

Designing and Validating an Effective Leadership Competency Model for Digital Marketing in the Iraqi Banking Sector

Ayad Hadi. Khaleel¹, Morteza. Movaghar^{2*}, Abolhassan. Hosseini³, Mohsen. Alizadeh Sani⁴

- ¹ Ph.D. Candidate, Department of Business Management, Faculty of Economics and Administrative Sciences, University of Mazandaran, Babolsar, Iran
- ² Associate Professor, Department of Business Management, Faculty of Economic and Administrative Sciences, University of Mazandaran, Babolsar, Iran
- ³ Professor, Department of Business Management, Faculty of Economics and Administrative Sciences, University of Mazandaran, Babolsar, Iran
- ⁴ Associate Professor, Department of Business Management, Faculty of Economics and Administrative Sciences, University of Mazandaran, Babolsar, Iran

* Corresponding author email address: m.movaghar@umz.ac.ir

E d i t o r	R e v i e w e r s
Azar Kafashpoor  Professor, Department of Educational Management and Human Resource Development, Ferdowsi University of Mashhad, Mashhad, Iran kafashpor@um.ac.ir	Reviewer 1: Abbas Monavarian  Professor, Management Department, Tehran University, Tehran, Iran. Email: amonavar@ut.ac.ir Reviewer 2: Rezvan Hosseingholizadeh  Associate Professor, Department of Educational Management and Human Resource Development, Ferdowsi University of Mashhad, Mashhad, Iran. Email: rhgholizadeh@um.ac.ir

1. Round 1

1.1. Reviewer 1

Reviewer:

The assertion that “Data collection continued until theoretical saturation was achieved” requires empirical substantiation. Merely stating that no new concepts emerged does not demonstrate saturation in a grounded-theory study. The authors should indicate approximately at which interview saturation was first observed, whether additional interviews were conducted to confirm it, and whether saturation referred to code saturation, category saturation, or theoretical saturation. A saturation table showing the emergence of new codes/categories across successive interviews would considerably strengthen methodological transparency and support the adequacy of the sample of 19 participants.

The qualitative data-collection process should be described in considerably greater operational detail. Although the manuscript states that semi-structured interviews were conducted, it does not report the interview setting, duration, recording procedure, language, transcription method, principal interview questions, probing strategy, or whether the interview guide was pilot tested. These details are particularly important because the proposed competency model is derived directly from interview data. The authors should provide examples of the major interview questions and explain how questions evolved during theoretical sampling. If interviews were conducted in Arabic and subsequently reported in English, translation and back-translation or equivalent procedures should also be reported to establish linguistic fidelity.

The procedures used to establish qualitative trustworthiness are too general. The manuscript states that “credibility, transferability, dependability, and confirmability criteria were considered,” but no concrete procedures are described. The authors should specify whether member checking, peer debriefing, prolonged engagement, negative-case analysis, audit trails, reflexive memoing, independent coding, or external auditing were undertaken. Simply naming Lincoln-and-Guba-type criteria is not sufficient evidence that they were achieved. The manuscript should explain exactly what procedure corresponded to each criterion and, where multiple analysts participated, how coding disagreements were resolved.

The measurement-model assessment is incomplete because discriminant validity is evaluated solely through the Fornell–Larcker criterion. The statement that “in all cases, the diagonal values are greater than the off diagonal elements” is insufficient under contemporary PLS-SEM practice, because Fornell–Larcker can fail to detect discriminant-validity problems. The authors should additionally report the heterotrait-monotrait ratio (HTMT), including appropriate threshold criteria and preferably bootstrapped confidence intervals. Particular attention should be paid to theoretically related constructs such as technical competencies, knowledge competencies, ethical competencies, marketing outcomes, and organizational outcomes. The presence of diagonal entries of exactly 1.000 for some constructs also deserves explanation because it may indicate single-item constructs or a table-construction issue; if single-item measures were used, this must be explicitly justified.

A serious internal inconsistency appears in the reporting of the contextual and intervening paths. In the text accompanying Figure 2, the manuscript reports “from contextual factors to the strategies (2.213)” and “from intervening factors to the strategies (0.569),” concluding that contextual factors are significant whereas intervening factors are not. However, the subsequent hypothesis-testing section reports exactly the reverse assignment: contextual conditions → strategies has $Z = 0.569$ and $\beta = 0.048$, whereas intervening conditions → strategies has $Z = 2.213$ and $\beta = 0.241$. This discrepancy affects a substantive conclusion of the study and must be resolved by checking the original PLS output. Figure 2, Table 10, the abstract, results narrative, discussion, and conclusion should then all be corrected to reflect one verified set of coefficients consistently.

Authors revised the manuscript and uploaded the new document.

1.2. Reviewer 2

Reviewer:

The grounded-theory analysis needs stronger demonstration of the transition from raw data to theoretical categories. The manuscript explains open, axial, and selective coding conceptually and then states that “50 sub-categories of the research were placed in 6 main categories,” but the analytic chain is not sufficiently visible. The authors should report the number of initial/open codes, explain how duplicate or conceptually overlapping codes were consolidated, specify the number of concepts/categories generated at each stage, and present representative participant quotations for every major category. At present, Table 1 provides only selected examples, some of which lack corresponding codes, making it difficult to assess whether the final six-part paradigm is genuinely grounded in the interview material.

The conceptual status of the “central phenomenon” requires clarification. The model identifies individual, social/interactional, technical, knowledge, and ethical competencies as components of the core phenomenon, while causal conditions, contextual conditions, intervening conditions, strategies, and consequences are positioned around it according to the Strauss–Corbin paradigm. However, these five competency domains appear themselves to constitute a multidimensional higher-order construct rather than a single substantive grounded-theory “core category.” The authors should explain the

theoretical rationale for treating these competency dimensions collectively as the central phenomenon and demonstrate how selective coding identified this phenomenon as the integrative category linking the entire model. Otherwise, the model risks reproducing the Strauss–Corbin template structurally without adequately establishing an emergent substantive theory.

Development of the quantitative instrument is underreported and represents a major reproducibility concern. The manuscript states that “a structured questionnaire was designed” from qualitative coding and relevant literature and “evaluated by academic experts and banking professionals to ensure content validity,” but it does not report the number of items, response scale, number of items allocated to each construct, expert-panel size, item-generation rules, item-reduction process, pilot testing, or quantitative indices of content validity. The authors should provide the complete instrument, preferably as a supplementary appendix, and report CVR/CVI or another defensible content-validation procedure. Without these data, readers cannot determine whether the quantitative measures faithfully operationalize the qualitative categories.

The quantitative sampling strategy requires substantial clarification. The manuscript reports a sample of 161 banking professionals but does not explain how the sample size was determined, how respondents were recruited, how many Iraqi banks participated, whether public/private or Islamic/conventional banks were represented, how questionnaires were distributed, the response rate, or the number of cases excluded during data screening. Because PLS-SEM sample adequacy should be justified on the basis of statistical power and model complexity rather than only the traditional “10-times rule,” the authors should provide an a priori or inverse-square-root/power-based sample-size justification. The sampling frame and geographic/institutional coverage are also necessary to evaluate representativeness and external validity.

There is a factual inconsistency in the reliability interpretation that must be corrected. The manuscript states, “all Cronbach’s alpha and composite reliability (CR) coefficients for the first-order latent variables are above the threshold value of 0.70,” whereas Table 3 reports Cronbach’s alpha values of 0.669 for application of digital technology, 0.621 for changing customer behavior and expectations, 0.632 for competitive pressure, and 0.697 for interaction-based competencies. Therefore, the statement is demonstrably incorrect. The authors should accurately acknowledge these coefficients, justify their acceptance in an exploratory scale-development context, and preferably report ρ_A in addition to Cronbach’s alpha and composite reliability, particularly because alpha is sensitive to the number of indicators and tau-equivalence assumptions.

The treatment of convergent validity should be reconsidered. The manuscript reports an AVE of 0.472 for “changing customer behavior and expectation,” below the conventional 0.50 criterion, and then argues that values above 0.40 are acceptable. This should not be presented as unequivocal confirmation of convergent validity. The authors should inspect the relevant indicator loadings, determine whether weak items account for the deficient AVE, and report whether removing poorly performing indicators improves construct validity without compromising content coverage. If the construct is retained with $AVE < 0.50$, the authors should explicitly characterize convergent validity as marginal and justify retention based on composite reliability and substantive considerations rather than claiming that the entire measurement model unambiguously satisfies the standard criterion.

Authors revised the manuscript and uploaded the new document.

2. Revised

Editor’s decision after revisions: Accepted.

Editor in Chief’s decision: Accepted.