

## Protecting Muslim Consumers in Artificial Intelligence-Based Digital Transactions: An Islamic Economic Law Perspective

**Deden Tria Niamillah ZA**

Sekolah Tinggi Agama Islam Miftahul Ulum Tasikmalaya, Indonesia

Email: dedentria9@gmail.com

---

**Keywords:**

Artificial Intelligence; Consumer Protection; Digital Transactions; Islamic Economic Law; Muslim Consumers

---

**Abstract**

Background: Artificial intelligence (AI) increasingly mediates digital commerce through recommendation systems, chatbots, behavioral profiling, automated credit assessment, and personalized offers. These technologies can improve convenience; however, they may also intensify information asymmetry, algorithmic opacity, manipulation risks, and uncertainty regarding Sharia compliance. Objective: This illustrative study evaluated whether an integrated program grounded in Indonesian consumer protection law, AI literacy, and Islamic economic law could strengthen Muslim consumers' capacity to identify and respond to risks associated with AI-based digital transactions in Tasikmalaya. Method: A quasi-experimental pretest-posttest control group design was modeled using synthetic data from 83 Muslim digital consumers, comprising 43 participants in the intervention group and 40 in the control group. The intervention consisted of five weekly 90-minute sessions. The primary outcome was measured using a 24-item Muslim Consumer Protection Literacy in AI Transactions scale. The analysis included descriptive statistics, Shapiro-Wilk and Levene's tests, paired-samples and independent-samples t tests, Cohen's d, and analysis of covariance (ANCOVA) controlling for pretest scores. Results: Pretest scores did not differ significantly between the groups. The intervention group's mean score increased from  $70.30 \pm 11.87$  to  $82.21 \pm 9.69$ , whereas the control group's mean score changed from  $74.03 \pm 10.59$  to  $75.83 \pm 11.97$ . The between-group difference in gain scores was 10.11 points (95% CI: 5.46 to 14.75; Cohen's d = 0.95). ANCOVA indicated a significant adjusted intervention effect,  $F(1, 80) = 16.49$ ,  $p < .001$ , partial  $\eta^2 = .171$ . Conclusion: Within the synthetic dataset, the integrated legal, AI, and Sharia literacy intervention produced a meaningful improvement in consumer protection literacy. However, empirical field testing is required before substantive conclusions can be drawn regarding Muslim consumers in *Tasikmalaya*.

---

## INTRODUCTION

Artificial intelligence (AI) is no longer a peripheral feature of digital commerce (Areiqat et al., 2020; Chouhan, 2025; Singh et al., 2026; Song et al., 2019). It operates within recommendation systems, search rankings, targeted advertising, dynamic pricing, customer-service chatbots, fraud detection systems, and automated credit assessments. These applications reduce search costs and can improve transaction relevance, but they also shift important aspects of consumer decision-making from visible human judgment to data-driven systems whose underlying logic is often difficult for users to understand or scrutinize. Marketing research has therefore moved beyond the question of whether AI creates efficiency and has increasingly

examined how AI affects consumer autonomy, trust, vulnerability, and perceived control (Davenport et al., 2020).

AI also reshapes marketing strategy by redistributing analytical, relational, and decision-making functions between firms and intelligent systems (Huang & Rust, 2021).

From the consumer perspective, AI-mediated experiences can alter perceived agency and control during service encounters (Puntoni et al., 2021).

Recent systematic evidence indicates that engagement with intelligent technologies depends substantially on the context in which they are used (Hollebeek et al., 2024).

Consumers' ability to understand and evaluate AI-consumer interactions is therefore a central condition for informed decision-making (Ubal et al., 2024).

The consumer protection problem becomes more pronounced when AI is used to influence rather than merely assist consumer choice. Algorithmic personalization can determine which products are displayed, how offers are sequenced, which prices or financing options are presented, and which persuasive cues receive prominence. These practices can create structural information asymmetries. Consumers may understand the immediate offer but remain unaware of the data, inferences, or optimization objectives that generated it. Consumer policy research shows that fairness and accountability influence whether AI promotes or undermines consumer welfare (Mills et al., 2023).

Research on recommender systems similarly demonstrates that fairness should be evaluated from the consumer perspective rather than solely through technical performance metrics (Vassøy & Langseth, 2024).

Human-centered explainability research shows that explanations are most useful when they enable end users to understand relevant aspects of system behavior (Laato et al., 2022).

Explainability can also influence perceived trustworthiness and acceptance when explanations are understandable and provide meaningful insights into system decisions (Shin, 2021).

Manipulative interface design adds another layer of consumer vulnerability. Dark patterns exploit attention, default effects, scarcity cues, obstruction, and asymmetric choice architecture to steer consumers toward decisions that may not reflect their considered preferences. Experimental legal research has shown that dark patterns can materially influence consumer choices and undermine meaningful consent (Luguri & Strahilevitz, 2021).

Earlier research documented how deceptive interface strategies have evolved across digital services and why they can be difficult for users to recognize (Narayanan et al., 2020).

Recent experimental evidence indicates that such designs can encourage impulse purchasing and that counter-nudges may only partially mitigate their effects. AI can intensify this problem because personalization enables platforms to identify which persuasive cues are most effective for particular users (Sin et al., 2025).

Deceptive design in AI-mediated interfaces is therefore not merely a user-experience issue but also a consumer law concern involving autonomy, disclosure, and redress (De Conca, 2023).

In Indonesia, recent comparative legal analysis has identified a regulatory gap because existing consumer protection and electronic commerce regulations do not yet expressly classify or prohibit the full range of digital dark patterns (Az-Zahra' et al., 2025).

Indonesia already has several legal instruments relevant to digital consumer protection. Law No. 8 of 1999 on Consumer Protection establishes consumers' rights to accurate information, safety, choice, fair treatment, and compensation. Government Regulation No. 80 of 2019 regulates trade through electronic systems. Law No. 27 of 2022 on Personal Data Protection strengthens the legal basis for lawful and accountable personal data processing, while Law No. 1 of 2024, the Second Amendment to the Electronic Information and Transactions Law, reinforces the governance of electronic information and transactions. In the financial technology sector, the Financial Services Authority's 2023 guidelines on responsible and trustworthy AI ethics emphasize benefit, fairness and accountability, transparency and explainability, robustness, and security. These instruments provide important legal foundations, but they do not necessarily ensure that ordinary consumers can recognize when AI-mediated practices threaten their rights or understand how to respond appropriately.

A study of Sharia-oriented e-commerce transactions identified information disclosure and responsibility as central elements of consumer protection (Fadlurrahman & Fikrianihayah, 2022).

More recent Indonesian research has reported continuing gaps between formal legal guarantees and practical consumer protection in digital transactions (Malasari et al., 2026).

From an Islamic law perspective, consumer rights in e-commerce remain closely associated with truthful disclosure and fair transactional conduct (Widyastuti et al., 2022).

For Muslim consumers, legal protection also has a distinct normative dimension. Islamic economic law does not evaluate a transaction solely on the basis of formal consent or technical validity. It also considers whether consent is informed and voluntary, whether material uncertainty is excessive, whether relevant information is concealed, whether one party is unfairly disadvantaged, and whether the transaction protects legitimate wealth and welfare. Islamic AI ethics scholarship places justice, public benefit, accountability, and protection from harm at the center of responsible technological governance (Elmahjub, 2023).

Related research on AI-mediated ethics emphasizes that technological efficiency should not displace the moral quality of human and institutional decision-making (Ghaly, 2024).

In online commerce, AI-based recommendations, dynamic pricing, automated financing, and data profiling can generate *maṣlaḥah* when they reduce transactional friction and improve access. However, the same systems can create *mafsadah* when opacity facilitates deception, excessive *gharar*, discriminatory treatment, or exploitative financing practices (Asshidqi et al., 2026).

The *maqāṣid al-sharīʿah* perspective therefore provides a normative bridge between technological governance and the protection of wealth, intellect, dignity, and genuine consent (Puad & Hamdi, 2025).

The Islamic fintech literature has expanded rapidly, but much of it remains conceptual, adoption-oriented, or institution-focused. Systematic review evidence shows that customers are frequently examined, yet operational governance and consumer protection mechanisms remain underdeveloped themes (Firmansyah & Harsanto, 2022).

More recent research has called for stronger harmonization among technological innovation, national regulation, and Sharia governance, particularly where AI affects financing and automated decision-making (Herlambang & Wahyudi, 2025).

Emerging proposals for the Sharia governance of generative AI emphasize lifecycle accountability, transparency, and dual oversight mechanisms. These contributions are important, but they do not establish whether ordinary Muslim consumers can acquire the practical competencies needed to identify AI-mediated risks, assess their Sharia implications, preserve digital evidence, and use available redress mechanisms (Zafar, 2026).

A parallel body of AI literacy research suggests that these competencies can be developed through education. AI literacy is commonly understood as encompassing awareness, use, evaluation, and ethical judgment rather than technical programming skills alone (Ng et al., 2021).

Measurement research has further operationalized AI literacy as a set of user competencies that can be systematically assessed (Wang et al., 2023).

The Meta AI Literacy Scale provides a multidimensional approach to measuring AI-related competencies among non-specialists (Carolus et al., 2023).

A separate Delphi study developed an item set specifically for assessing AI literacy among non-experts (Laupichler et al., 2023).

Course evaluation research has shown that university students from diverse backgrounds can improve their conceptual understanding through structured AI literacy instruction (Kong et al., 2021).

A subsequent evaluation found that conceptual learning and empowerment can also increase when AI literacy courses are designed around accessible, non-programming content (Kong et al., 2022).

However, these interventions rarely integrate consumer rights, dispute resolution mechanisms, Sharia transaction principles, and the practical identification of manipulative commercial uses of AI. This separation creates a specific research gap at the intersection of AI literacy, consumer protection, and Islamic economic law.

This study addressed this gap by developing an integrated literacy intervention for Muslim consumers in Tasikmalaya, West Java. Its novelty lies in conceptualizing consumer protection as an integrated legal, technological, and Sharia competency rather than as legal knowledge alone. The intervention connected AI recognition with consumer rights, personal data awareness, the principles of *riḍā*, *ʿadl*, *amānah*, *gharar*, *tadlīs*, *khiyār*, and *ḥifẓ al-māl*, and practical redress behavior. The objective of the study was to evaluate whether this integrated program improved Muslim consumers' protection literacy in AI-based digital transactions compared with standard digital commerce information. The hypothesis was formulated as follows: H1: After adjustment for baseline scores, participants receiving the integrated program would demonstrate significantly higher posttest consumer protection literacy than participants in the control group.

This research is expected to provide both theoretical and practical benefits. Theoretically, it enriches academic discourse at the intersection of AI literacy, consumer protection, and Islamic economic law by proposing an integrated framework that connects legal rights, technological awareness, and Sharia principles. Practically, the findings may serve as a reference for policymakers in designing consumer education programs, for Islamic economic institutions in developing Sharia-compliant digital transaction guidelines, for consumer organizations in formulating advocacy strategies, and for future researchers conducting

empirical studies on consumer protection in AI-mediated digital commerce, particularly in Muslim-majority societies such as Indonesia.

## **METHOD**

### **Research Design**

This manuscript employed a quasi-experimental pretest-posttest control group design using a synthetic dataset. No real participants were recruited, and no field observations were conducted. The modeled design used nonrandomized cluster allocation to simulate a community-based intervention setting. Analysis of covariance (ANCOVA) was used as the primary analytical method to compare post-intervention outcomes between groups while controlling for baseline scores.

### **Participants and Sampling**

The illustrative sample comprised 83 Muslim adult digital consumers aged 18 to 55 years. Forty-three were assigned to the intervention group and 40 to the control group. Eligibility criteria were: self-identification as Muslim; residence in Tasikmalaya, West Java; age of at least 18 years; at least two digital purchases or financial transactions during the previous three months; and prior exposure to at least one AI-mediated feature such as product recommendations, automated chat support, personalized advertising, risk scoring, or algorithmically curated offers. Individuals whose professional work centered on AI engineering, consumer-law practice, or Sharia legal adjudication were excluded to avoid ceiling effects in the literacy outcome. The modeled sampling strategy was purposive community recruitment followed by allocation by recruitment cluster rather than individual randomization.

### **Research Setting**

The hypothetical field setting was Tasikmalaya, West Java, Indonesia. Recruitment was modeled across community-based learning venues serving adult digital consumers with heterogeneous educational and occupational backgrounds. Tasikmalaya was selected because the research question concerns consumer protection in a Muslim-majority social setting where e-commerce, digital payments, platform financing, and algorithmically mediated product discovery are increasingly normalized. The study was not tied to a named institution because no actual fieldwork occurred.

### **Intervention and Experimental Procedure**

The intervention was a five-week Muslim Consumer Protection in AI Transactions program, delivered in one 90-minute session per week. Total contact time was 7.5 hours. Week 1 introduced common consumer-facing AI applications and taught participants to identify recommendation engines, ranking systems, chatbots, automated profiling, dynamic pricing, and algorithmic credit assessment. Week 2 covered Indonesian consumer rights, electronic-commerce duties, personal-data protection, evidence preservation, complaints, refunds, and dispute pathways. Week 3 translated Islamic economic law principles into practical transaction checks, including *ridha*, *'adl*, *amanah*, *gharar*, *tadlis*, *riba*, *khiyar*, *maṣlaḥah*, and *hifz al-mal*. Week 4 used simulated marketplace cases involving dark patterns, personalized recommendations, opaque fees, automated financing, halal claims, and chatbot responses. Week 5 required participants to apply a structured decision checklist and prepare a mock complaint supported by screenshots, transaction records, and a concise legal-Sharia justification. The control group received a short standard information sheet on general safe

online shopping without the integrated AI, legal, and Sharia training. A delayed-access model was planned so that the control group could receive the full module after posttesting.

### **Instruments and Measurement**

The primary outcome was the 24-item Muslim Consumer Protection Literacy in AI Transactions (MCPL-AI) scale, designed for this illustrative study. Items used a five-point response format from 1 (strongly disagree or unable to identify) to 5 (strongly agree or consistently able to identify), producing a total score from 24 to 120. Higher scores indicated stronger consumer-protection literacy. The scale contained four six-item domains: (1) recognition of consumer rights and redress; (2) AI transparency, data use, and algorithmic-risk awareness; (3) Sharia risk identification, including gharar, tadlis, riba, informed ridha, and khiyar; and (4) protective decision behavior, such as verifying disclosures, preserving evidence, challenging suspicious recommendations, and selecting complaint channels. Item construction drew conceptually on AI literacy dimensions that emphasize understanding, evaluation, application, and ethics (Ng et al., 2021).

The competence-oriented structure was also informed by measurement work that treats AI literacy as an assessable user capability. The legal-Sharia domains were added specifically to fit the research question (Wang et al., 2023).

### **Validity and Reliability**

For illustration, content validation was modeled through a five-expert review matrix covering consumer law, Islamic economic law, digital commerce, AI governance, and measurement. The simulated scale-level content validity index was  $S\text{-}CVI/Ave = .92$ , indicating acceptable conceptual coverage for a draft instrument. Because no real expert panel was convened, this coefficient must not be represented as empirical validation. Internal consistency calculated from the synthetic item-level responses was Cronbach's  $\alpha = .784$  at pretest and  $\alpha = .785$  at posttest. These values are adequate for an exploratory educational outcome but do not replace full construct validation. A field study should add cognitive interviews, confirmatory factor analysis, convergent validity, and test-retest reliability before the scale is used for high-stakes inference. The multidimensional measurement strategy is consistent with recent scale development for general AI literacy (Carolus et al., 2023).

Its focus on non-expert competence also aligns with Delphi-based work designed specifically for populations without specialist AI training (Laupichler et al., 2023).

### **Data Collection Procedure**

The modeled sequence consisted of eligibility screening, informed consent, baseline questionnaire completion, five weeks of intervention or control exposure, and immediate posttesting within three days after the final session. Identical MCPL-AI items were used at pretest and posttest. Demographic variables included age, sex, educational attainment, monthly digital-transaction frequency, repeated exposure to AI-mediated platform features, and prior experience submitting a digital consumer complaint. Participant codes, rather than names, were assumed for analysis. In a future real study, data collection should be conducted by personnel who are not responsible for intervention delivery when feasible to reduce expectancy bias.

### **Ethical Considerations**

No actual human-subject research was conducted for this manuscript, and no identifiable or private participant data exist. Consequently, there is no genuine ethics approval number or consent record to report. The protocol is presented as a model for future

implementation. Before field deployment, the investigators must obtain approval from an appropriate institutional ethics committee, secure written informed consent, explain voluntary participation and withdrawal rights, minimize collection of unnecessary personal data, and establish secure data-retention procedures. The study should also avoid asking participants to disclose account credentials, full transaction identifiers, or other sensitive digital information during case exercises.

## 2.9 Statistical Analysis

Analyses were modeled in two stages. First, descriptive statistics summarized participant characteristics and MCPL-AI scores. Baseline group comparability was examined using an independent-samples t-test for age and chi-square tests for categorical variables. Shapiro-Wilk tests assessed score normality within each group, while Levene's test assessed homogeneity of variance. Within-group change was tested with paired-samples t-tests. Between-group differences in gain scores were assessed using an independent-samples t-test, with Cohen's d and a bootstrap 95% confidence interval used to quantify standardized effect magnitude. Second, a one-way ANCOVA modeled posttest MCPL-AI as the dependent variable, study group as the fixed factor, and pretest MCPL-AI as the covariate. Homogeneity of regression slopes was examined by testing the group-by-pretest interaction. Partial eta squared (partial  $\eta^2$ ) quantified the adjusted group effect. All tests were two-sided with  $\alpha = .05$ . The simulated analysis retained complete data for all 83 cases, so no imputation was required.

## RESULTS AND DISCUSSIONS

### Characteristics of the Illustrative Sample

The synthetic sample contained 43 participants in the intervention group and 40 in the control group. Mean age was 33.65 years (SD = 6.92) in the intervention group and 33.03 years (SD = 7.99) in the control group,  $t(81) = 0.38$ ,  $p = .703$ . Women represented slightly more than half of both groups. Educational attainment, transaction frequency, repeated exposure to AI-mediated features, and prior complaint experience were also similar. None of the modeled baseline characteristic comparisons indicated a statistically meaningful group imbalance (Table 1).

**Table 1.** Baseline characteristics of the synthetic sample

| Characteristic                                   | Intervention (n = 43) | Control (n = 40) | p value |
|--------------------------------------------------|-----------------------|------------------|---------|
| Age, years, mean $\pm$ SD                        | 33.65 $\pm$ 6.92      | 33.03 $\pm$ 7.99 | .703    |
| Male, n (%)                                      | 18 (41.9)             | 16 (40.0)        | .863    |
| Female, n (%)                                    | 25 (58.1)             | 24 (60.0)        |         |
| Senior secondary or below, n (%)                 | 16 (37.2)             | 15 (37.5)        | .994    |
| Diploma, n (%)                                   | 9 (20.9)              | 8 (20.0)         |         |
| Bachelor's degree or higher, n (%)               | 18 (41.9)             | 17 (42.5)        |         |
| 1-3 digital transactions/month, n (%)            | 10 (23.3)             | 11 (27.5)        | .905    |
| 4-8 digital transactions/month, n (%)            | 18 (41.9)             | 16 (40.0)        |         |
| >8 digital transactions/month, n (%)             | 15 (34.9)             | 13 (32.5)        |         |
| Repeated exposure to AI-mediated features, n (%) | 37 (86.0)             | 33 (82.5)        | .657    |
| Prior digital consumer complaint, n (%)          | 7 (16.3)              | 5 (12.5)         | .625    |

Source: Synthetic data generated for illustrative modeling purposes (2026)

## Descriptive Statistics and Pretest-Posttest Change

At baseline, the intervention group had a mean MCPL-AI score of 70.30 (SD = 11.87), compared with 74.03 (SD = 10.59) in the control group. The 3.72-point baseline difference was not significant,  $t(81) = -1.50$ ,  $p = .137$ . At posttest, the intervention mean increased to 82.21 (SD = 9.69), while the control mean was 75.83 (SD = 11.97). Within the intervention group, the mean increase of 11.91 points was significant,  $t(42) = 8.07$ ,  $p < .001$ . The control group's mean increase of 1.80 points was not significant,  $t(39) = 0.98$ ,  $p = .331$  (Table 2).

**Table 2.** MCPL-AI pretest, posttest, and within-group change

| Group        | n  | Pretest mean $\pm$ SD | Posttest mean $\pm$ SD | Gain mean $\pm$ SD | Paired t, p                 |
|--------------|----|-----------------------|------------------------|--------------------|-----------------------------|
| Intervention | 43 | 70.30 $\pm$ 11.87     | 82.21 $\pm$ 9.69       | 11.91 $\pm$ 9.68   | $t(42) = 8.07$ , $p < .001$ |
| Control      | 40 | 74.03 $\pm$ 10.59     | 75.83 $\pm$ 11.97      | 1.80 $\pm$ 11.56   | $t(39) = 0.98$ , $p = .331$ |

Source: Synthetic data generated for illustrative modeling purposes (2026)

## Assumption Testing

Shapiro-Wilk tests did not indicate significant departures from normality for pretest, posttest, or gain scores in either group. Levene's tests were also non-significant for pretest, posttest, and gain scores, supporting the equal-variance assumption for the reported t-tests. For ANCOVA, the pretest-by-group interaction was not significant,  $F(1,79) = 0.05$ ,  $p = .823$ , which supported the homogeneity-of-regression-slopes assumption (Table 3).

**Table 3.** Statistical assumption checks

| Assumption test                    | Intervention                  | Control                 | Interpretation            |
|------------------------------------|-------------------------------|-------------------------|---------------------------|
| Shapiro-Wilk, pretest              | $W = .966$ , $p = .232$       | $W = .952$ , $p = .092$ | Normality not rejected    |
| Shapiro-Wilk, posttest             | $W = .982$ , $p = .707$       | $W = .976$ , $p = .559$ | Normality not rejected    |
| Shapiro-Wilk, gain                 | $W = .972$ , $p = .381$       | $W = .966$ , $p = .257$ | Normality not rejected    |
| Levene, pretest                    | $F = 0.39$ , $p = .536$       |                         | Equal variance supported  |
| Levene, posttest                   | $F = 1.78$ , $p = .186$       |                         | Equal variance supported  |
| Levene's gain                      | $F = 0.29$ , $p = .590$       |                         | Equal variance supported  |
| Pretest $\times$ group interaction | $F(1,79) = 0.05$ , $p = .823$ |                         | Parallel slopes supported |

Source: Synthetic data generated for illustrative modeling purposes (2026)

## Between-Group Difference and Effect Size

The posttest difference favored the intervention group,  $t(81) = 2.68$ ,  $p = .009$ . More importantly, the mean gain was 10.11 points greater in the intervention group than in the control group, 95% CI [5.46, 14.75],  $t(81) = 4.33$ ,  $p < .001$ . The standardized gain-score effect was Cohen's  $d = 0.95$ , with a bootstrap 95% CI of approximately [0.54, 1.44]. This magnitude indicates a large educational effect in the synthetic dataset, but it should not be interpreted as an observed population effect because the underlying cases are simulated.

**Table 4.** Between-group comparisons and effect size

| Outcome                          | Mean difference | 95% CI        | Test                      | Effect size                                 |
|----------------------------------|-----------------|---------------|---------------------------|---------------------------------------------|
| Pretest, intervention - control  | -3.72           | Not primary   | $t(81) = -1.50, p = .137$ | -                                           |
| Posttest, intervention - control | 6.38            | -             | $t(81) = 2.68, p = .009$  | -                                           |
| Gain, intervention - control     | 10.11           | 5.46 to 14.75 | $t(81) = 4.33, p < .001$  | Cohen's $d = 0.95$<br>(95% CI 0.54 to 1.44) |

Source: Synthetic data generated for illustrative modeling purposes (2026)

## ANCOVA

ANCOVA was used as the primary analysis because the intervention group began with a somewhat lower, although non-significant, pretest mean. Pretest score was a significant covariate,  $F(1,80) = 32.77, p < .001$ . After adjustment, study group remained significant,  $F(1,80) = 16.49, p < .001$ , partial  $\eta^2 = .171$ . At the overall pretest mean, the adjusted posttest mean was 83.14 (95% CI 80.33 to 85.95) for the intervention group and 74.82 (95% CI 71.91 to 77.74) for the control group. The adjusted group difference was 8.32 points, 95% CI [4.24, 12.39] (Table 5).

**Table 5.** ANCOVA for posttest MCPL-AI score

| Source          | SS      | df | F     | p      | Partial $\eta^2$ |
|-----------------|---------|----|-------|--------|------------------|
| Pretest MCPL-AI | 2770.62 | 1  | 32.77 | < .001 | -                |
| Study group     | 1394.19 | 1  | 16.49 | < .001 | .171             |
| Error           | 6764.27 | 80 |       |        |                  |

*Adjusted means: intervention = 83.14 (95% CI 80.33 to 85.95); control = 74.82 (95% CI 71.91 to 77.74). Adjusted difference = 8.32 points (95% CI 4.24 to 12.39). Synthetic data only.*

Source: Synthetic data generated for illustrative modeling purposes (2026)

The synthetic analysis supports the study hypothesis within the limits of an illustrative dataset. Participants modeled as receiving the integrated program showed a substantially larger improvement in consumer-protection literacy than those receiving standard digital-shopping information. The ANCOVA result is particularly relevant because the intervention group began with a modestly lower baseline mean. Adjustment for that difference still produced a meaningful group effect. The result should be read as a demonstration of what a plausible field pattern could look like, not as evidence that the program has already been effective in Tasikmalaya.

The modeled effect is theoretically credible because the intervention does not treat AI risk as a purely technical topic. Consumers encounter AI through commercial decisions. A recommendation, personalized price, chatbot response, or automated financing offer becomes legally important when it changes the information available to the consumer, obscures a material term, processes personal data, or limits realistic choice. Research on consumer experience with AI shows that perceived control and agency are shaped by how people understand the role of intelligent systems in the decision process (Puntoni et al., 2021).

A more recent framework similarly argues that consumer responses depend on how users interpret the broader AI-consumer interface (Ubal et al., 2024).

Explainability research indicates that disclosure is useful when it helps users construct a workable mental model of system behavior (Laato et al., 2022).

Perceived trust and acceptance also improve when explanations are understandable and causally meaningful (Shin, 2021).

The intervention operationalizes these insights by teaching consumers to ask practical questions: Was AI used? What data may have influenced the offer? Can the recommendation be challenged? Is a price or financing term transparent? What evidence should be retained?

The improvement is also consistent with prior AI literacy interventions. Learners from diverse disciplinary backgrounds can improve AI conceptual understanding without programming prerequisites (Kong et al., 2021).

A later evaluation found additional gains when course design emphasized conceptual building and learner empowerment (Kong et al., 2022).

The present model extends that pedagogical logic into a consumer-protection setting. It does not require participants to understand model architecture. Instead, it builds recognition, evaluation, and response skills. This emphasis is consistent with AI literacy frameworks that treat ethical judgment and critical evaluation as core competencies (Ng et al., 2021).

Competence-based measurement research also supports evaluating whether users can understand and apply AI-related knowledge in practical settings. In practical terms, a consumer does not need to reconstruct a recommender algorithm to recognize a disclosure problem or question a suspiciously personalized offer (Wang et al., 2023).

The dark-pattern component provides a second mechanism. Digital consumer vulnerability is partly behavioral. Interface design can convert limited attention, urgency, social proof, and default bias into commercial advantage. Experimental evidence indicates that dark patterns can increase impulsive purchasing, while interventions can reduce their influence, although effects vary by design and context (Sin et al., 2025).

Legal scholarship likewise emphasizes that manipulative interfaces can undermine authentic preference formation and meaningful consent. This matters from an Islamic economic law perspective because *ridha* should not be reduced to a mechanical click. Consent that is produced through concealment, misleading choice architecture, or material uncertainty is normatively weaker than consent formed after clear disclosure and a genuine opportunity to choose (Luguri & Strahilevitz, 2021).

Islamic economic law deepens the interpretation of consumer protection by connecting transactional legality with substantive fairness. *Gharar* is relevant when essential features, risks, or obligations remain materially uncertain. *Tadlis* is relevant when information is concealed or represented deceptively. *Khiyar* provides a conceptual basis for meaningful opportunity to reconsider or rescind under recognized conditions. The protection of wealth, *hifz al-mal*, directs attention to financial harm, while *maṣlaḥah* requires technology to advance welfare rather than merely platform efficiency. Contemporary Islamic AI ethics similarly emphasizes justice, welfare, and accountability as evaluative criteria (Elmahjub, 2023).

Recent scholarship on online trade argues that AI may produce both benefit and harm depending on whether it strengthens transparency and access or instead facilitates opaque persuasion and exploitation. Under this approach, Sharia compliance should be assessed not only at the payment-contract level but also across the informational and algorithmic environment in which the contract is formed (Asshidqi et al., 2026).

This perspective also reveals a regulatory problem. Indonesia's consumer-protection architecture contains broad rights to information, fairness, security, and compensation, and its data-protection regime creates additional obligations around personal information. However, AI-mediated commerce creates specific questions that general rules do not always answer clearly: whether consumers should be told that recommendations or prices are algorithmically personalized; what level of explanation is required for automated profiling; how responsibility should be allocated when an AI agent gives misleading transactional information; and how platforms should prevent personalized dark patterns. Indonesian comparative research has already identified the lack of explicit dark-pattern regulation as a weakness (Az-Zahra' et al., 2025). The Financial Services Authority's responsible-AI principles provide a useful sectoral direction, but consumer-facing AI governance across general e-commerce remains fragmented.

The Sharia dimension does not replace positive law. It adds normative criteria that can inform compliance design, consumer education, and institutional governance. A *maqāṣid*-based approach can strengthen consumer safeguards in digital commerce by linking legal protection to welfare, justice, and preservation of legitimate interests (Puad & Hamdi, 2025).

Regulatory scholarship on AI-enabled Islamic finance emphasizes the need to harmonize national regulation and Sharia guidance (Herlambang & Wahyudi, 2025).

Lifecycle governance proposals for generative AI stress continuous Sharia oversight rather than one-time compliance review. Together, these strands support a layered model in which positive law defines enforceable rights and duties, AI governance supplies technical accountability principles, and Islamic economic law evaluates whether the transaction process preserves justice, transparency, welfare, and legitimate consent (Zafar, 2026).

The practical implication is that consumer protection policy should include competence-building, not regulation alone. A legal right has limited preventive value when consumers cannot recognize the triggering practice, preserve evidence, or identify the relevant complaint route. Community-level training can therefore complement platform duties and enforcement. In Tasikmalaya, an empirically tested version of the program could be delivered through universities, Islamic economic institutions, consumer organizations, community learning centers, and digital-finance outreach programs. The curriculum should remain scenario-based. Cases involving AI recommendations, personalized discounts, pay-later offers, chatbot misstatements, halal claims, consent interfaces, and refund obstruction are more likely to produce transferable judgment than abstract lectures on statutory provisions.

The theoretical implication concerns the construct of consumer literacy itself. Standard digital literacy focuses on access and functional competence, while AI literacy emphasizes understanding, use, evaluation, and ethics. Consumer-protection literacy adds legal agency, and Islamic economic law adds normative transaction assessment. Integrating these domains produces a broader concept of protected digital agency: the capacity to understand algorithmically mediated choice, evaluate legal and Sharia risks, make an informed decision, and seek redress when necessary. This construct may be useful beyond Indonesia, especially in Muslim-majority markets where AI-driven commerce grows faster than consumer-facing governance and Sharia guidance.

### **Strengths and Limitations**

The manuscript has several design strengths. It specifies a clear comparison group, uses pretest adjustment rather than relying only on raw gain scores, reports standardized and adjusted

effect sizes, and integrates legal, technological, behavioral, and Sharia dimensions into one intervention. The outcome instrument is also designed around actionable consumer competencies rather than general attitudes toward technology. The use of scenario-based learning is supported by evidence that consumer vulnerability often emerges during concrete interface interactions rather than through abstract lack of knowledge.

The limitations are substantial and must be explicit. First, all data are synthetic. The reported p-values, confidence intervals, reliability coefficients, and effect sizes demonstrate statistical coherence only and provide no empirical evidence about Muslim consumers in Tasikmalaya. Second, the modeled design is quasi-experimental, so even a future field replication would remain vulnerable to selection effects unless allocation procedures are strengthened. Third, the MCPL-AI scale is newly proposed and requires real psychometric validation. Fourth, immediate posttesting cannot establish whether learning persists or changes actual marketplace behavior. Fifth, a single-location study would limit external validity across regions with different levels of digital access, educational attainment, platform use, and Islamic economic literacy. Finally, the legal and technological environment changes quickly. Future studies should update case materials when platform practices, AI capabilities, OJK rules, consumer law, data-protection implementation, or relevant Sharia guidance changes.

## CONCLUSION

This Scopus-style manuscript modeled a quasi-experimental evaluation of an integrated consumer protection literacy program for Muslim digital consumers in Tasikmalaya. Within the synthetic dataset, participants who received five weeks of combined AI literacy, consumer law, and Islamic economic law instruction demonstrated greater improvement in MCPL-AI scores than those in the control group. The adjusted between-group difference remained statistically significant after controlling for baseline literacy, with a large standardized effect for gain scores and a moderate-to-large ANCOVA effect. However, these results were illustrative and should not be interpreted as empirical field findings.

Substantively, the study proposed that Muslim consumer protection in AI-based digital transactions should be grounded in three interconnected dimensions: enforceable consumer and data protection rights, practical AI literacy, and Islamic economic principles that safeguard informed consent, fairness, transparency, and the protection of wealth. In practice, regulators, digital platforms, consumer organizations, universities, and Islamic economic institutions should evaluate concise scenario-based educational interventions alongside stronger mechanisms for platform accountability. Future research should collect empirical data from Muslim consumers in Tasikmalaya, validate the MCPL-AI scale, apply more robust allocation procedures where feasible, incorporate behavioral outcome measures, and include delayed follow-up assessments. Such research would help determine whether the patterns observed in the synthetic dataset can be replicated under real-world consumer conditions.

## REFERENCE

Areiqat, A. Y., Hamdan, A., Alheet, A. F., & Alareeni, B. (2020). Impact of artificial intelligence on E-commerce development. *International Conference on Business and Technology*, 571–578.

- Asshidqi, A. F., Muhajir, & Riyantoro, S. F. (2026). Masalah dan mafsadah pemanfaatan artificial intelligence dalam jual beli online: Perspektif hukum ekonomi syariah. *Kartika: Jurnal Studi Keislaman*, 6(2), 3249–3267. <https://doi.org/10.59240/kjsk.v6i2.724>
- Az-Zahra', P. N., Nurlaily, & Agustianto. (2025). Regulating dark patterns in Indonesian e-commerce: Comparative lessons from South Korea and the EU. *Journal of Judicial Review*, 27(2), 421-452. <https://doi.org/10.37253/jjr.v27i2.11304>
- Carolus, A., Koch, M. J., Straka, S., Latoschik, M. E., & Wienrich, C. (2023). MAILS: Meta AI literacy scale: Development and testing of an AI literacy questionnaire based on well-founded competency models and psychological change- and meta-competencies. *Computers in Human Behavior: Artificial Humans*, 1(2), 100014. <https://doi.org/10.1016/j.chbah.2023.100014>
- Chouhan, M. A. (2025). E-Commerce and Digital Transactions in the Age of Artificial Intelligence. *Center for Development Economic*, 12(26), 30–42.
- Davenport, T., Guha, A., Grewal, D., & Bressgott, T. (2020). How artificial intelligence will change the future of marketing. *Journal of the Academy of Marketing Science*, 48, 24–42. <https://doi.org/10.1007/s11747-019-00696-0>
- De Conca, S. (2023). The present looks nothing like the Jetsons: Deceptive design in virtual assistants and the protection of the rights of users. *Computer Law & Security Review*, 51, 105866. <https://doi.org/10.1016/j.clsr.2023.105866>
- Elmahjub, E. (2023). Artificial intelligence (AI) in Islamic ethics: Towards pluralist ethical benchmarking for AI. *Philosophy & Technology*, 36, 73. <https://doi.org/10.1007/s13347-023-00668-x>
- Fadlurrahman, I., & Fikrianihayah, A. N. (2022). Consumer protection in e-commerce transactions through the Shopee application according to Sharia economic law. *Al-Muamalat: Jurnal Ekonomi Syariah*, 9(1), 14–19. <https://doi.org/10.15575/am.v9i1.13920>
- Firmansyah, E. A., & Harsanto, B. (2022). Islamic fintech research: Systematic review using mainstream databases. *Etikonomi*, 21(2), 355-368. <https://doi.org/10.15408/etk.v21i2.24602>
- Ghaly, M. (2024). What makes work 'good' in the age of artificial intelligence (AI)? Islamic perspectives on AI-mediated work ethics. *The Journal of Ethics*, 28, 429–453. <https://doi.org/10.1007/s10892-023-09456-3>
- Herlambang, H., & Wahyudi, A. (2025). Digital transformation of Sharia finance: Challenges of harmonizing regulations and fatwas in the application of artificial intelligence in Sharia fintech financing in Indonesia. *Jurnal Iqtisaduna*, 11(2), 499-511. <https://doi.org/10.24252/iqtisaduna.v11i2.62238>
- Hollebeek, L. D., Menidjel, C., Sarstedt, M., Jansson, J., & Urbonavicius, S. (2024). Engaging consumers through artificially intelligent technologies: Systematic review, conceptual model, and further research. *Psychology & Marketing*, 41(4), 880–898. <https://doi.org/10.1002/mar.21957>
- Huang, M.-H., & Rust, R. T. (2021). A strategic framework for artificial intelligence in marketing. *Journal of the Academy of Marketing Science*, 49, 30–50. <https://doi.org/10.1007/s11747-020-00749-9>
- Kong, S.-C., Cheung, W. M.-Y., & Zhang, G. (2021). Evaluation of an artificial intelligence literacy course for university students with diverse study backgrounds. *Computers and Education: Artificial Intelligence*, 2, 100026. <https://doi.org/10.1016/j.caeai.2021.100026>
- Kong, S.-C., Cheung, W. M.-Y., & Zhang, G. (2022). Evaluating artificial intelligence literacy courses for fostering conceptual learning, literacy, and empowerment in university students: Refocusing on conceptual building. *Computers in Human Behavior Reports*, 7, 100223. <https://doi.org/10.1016/j.chbr.2022.100223>

- Laato, S., Tiainen, M., Najmul Islam, A. K. M., & Mäntymäki, M. (2022). How to explain AI systems to end users: A systematic literature review and research agenda. *Internet Research*, 32(7), 1–31. <https://doi.org/10.1108/INTR-08-2021-0600>
- Laupichler, M. C., Aster, A., & Raupach, T. (2023). Delphi study for the development and preliminary validation of an item set for the assessment of non-experts' AI literacy. *Computers and Education: Artificial Intelligence*, 4, 100126. <https://doi.org/10.1016/j.caeai.2023.100126>
- Luguri, J., & Strahilevitz, L. J. (2021). Shining a light on dark patterns. *Journal of Legal Analysis*, 13(1), 43–109. <https://doi.org/10.1093/jla/laaa006>
- Malasari, R., Darwis, M., & Said, Z. (2026). Consumer protections in digital transactions: A study of Indonesian e-commerce through positive and Islamic law. *Milkiyah: Jurnal Hukum Ekonomi Syariah*, 5(1), 13–22. <https://doi.org/10.46870/milkiyah.v5i1.1638>
- Mills, S., Costa, S., & Sunstein, C. R. (2023). AI, behavioral science, and consumer welfare. *Journal of Consumer Policy*, 46, 387–400. <https://doi.org/10.1007/s10603-023-09547-6>
- Narayanan, A., Mathur, A., Chetty, M., & Kshirsagar, M. (2020). Dark patterns: Past, present, and future. *Communications of the ACM*, 63(9), 42–47. <https://doi.org/10.1145/3409116>
- Ng, D. T. K., Leung, J. K. L., Chu, S. K. W., & Qiao, M. S. (2021). Conceptualizing AI literacy: An exploratory review. *Computers and Education: Artificial Intelligence*, 2, 100041. <https://doi.org/10.1016/j.caeai.2021.100041>
- Puad, N. A. M., & Hamdi, A. S. (2025). Maqasid Shariah and consumer protection in e-commerce: Strengthening legal safeguards in Indonesia's digital economy. *International Journal of Islamic Economics and Finance Research*, 1, 64–75. <https://doi.org/10.53840/ijiefer222>
- Puntoni, S., Reczek, R. W., Giesler, M., & Botti, S. (2021). Consumers and artificial intelligence: An experiential perspective. *Journal of Marketing*, 85(1), 131–151. <https://doi.org/10.1177/0022242920953847>
- Shin, D. (2021). The effects of explainability and causability on perception, trust, and acceptance: Implications for explainable AI. *International Journal of Human-Computer Studies*, 146, 102551. <https://doi.org/10.1016/j.ijhcs.2020.102551>
- Singh, V., Sharma, R., & Rawat, S. R. (2026). Charting the knowledge landscape of artificial intelligence driven personalization and consumer purchase intention in e-commerce: A bibliometric analysis. *International Journal of Business and Management (IJBM)*, 5(1), 1041–1070.
- Sin, R., Harris, T., Nilsson, S., & Beck, T. (2025). Dark patterns in online shopping: Do they work and can nudges help mitigate impulse buying? *Behavioral Public Policy*, 9(1), 61–87. <https://doi.org/10.1017/bpp.2022.11>
- Song, X., Yang, S., Huang, Z., & Huang, T. (2019). The application of artificial intelligence in electronic commerce. *Journal of Physics: Conference Series*, 1302(3), 32030.
- Ubal, V. O., Lisjak, M., & Mende, M. (2024). Cracking the consumers' code: A framework for understanding the artificial intelligence-consumer interface. *Current Opinion in Psychology*, 58, 101832. <https://doi.org/10.1016/j.copsyc.2024.101832>
- Vassøy, B., & Langseth, H. (2024). Consumer-side fairness in recommender systems: A systematic survey of methods and evaluation. *Artificial Intelligence Review*, 57, 101. <https://doi.org/10.1007/s10462-023-10663-5>
- Wang, B., Rau, P.-L. P., & Yuan, T. (2023). Measuring user competence in using artificial intelligence: Validity and reliability of artificial intelligence literacy scale. *Behavior & Information Technology*, 42(9), 1324–1337. <https://doi.org/10.1080/0144929X.2022.2072768>

- Widyastuti, E. S., Kamila, T. R., & Saputra, P. A. A. (2022). Perlindungan konsumen dalam transaksi e-commerce: Suatu perspektif hukum Islam. *Milkiyah: Jurnal Hukum Ekonomi Syariah*, 1(2), 43–50. <https://doi.org/10.46870/milkiyah.v1i2.208>
- Zafar, M. B. (2026). Shariah governance standard on generative AI for Islamic financial institutions. *Journal of Islamic Accounting and Business Research*. Advance online publication. <https://doi.org/10.1108/JIABR-03-2025-0141>