

Comparison of the Effectiveness of Rumination-Focused Cognitive Behavioral Therapy and Classical Cognitive Behavioral Therapy on Rumination and Positive Beliefs About Rumination in Individuals With Obesity: A Single-Case Multiple-Baseline Study

Hasan. Mashhadizade¹, Azam. Davoodi^{1*}, Ghasem. Naziri¹, Najmeh. Fath¹

¹ Department of Clinical Psychology, Shi.C., Islamic Azad University, Shiraz, Iran

* Corresponding author email address: azam.davoodi@iau.ac.ir

E d i t o r

Seyed Hamid Atashpour
Associate Professor, Department of
Psychology, Isfahan (Khorasgan)
Branch, Islamic Azad University,
Isfahan, Iran
hamidatashpour@gmail.com

R e v i e w e r s

Reviewer 1: Hajar Torkan
Assistant Professor, Department of Psychology, Islamic Azad University, Isfahan
Branch (Khorasgan), Isfahan, Iran. h.torkan@khuisf.ac.ir
Reviewer 2: Farhad Namjoo
Department of Psychology and Counseling, KMAN Research Institute, Richmond
Hill, Ontario, Canada. Email: farhadnamjoo@kmanresce.ca

1. Round 1

1.1. Reviewer 1

Reviewer:

Table 1 reports baseline lengths of “14 days,” “21 days,” and “28 days,” corresponding to four, five, and six baseline assessment points; however, the manuscript later states that “the exact time intervals between assessment occasions were not documented in the available records.” These statements create an apparent methodological inconsistency. If exact measurement intervals were not documented, the basis for confidently assigning baseline durations of 14, 21, and 28 days needs to be explained. Conversely, if baseline duration was prospectively specified, it should be possible to state the scheduled assessment frequency. Please reconcile these descriptions and provide the planned and actual assessment schedule as accurately as possible, because temporal spacing is essential in evaluating multiple-baseline designs.

The intervention description states that “both interventions were delivered in ten weekly 60-minute sessions” but that “the intervention phase comprised eight assessment occasions” that “did not correspond one-to-one with the ten therapy sessions.” This is an important methodological feature that requires much greater precision. Please indicate, if recoverable, after which sessions the eight assessments occurred or at least explain the assessment rule used prospectively. Without knowing whether

assessments occurred, for example, after Sessions 1, 2, 4, 5, 6, 8, 9, and 10, the timing and immediacy of treatment effects cannot be evaluated reliably. Because one of the principal claims is that positive beliefs declined earlier in RFCBT and were followed by accelerated declines in rumination, the temporal mapping between treatment exposure and outcome measurement is especially important.

The statement that “baseline scores were high, showed little variability, and lacked a systematic downward trend” should be supported by explicit participant-level evidence rather than qualitative assertion alone. Because stability of the baseline is fundamental to causal inference in a multiple-baseline design, the manuscript should report baseline slopes, variability ranges, or another transparent numerical characterization for each participant and each outcome. Visual inspection should remain central, but the authors should present sufficient information for readers to verify the claim of baseline stability. This is particularly important because Tau-U calculations are sensitive to baseline trends, and the manuscript later reports differences between the two treatments that may partly reflect trend correction.

The manuscript reports group-level NAP and Tau-U values obtained by “calculating the arithmetic mean of the values for the three participants in each treatment condition.” This aggregation requires stronger methodological justification. Single-case effect sizes are fundamentally participant-level estimates, and simple arithmetic averaging may obscure differences in the number of baseline observations, sampling variability, and response heterogeneity across cases. The authors appropriately call these values “descriptive summaries,” but the manuscript should further explain why averaging was chosen rather than reporting participant-level estimates as the primary evidence or using established meta-analytic procedures for aggregating single-case effect sizes. At minimum, conclusions regarding comparative treatment superiority should rely primarily on replicated individual effects rather than these unweighted group averages.

Response: Revised and uploaded the manuscript.

1.2. Reviewer 2

Reviewer:

The treatment description is informative but insufficiently detailed to establish that the two interventions were comparable in therapeutic intensity and distinguishable in content. The sentence “Both interventions were structurally matched with respect to session duration, homework assignments, and the repeated-assessment schedule” should be supported by more explicit information regarding therapist contact, homework quantity, adherence monitoring, treatment manuals, and procedures used to minimize contamination between classical CBT and RFCBT. Because the same therapist appears to have delivered both interventions, allegiance and cross-contamination effects are realistic concerns. A supplementary session-by-session protocol table specifying objectives, techniques, exercises, and homework for each condition would make the contrast between treatments much more transparent.

The manuscript appropriately acknowledges that “treatment fidelity was also not assessed systematically,” but this limitation is particularly consequential because treatment differentiation is central to the article’s primary conclusion. The authors should clarify whether sessions were audio- or video-recorded, whether checklists were completed, whether supervision was provided, and whether any informal fidelity procedures occurred even if no formal fidelity coefficient was calculated. In the absence of independent fidelity assessment, it cannot be verified that RFCBT-specific techniques were delivered consistently or that classical CBT remained free of explicit rumination-focused strategies. This issue should be discussed not merely as a general limitation but as a direct threat to interpretation of between-condition differences.

The Measures section describes the Ruminative Responses Scale and Positive Beliefs About Rumination Scale but provides no reliability estimates from the present study. The sentence concerning the Persian measures—“Evidence supporting the factorial structure, validity, and reliability of Persian-language measures ... has also been reported”—establishes prior psychometric support but does not substitute for describing measurement performance in the current sample. Traditional internal-consistency coefficients are difficult to interpret with only six individuals, so the authors should acknowledge that sample-based reliability cannot be estimated meaningfully and emphasize prior validation evidence. More importantly, the

manuscript should explain whether repeated administration at many assessment points could have produced testing, reactivity, or response-memory effects, particularly for a brief nine-item metacognitive scale.

The operationalization of “final baseline” in Tables 2 and 3 requires clarification. The analysis section states that percentage improvement was calculated using “the mean of the final two baseline data points,” yet Tables 2 and 3 present a column labeled “Final Baseline” with single values such as 74.00, 78.00, and 76.00. The manuscript should explicitly state that these values represent the mean of the final two baseline measurements rather than an individual final observation, if that is indeed the case. The current terminology could mislead readers into believing that a single final baseline score was used. I recommend relabeling the column as “Mean of Final Two Baseline Observations” throughout.

The calculation of percentage improvement is understandable, but the manuscript should justify why the “mean of the final two baseline data points” was selected as the reference rather than the entire baseline mean. In multiple-baseline single-case research, use of only two terminal observations may increase sensitivity to local variability and can potentially exaggerate or attenuate improvement estimates. This is especially relevant because the baseline phases contained four to six observations. Please explain the methodological rationale for this choice and, ideally, provide a sensitivity analysis comparing percentage improvement calculated from the entire baseline mean with that based on the final two baseline observations. Such an analysis would demonstrate whether the reported treatment advantage is robust to the choice of baseline reference.

Response: Revised and uploaded the manuscript.

2. Revised

Editor’s decision after revisions: Accepted.

Editor in Chief’s decision: Accepted.