



MEASURING ANTICIPATED SATISFACTION: A FULL-SAMPLE APPROACH TO TECHNOLOGY SATISFACTION SCALES IN PRE- ADOPTION CONTEXTS

Ajay Kumar¹, Dr.Balgopal Singh²

¹Research Scholar, Faculty of Management Studies, (FMS-WISDOM), Banasthali Vidyapith,
Rajasthan, India, ajaykbokaro@gmail.com, ORCID id: 0009-0001-6849-5981

²Associate Professor, Faculty of Management, Studies (FMS-WISDOM), Banasthali Vidyapith, Rajasthan,
India, balgopalsingh@banasthali.in, ORCID id: 0000-0001-8402-9617

Article History

Received : 2026-07-22

Revised : 2026-08-25

Accepted : 2026-09-15

Published : 2026-09-22

Abstract

Satisfaction is often only measured among current users of a technology since standard survey questions (e.g., 'I am satisfied with...') assume prior experience. This convention necessarily excludes any participant who has not yet adopted the technology — usually the majority of a sampled population in early technology-diffusion scenarios — and prevents direct, one-sample comparison of projected satisfaction across firms with different adoption statuses. This study presents and substantiates a comprehensive alternative: utilizing a singular four-item projected-satisfaction scale across an entire cross-sectional sample irrespective of adoption status, by administering identical, forward-looking satisfaction items to every respondent.. The methodology relies on Oliver's (1980) expectation-confirmation theory (ECT), which considers pre-consumption expectations as a separate, quantifiable cognitive state against which satisfaction assessments are made, and on the later adaptation of the expectation-confirmation model to technological environments. Utilizing survey data from 450 micro, small, and medium enterprises (195 users and 255 non-users of a governmental e-procurement platform), the SAT scores obtained from adopters and non-adopters demonstrate evidence of cross-group measurement comparability: reliability is notably high and nearly uniform across subgroups ($\alpha = .942$ for users, $\alpha = .936$ for non-users, $\alpha = .946$ combined), a single-factor structure accounts for a similar proportion of variance in both cohorts (85.3% vs. 84.0%), and the scale's correlations with four theoretically relevant constructs (attitude, perceived relative advantage, performance expectancy, and awareness) show no significant differences between users and non-users (all $|z| < 1.3$, $p > .20$). Simultaneously, the average projected satisfaction level is higher among adopters ($M = 3.53$) than among non-adopters ($M = 2.65$, $p < .001$), demonstrating that the scale continues to effectively differentiate between groups rather than devolving into a mere artifact of shared method variance. Collectively, these findings endorse a comprehensive, anticipated-framing methodology for assessing satisfaction, serving as a versatile research instrument for scholars examining groups with varied adoption statuses, restoring statistical efficacy and facilitating comparisons of projected satisfaction across firms with different adoption statuses within a unified cross-sectional framework.

Keywords: anticipated satisfaction; expectation-confirmation theory; scale measurement; pre-adoption research; non-adopters; technology satisfaction; survey methodology

1. Introduction

Research on technology satisfaction encounters a structural sampling issue that is seldom identified as such. The conventional satisfaction query prompts respondents to assess their interaction with a technology ('I am content with X'; 'X has fulfilled my expectations'), hence it can only be answered, in a literal context, by those who have really used the technology. In cross-sectional analyses of technology dissemination — which must encompass entities at all stages from ignorance to complete adoption — this standard necessitates a decision: either limit the satisfaction evaluation to the adopter subset, excluding the data of non-adopters and frequently a significant portion of the overall sample, or eliminate satisfaction from the investigation entirely for all except verified users.

Both alternatives are expensive. Limiting the analysis to adopters diminishes statistical power, introduces a self-selection bias (only firms that opted to adopt can be assessed, and among them, response patterns may differ systematically from those of non-adopters), and precludes an entire category of research inquiry: how do the anticipated satisfaction levels of potential adopters contrast with the satisfaction that adopters ultimately report? Responding to that inquiry often requires a longitudinal panel — monitoring the same companies from prior to adoption to subsequent stages — which is costly and time-consuming to implement, and is seldom accessible in practical MSME or e-governance studies.

This study presents a more straightforward, theoretically sound option that eliminates the need for panel design: identical forward-looking, projected-satisfaction items were administered to all respondents regardless of adoption status. All respondents — whether or not they have adopted the technology — are asked to assess the same forward-looking, projected satisfaction item, informed by their understanding of the technology and its intended use. This single formulation accesses the same underlying evaluative framework — a positive vs. negative overall assessment of the technology in relation to anticipated outcomes — for adopters and non-adopters alike, rather than posing conceptually different questions to the two groups.

The conceptual justification for seeing a projected assessment as a valid gauge of the same fundamental construct as a realized evaluation is derived from expectation-confirmation theory (ECT). The subsequent sections of this document (a) delineate the theoretical justification, (b) elucidate the comprehensive measurement methodology employed in a survey of 450 MSMEs concerning a government e-procurement platform, and (c) present a series of evidence of cross-group measurement comparability assessments — reliability, dimensionality, and nomological validity — contrasting projected-satisfaction responses between adopters and non-adopters. The objective is not to provide a significant discovery on the technology in issue, but to authenticate and record a measuring methodology that future scholars examining mixed adopter/non-adopter demographics may replicate.

2. Theoretical Grounding

2.1 Expectation as a distinct, pre-consumption cognitive state

Oliver's (1980) expectation-confirmation model posits that satisfaction is not merely a reflection of product or service performance, but rather the result of juxtaposing that performance with pre-existing expectations: consumers establish expectations prior to consumption, engage with the product or service, assess the actual performance against those expectations (leading to confirmation or disconfirmation), and ultimately derive a satisfaction evaluation based on this comparison. This model basically says that pre-consumption expectations are a unique cognitive state that can be measured independently, and that they are developed before and independent of the subsequent satisfaction evaluation against which they will be judged. Thus, they are not just a measurement artifact but rather a significant contributor to the satisfaction process.

This is the conceptual basis for the assessment of major aspects of happiness before consumption. If expectation is an independent, measurable state, with a particular causal role in the final assessment of satisfaction, then asking a non-adopter about their expected satisfaction is not a categorical error. It is an attempt to measure one of the two identified components (expectation) of the process described by ECT, expressed in the same evaluative terms (the specific satisfaction items) that will later be used, in validated form, to assess the resulting judgment among adopters.

2.2 Extension to technology contexts and continuation

Bhattacharjee (2001) extended the expectation-confirmation model to information-technology continuation and showed that post-adoption satisfaction, which is influenced by the confirmation of pre-adoption expectations, is an important antecedent of a user's intention to continue using a technology. The literature is

clear in distinguishing pre-adoption expectations and post-adoption satisfaction evaluation as separate but related constructs within an integrated theoretical framework. It emphasizes that a well-defined measure of 'expected satisfaction', measured before adoption, is a theoretically valid construct in its own right, and not just a lesser form of the post-adoption evaluation.

Later refinements of the confirmation/disconfirmation paradigm (e.g., Spreng, MacKenzie, & Olshavsky, 1996) have separated expectations of specific attributes from general or global expectations of a product's desirability and have shown that each contributes to later satisfaction judgments through partially different routes. The relevant implication for our present purpose is more circumscribed: whatever the precise mechanism, the literature agrees on the need to treat pre-consumption expectations as genuinely meaningful and empirically manageable, and as such, to ask questions to non-adopters about their expected satisfaction along the lines of the post-adoption scale.

2.3 From theory to instrument design

Converting this warrant into a tool requires a single design decision: the item stem is framed in a forward-looking, projected register (e.g., 'our organization believes GeM will meet its needs') and administered identically to every respondent — adopter and non-adopter alike — rather than varying in tense or epistemic stance by adoption status. This uniform projected framing has the added advantage of eliminating any framing-based confound between the two groups, since both answer the same item. This bears resemblance to the conditional or hypothetical-scenario framing employed in survey methodology to draw out assessments of unexperienced situations, yet it is distinctly rooted in ECT's conceptualization of expectation as a genuine precursor construct rather than an improvised modification.

3. Method

3.1 Sample and context

The data is collected from a cross sectional survey of 450 micro, small and medium companies (MSMEs) operating in Delhi National Capital Region in relation to India's Government e-Marketplace (GeM) a digital public procurement platform. During the study period, 195 companies (43.3%) had an active GeM seller account (adopters) and 255 companies (56.7%) did not have one (non-adopters). This setup – a strong majority of non-adopters – is typical of early to mid-stage technology diffusion demographics and is precisely the scenario where a traditional satisfaction metric confined to adopters would miss the overwhelming majority of the sample.

3.2 Instrument

All the sample was assessed with a four-item satisfaction measure (SAT1–SAT4), on a five-point Likert scale (1 = strongly disagree, 5 = strongly agree). All four items measured the same underlying evaluative dimension, overall prospective satisfaction with GeM relative to expectations, and were worded with the same projected framing to all respondents (e.g., overall, our organization believes GeM will meet its needs). This single item framing was uniformly applied throughout the entire sample in line with the treatment of Satisfaction as a projected, rather than strictly post-usage, construct (Section 2).

3.3 Comparison constructs

To evaluate nomological validity, satisfaction ratings were correlated with four associated constructs assessed on the same participants: Attitude toward GeM (4 items, $\alpha = .950$), Perceived Relative Advantage (7 items, $\alpha = .913$), Performance Expectancy (3 items, $\alpha = .825$), and Awareness of GeM (5 items, $\alpha = .902$). Satisfaction among adopters was also associated with a 4-item Adoption Extent scale ($\alpha = .936$) that evaluates the intensity of platform utilization, a concept lacking a counterpart for non-adopters, thus functioning as a verification of the discriminative behavior within the adopter subgroup rather than a comparative analysis across groups.

3.4 Analytical approach

Four categories of evidence were employed to evaluate whether the SAT scores obtained from non-adopters and adopters function as psychometrically equivalent measures of the same fundamental construct: (i) internal-consistency reliability (Cronbach's α) calculated independently for each subgroup and compared to the aggregated sample value; (ii) corrected item-total correlations (CITC) within each subgroup to eliminate any single item exhibiting irregular behavior in one framing; (iii) the proportion of item variance elucidated

by a singular principal component, calculated separately for each subgroup, serving as a rough assessment of configural (structural) equivalence; and (iv) the equality of the scale's correlations with the four comparative constructs across groups, assessed using Fisher's *r*-to-*z* transformation for independent correlations. The two-sample Welch *t*-test was used to evaluate whether projected-satisfaction levels differed between adopters and non-adopters, and Levene's test was used to evaluate the comparability of variances of the responses between the groups. A significant mean difference, with comparable variances, would suggest that the scale is effective in differentiating between groups without sacrificing the measurement precision in either group.

4. Results

4.1 Reliability is preserved across adoption-status subgroups

Cronbach's alpha was .946 for the combined sample (*N* = 450), .942 for adopters (*n* = 195) and .936 for non-adopters (*n* = 255), both responding to the same anticipated-satisfaction items — all significantly exceeding standard benchmarks for exceptional reliability, with differences of no more than one hundredth of a point. The adjusted item-total correlations were consistently elevated and closely aligned across the two categories (adopters: .85–.87 across items; non-adopters: .84–.86 across items), with no item exhibiting a significantly poorer association with the total score in either context (Table 1).

Table 1. Corrected item-total correlations (CITC) and reliability, by adoption status (pooled-sample $\alpha = .946$)

Item	CITC — adopters (<i>n</i> =195)	CITC — non-adopters (<i>n</i> =255)
SAT1	0.870	0.850
SAT2	0.872	0.847
SAT3	0.852	0.838
SAT4	0.855	0.860
Cronbach's α	0.942	0.936

4.2 A single factor structure holds in both groups

A solitary component (the initial unrotated main component) accounted for 85.3% of item variation in adopters and 84.0% in non-adopters, in contrast to 86.1% for the combined sample — the variance explained varied by just 1.3 percentage points between the two groups. The component loadings were similarly consistent in size across items and groups (Table 2), providing approximate support for configural equivalence: both the adopter and non-adopter subgroups exhibit the same one-factor structure for the SAT scale.

Table 2. First principal-component loadings and variance explained, by group

Group	<i>n</i>	SAT1	SAT2	SAT3	SAT4	Variance explained (1st PC)
Adopters	195	1.16	1.19	1.15	1.15	85.3%
Non-adopters	255	1.19	1.18	1.16	1.20	84.0%
Pooled sample	450	1.26	1.25	1.25	1.25	86.1%

4.3 Mean satisfaction differs by group, but measurement precision does not

Non-adopters reported lower anticipated satisfaction (*M* = 2.65, *SD* = 1.18) than adopters did (*M* = 3.53, *SD* = 1.16), Welch's *t*(415.4) = 7.93, *p* < .001. This pattern was found for all four individual measures (all item-level *t* > 6.8, *p* < .001; Table 3). Levene's test showed no significant difference in variance of responses between the two groups (*F* = 0.14, *p* = .71), indicating that non-adopters were neither more nor less consistent in their evaluations than adopters. The two groups differ in central tendency but not in variability of measurement, which is what we would expect if the scale is successfully differentiating between two genuinely distinct evaluative conditions (projected

satisfaction reported by adopters and non-adopters, who differed in their adoption status and therefore in the extent of their direct platform experience) rather than adoption status simply being associated with more erratic data.

Table 3. Item and composite means: adopter vs. non-adopter satisfaction

Item	Adopters M	Non-adopters M	t	p
SAT1	3.49	2.61	7.30	< .001
SAT2	3.53	2.69	6.87	< .001
SAT3	3.61	2.66	7.86	< .001
SAT4	3.49	2.62	7.16	< .001
Composite	3.53	2.65	7.93	< .001

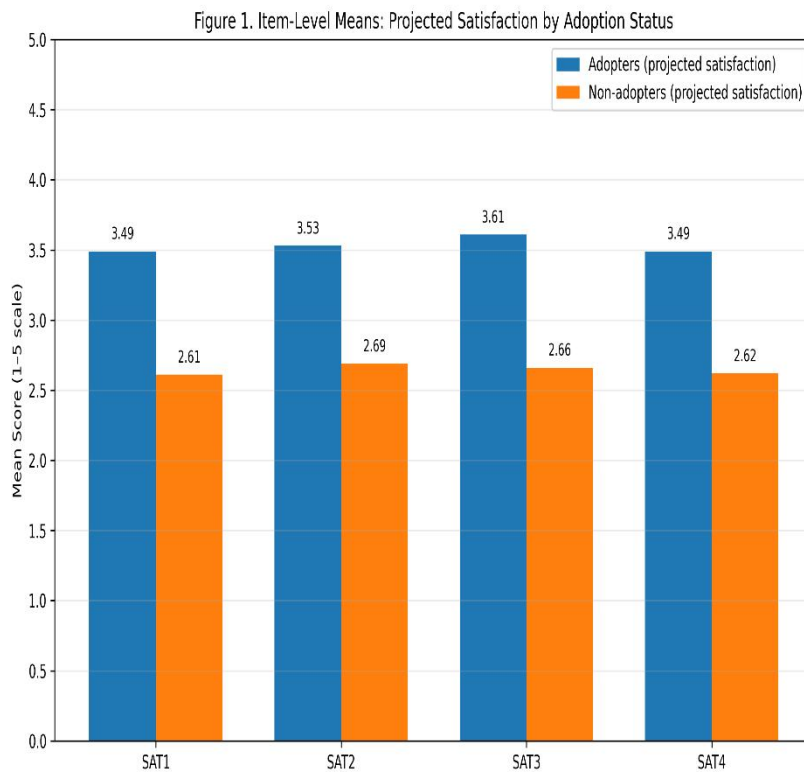


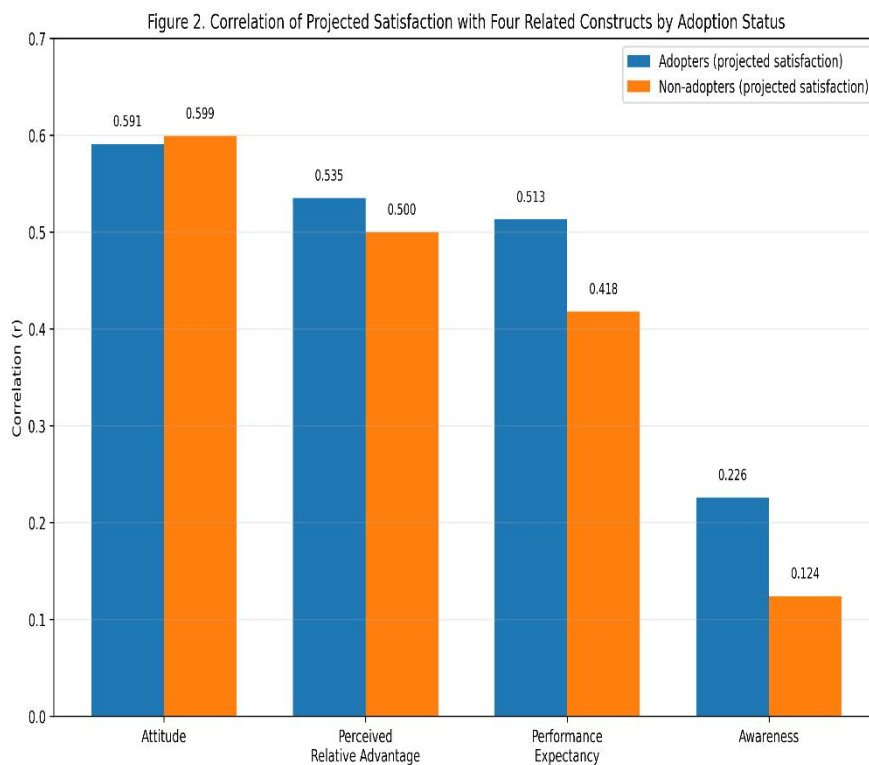
Figure 1. Item-level means (± 1 SE) for satisfaction among adopters versus non-adopters.

4.4 Nomological validity holds equally in both groups

If projected satisfaction is conceptualized as a common evaluative construct across firms with different adoption statuses, the association patterns of satisfaction with theoretically related variables should be similar between adopters and non-adopters, regardless of the absolute levels of the construct across groups. The results indicated that there were no statistically significant differences between groups in the relationship between satisfaction and the different constructs: Attitude (adopters $r = .591$, non-adopters $r = .599$), Perceived Relative Advantage ($r = .535$ vs. $r = .500$), Performance Expectancy ($r = .513$ vs. $r = .418$) and Awareness ($r = .226$ vs. $r = .124$) (all $|z| \leq 1.27$, all $p \geq .21$; Table 4, Figure 2). As a further group-specific assessment, satisfaction was moderately correlated with Adoption Extent among adopters ($r = .539$). This construct has no equivalent for non-adopters and therefore provides additional construct-related evidence within the adopter subgroup

Table 4. Fisher r-to-z comparison of satisfaction's correlations with related constructs, adopters vs. non-adopters

Comparison construct	r — adopters	r — non-adopters	z	p
Attitude	0.591	0.599	-0.12	.901
Perceived Relative Advantage	0.535	0.500	0.51	.610
Performance Expectancy	0.513	0.418	1.27	.205
Awareness	0.226	0.124	1.10	.274

**Figure 2. Correlation of satisfaction with four related constructs, compared across adopters and non-adopters .**

5. Discussion

Four distinct lines of evidence align on a singular conclusion: the SAT scores obtained from non-adopters... as the SAT scores obtained from adopters — both captured with the same anticipated-satisfaction items. The reliability is elevated and almost uniform across both adoption-status groups; a singular factor structure explains a similar proportion of item variance in each group; and the scale's correlations with four independently assessed, theoretically pertinent constructs are statistically equivalent across groups. Simultaneously, the scale exhibits a lack of insensitivity to adoption status: non-adopters' anticipated satisfaction is markedly lower than adopters' anticipated satisfaction, without any concomitant decline in measurement accuracy, suggesting that the scale effectively identifies a genuine and theoretically anticipated disparity in central tendency between the anticipated satisfaction of firms without direct platform experience and that of firms with it— aligning with the premise of ECT that expectation and validated post-usage satisfaction are interconnected yet distinct states.

The orientation of this discrepancy (non-adopters lower than adopters) is noteworthy in itself. It implies that,

in this context, companies that have not yet embraced the platform reported lower projected satisfaction than companies that had adopted it.— a trend more aligned with an information or salience deficit in the pre-adoption group's expectations rather than with exaggerated claims by adopters. Due to the current design being cross-sectional, this discrepancy cannot be validated against the subsequent satisfaction of the same companies; a true longitudinal replication — reinterviewing non-adopters who subsequently adopt — serves as the most robust confirmation that a firm's anticipated satisfaction before adoption is predictive of its own satisfaction after adoption, as ECT suggests.

The pragmatic significance of the methodology shown above is clear-cut. Limiting the current sample to only adopters would have excluded 255 out of 450 respondents (56.7% of the sample) from any analysis based on satisfaction, and would have rendered it unfeasible, within a single cross-sectional wave, to inquire whether the anticipated satisfaction of non-adopters aligns with the anticipated satisfaction reported by firms that have already adopted. The comprehensive sample, projected-framing methodology retrieves the entire sample for satisfaction-oriented evaluation, and — as demonstrated here to be psychometrically comparable measurement properties across adopters and non-adopters — allows for direct comparison of projected satisfaction between firms with different adoption statuses within a single cross-sectional survey without necessitating panel data.

6. Implications for Research Practice

Researchers examining technology or platform uptake in demographics with a significant proportion of non-adopters should consider administering a single, uniformly worded anticipated-satisfaction scale to the entire sample — rather than the conventional experience-based item, which only adopters can meaningfully answer — instead of limiting satisfaction assessment solely to adopters.

Researchers must first conduct equivalence assessments for their own instrument before combining or comparing satisfaction scores across adopter and non-adopter subgroups: subgroup reliability, subgroup dimensionality (for instance, through a single-factor variance-explained analysis), and the equality of the scale's correlations with theoretically related constructs across groups. None of these assessments need specialized software; all were calculated using ordinary descriptive and correlational statistics.

A notable disparity in mean satisfaction levels between adopters and non-adopters should not, in isolation, be interpreted as counterevidence to equivalence — as long as reliability, dimensionality, and nomological correlations are consistent across groups, such a mean difference indicates that the scale is effectively distinguishing between two distinct yet interconnected evaluative conditions, precisely as ECT suggests.

This equivalency assessment should, if possible, be accompanied by a modest longitudinal follow-up sample (re-surveying non-adopters post-adoption choice) to directly evaluate the most robust version of the expectation-confirmation hypothesis.

7. Limitations

This is an exercise in measurement validation, rather than a comprehensive evaluation of GeM adoption results, and several constraints are pertinent to that objective. Initially, equivalence was evaluated indirectly via reliability, approximate dimensionality (variance explained by a single component), and comparisons of nomological correlations; a comprehensive multi-group confirmatory factor analysis incorporating formal measurement-invariance testing (configural, metric, and scalar invariance) would yield a more stringent assessment of equivalence than the methodologies employed here, which were selected for their reproducibility without the need for specialized structural-equation-modeling software. Secondly, the precise phrasing of the expected-framing items is presented here in broad terms instead of being quoted verbatim from the field instrument; researchers utilizing this method ought to conduct pilot tests of their own item formulations for clarity among non-adopters, who may differ in their ability to envision a platform they have not yet encountered. Thirdly, due to the cross-sectional nature of the design, the identified gap between non-adopters' and adopters' anticipated satisfaction cannot be validated... since these anticipated-satisfaction scores were reported by different firms rather than the same firms surveyed before and after adoption rather than the same firms at two distinct time intervals. Fourth, the methodology presupposes that non-adopters possess adequate indirect knowledge of the technology (via awareness, word of mouth, or marketing) to develop a significant anticipated judgment; in contexts where non-adopters are completely oblivious to the technology's existence, an anticipated-satisfaction query may not yield a meaningful response and should be preceded by a fundamental awareness assessment, as demonstrated in the broader instrument of the current study.

8. Conclusion

Assessment of satisfaction should not be limited to the existing users of a technology. Rooted in the expectation-confirmation theory's conceptualization of pre-consumption expectations as a separate, quantifiable entity, a single, uniformly worded anticipated-satisfaction scale — administered identically to adopters and non-adopters alike — demonstrated consistent reliability, dimensionality, and nomological validity across both perspectives while effectively distinguishing between them. For researchers engaging with populations that comprise a significant proportion of non-adopters, this comprehensive measurement approach provides an effective means to regain statistical power, circumvent selection bias limited to adopters, and facilitate direct comparisons of projected satisfaction between adopters and non-adopters within a single cross-sectional survey wave

References

1. Oliver, R. L. (1980). A cognitive model of the antecedents and consequences of satisfaction decisions. *Journal of Marketing Research*, 17(4), 460–469.
2. Bhattacharjee, A. (2001). Understanding information systems continuance: An expectation-confirmation model. *MIS Quarterly*, 25(3), 351–370.
3. Spreng, R. A., MacKenzie, S. B., & Olshavsky, R. W. (1996). A reexamination of the determinants of consumer satisfaction. *Journal of Marketing*, 60(3), 15–32.
4. Venkatesh, V., Morris, M. G., Davis, G. B., & Davis, F. D. (2003). User acceptance of information technology: Toward a unified view. *MIS Quarterly*, 27(3), 425–478.
5. Rogers, E. M. (2003). *Diffusion of Innovations* (5th ed.). Free Press.