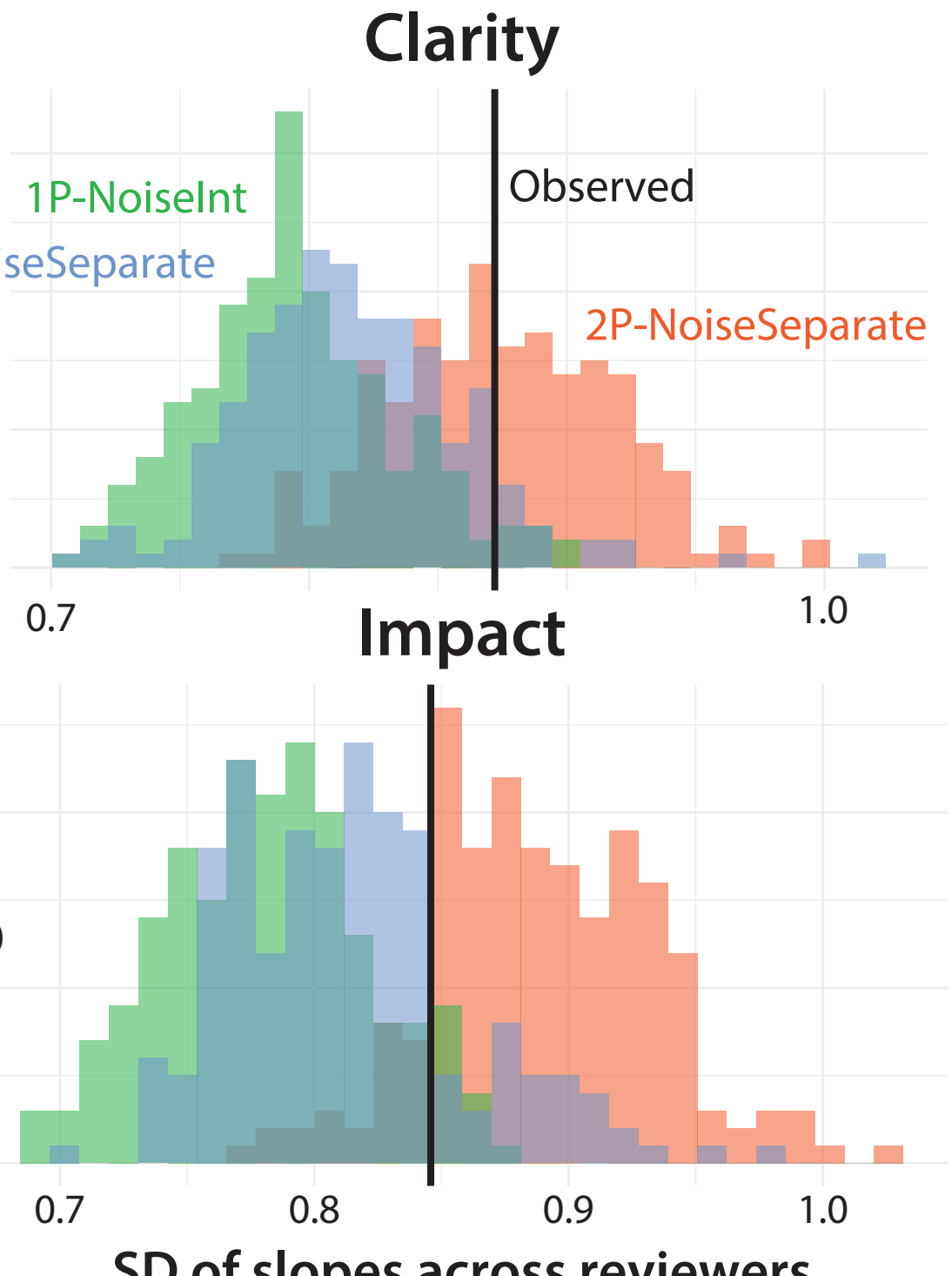
Posterior Predictives



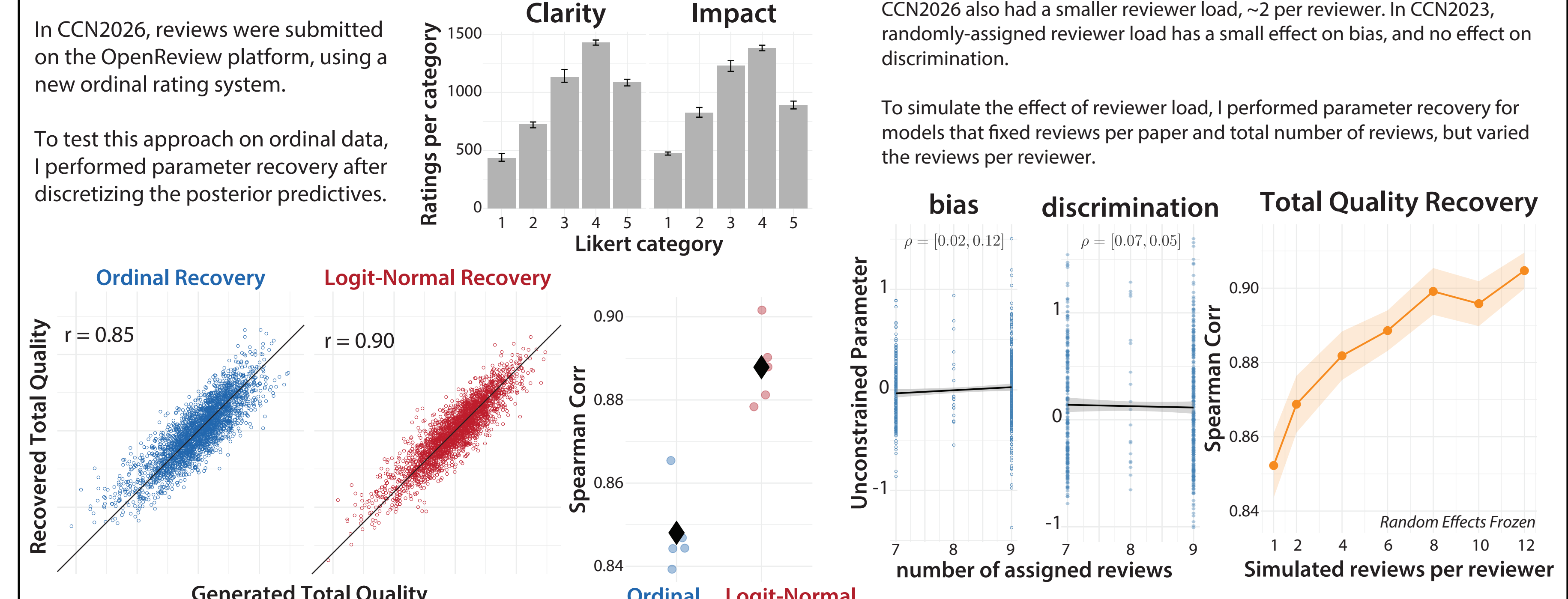
Parameter Posteriors



Discrimination Variability



Applicability to CCN 2026



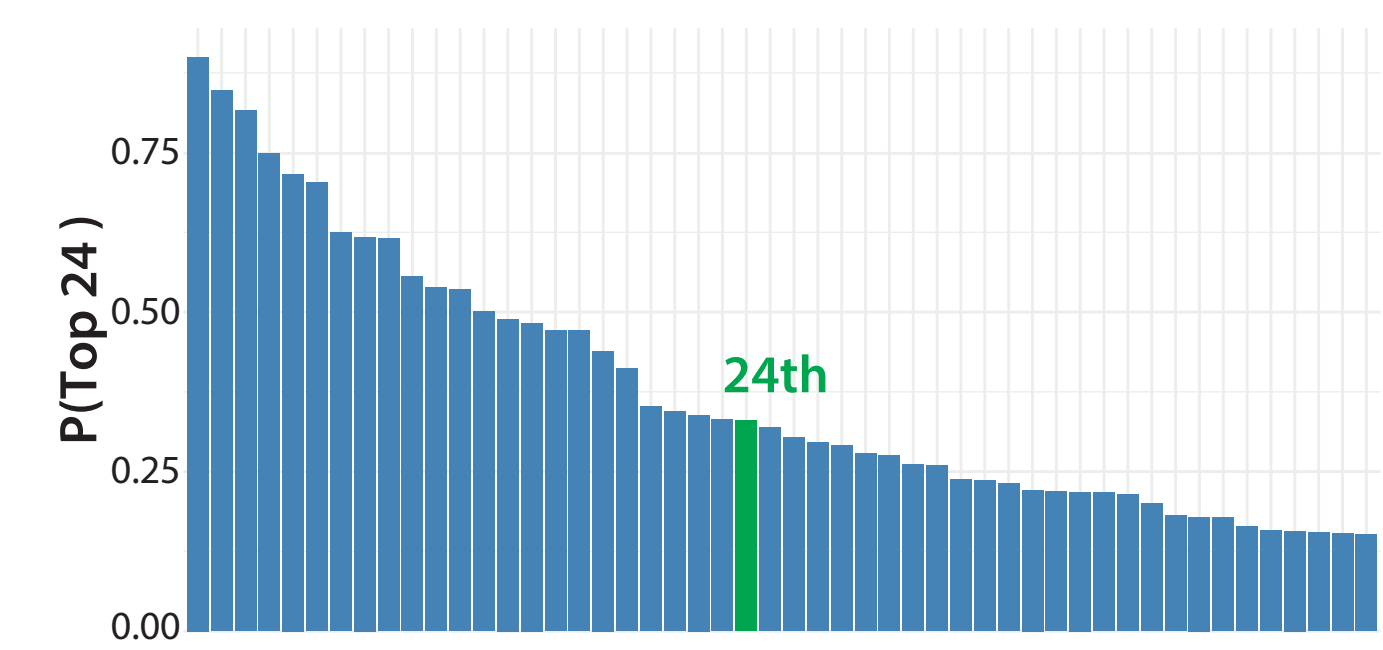
Paper Selection

How can we use quality posteriors to select papers for talks? Consider choosing 24 contributed talks, like in CCN2023.

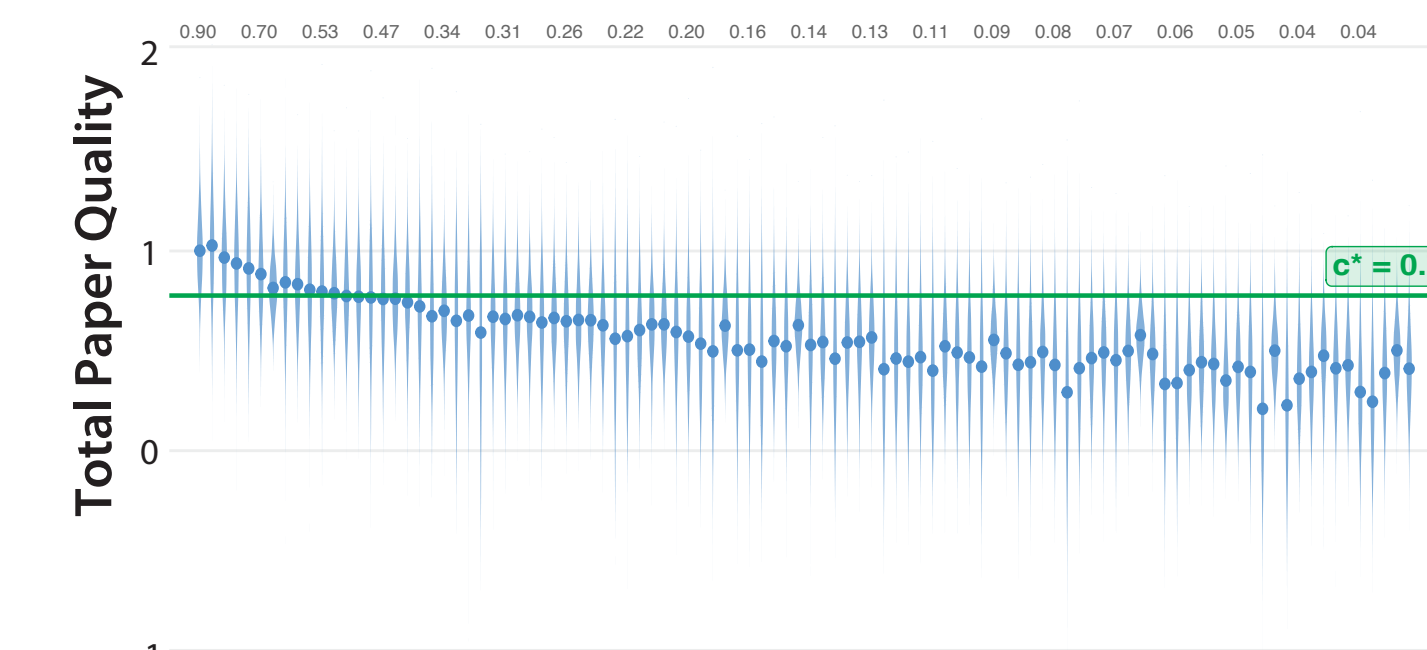
NeurIPS approach: select papers based P(top 24). However, there is high uncertainty, so adjacent papers are equally deserving of a talk.

Proposal: select papers with an calibrated lottery. Randomly select papers, weighted by their exceedance probability above an self-consistent threshold (the threshold where summed exceedance equals the number of slots). Differences in P(Accept) smoothly tract difference in quality.

High Uncertainty in Paper Rankings



Calibrated threshold: where summed posterior exceedance equals the number of slots



Weighted Lottery: probability of selection equals the posterior exceedance above the threshold

