



## Original Research

# Interpretable Models over Customer 360 Data for Stakeholder Trust and Governance

Muhammad Saad Khan<sup>1</sup>, Adeel Raza Siddiqui<sup>2</sup> and Hamza Nadeem Qureshi<sup>3</sup>

<sup>1</sup>Rawalpindi College of Management Sciences, Department of Business Administration, 54 Murree Road, Rawalpindi 46000, Pakistan.

<sup>2</sup>Karachi Institute of Finance and Economics, Department of Accounting and Finance, 22 Shahrah-e-Faisal, Karachi 75530, Pakistan.

<sup>3</sup>Lahore School of Business Innovation, Department of Management Studies, 81 Ferozepur Road, Lahore 54000, Pakistan.

## Abstract

Organisations in many sectors increasingly rely on integrated customer data platforms that consolidate operational, transactional, behavioural, and third-party information into so-called Customer 360 representations. These holistic profiles enable predictive and prescriptive analytics for marketing, service, risk, and compliance use cases. At the same time, regulatory expectations, internal model risk standards, and public scrutiny place growing emphasis on interpretability, governance, and demonstrable fairness of data-driven decisions derived from such systems. This paper examines interpretable modelling approaches over Customer 360 data with a focus on their suitability for supporting stakeholder trust and governance practices across business, technical, and oversight functions. The discussion characterises Customer 360 data along dimensions of heterogeneity, temporal structure, sparsity, and data quality, and analyses how these characteristics interact with transparency requirements and explanation methods. The paper surveys model classes that are inherently interpretable, as well as post-hoc explanation techniques, and assesses their strengths and limitations under typical Customer 360 workloads, including propensity scoring, retention prediction, next-best-action ranking, and early-warning signals. Particular attention is given to the alignment between model explanations, organisational decision processes, and the information needs of different stakeholders such as product owners, legal and compliance teams, auditors, and affected customers. The paper further sketches an architectural perspective on integrating interpretable models into Customer 360 platforms, touching on data lineage, policy enforcement, monitoring, and documentation artefacts. An evaluation framework based on quantitative and qualitative criteria is proposed to reason about model performance, stability, comprehensibility, and governance readiness in a coherent manner. The paper concludes with observations on practical trade-offs and open questions in operationalising interpretable Customer 360 modelling at scale.

## 1. Introduction

Customer analytics has evolved from isolated campaign reporting and simple segmentation toward integrated, real-time decisioning across the customer lifecycle [1]. This evolution is driven by the emergence of Customer 360 platforms that combine data from transactional systems, digital touchpoints, physical interactions, and external sources into unified customer views. Such views typically encode demographic attributes, product holdings, channel preferences, behavioural signals, and longitudinal event histories. Models built on top of these representations influence decisions such as whom to contact, what offer to propose, how to prioritise service queues, and when to intervene in potential churn or risk situations. As these decisions increasingly affect access to services and shape customer experience, interpretability of the underlying models and transparency of the decision logic become central concerns.

Interpretability is not a single, universally defined property, but rather a collection of qualities related to the ability of stakeholders to form an adequate mental model of how predictions are produced [2]. For some stakeholders, interpretability may mean simple global relationships between features and outcomes

that can be articulated in natural language. For others, it may refer to local explanations for individual decisions, such as identifying the main contributing factors to a propensity score. In regulated domains, interpretability can also be tied to specific obligations, for example providing reasons for adverse decisions, demonstrating non-discrimination, or evidencing appropriate use of personal data. In the context of Customer 360 data, these requirements intersect with complex data pipelines, heterogeneous feature spaces, and automated decision flows, which introduce distinct technical and organisational challenges.

This paper focuses on interpretable models over Customer 360 data with a particular emphasis on stakeholder trust and governance [3]. Rather than framing interpretability solely as a technical property of algorithms, the discussion treats it as a socio-technical attribute of an entire decisioning system. Models operate within data pipelines, orchestration layers, and oversight processes, and their explanations become meaningful only relative to these surrounding structures. For instance, a feature importance plot might be technically accurate but practically unhelpful if it refers to engineered features whose provenance and semantics are not well understood outside the data science team. Similarly, a model card or documentation artefact may satisfy internal templates yet remain misaligned with the questions auditors or regulators actually ask.

Customer 360 data introduces several domain-specific considerations for interpretability. The feature space is often wide and mixed-type, combining categorical variables, numeric metrics, text-derived signals, and time-aggregated behavioural indicators [4]. Many features are generated through nested aggregation or embedding pipelines, sometimes from semi-structured or unstructured sources. Temporal dynamics are central, as recent events typically carry different predictive weight from older ones, and as seasonality or campaign effects may modulate behaviour. In addition, Customer 360 models are commonly deployed in high-throughput environments, where explanations must be produced with constrained latency and integrated seamlessly into decisioning workflows and user interfaces. These factors influence which modelling approaches are feasible and how explanation techniques can be applied in practice.

Stakeholder trust in Customer 360 models depends not only on the ability to explain individual predictions, but also on consistent governance practices across the model lifecycle [5]. Governance encompasses model inventory management, approval workflows, versioning, monitoring for drift and performance degradation, access control over sensitive features, and traceable documentation of changes. Interpretable models can facilitate these activities by making assumptions explicit and exposing decision logic in forms amenable to review, challenge, and audit. However, interpretable approaches may also introduce trade-offs in terms of predictive performance, robustness, or operational complexity. Consequently, practitioners often face decisions about where along a continuum between full transparency and black-box performance their applications should reside, given regulatory expectations, business objectives, and system constraints.

The remainder of this paper explores these issues in a structured manner. It characterises the structure and governance context of Customer 360 data, outlines modelling paradigms with interpretability properties, discusses stakeholder-oriented trust and governance requirements, and sketches architectural patterns for deploying interpretable models within Customer 360 platforms [6]. An evaluation perspective is introduced to reason about quantitative and qualitative metrics relevant to interpretability and governance, followed by illustrative scenario-driven analyses. The paper does not attempt an exhaustive survey of algorithms, but instead focuses on the interaction between modelling choices, data characteristics, and organisational processes that shape how interpretability is achieved and used.

## 2. Customer 360 Data Landscape and Challenges

Customer 360 data generally arises from the integration of multiple upstream systems, each optimised for particular operational processes. Core transaction systems capture product usage, billing, and payments. Customer relationship management tools record interactions with sales and service teams [7]. Digital platforms provide clickstream logs, app events, and web forms. Contact centre platforms store

call metadata and transcripts. External data providers contribute demographic information, geospatial attributes, and credit or risk scores. The Customer 360 platform acts as a consolidation layer, reconciling identities, resolving conflicts across sources, and constructing harmonised entities that represent individual customers or households.

From a modelling perspective, the resulting data exhibits heterogeneity at several levels [8]. Attribute modalities vary, spanning binary flags, high-cardinality categorical variables, continuous measures with different scales, and text or image-derived representations. Temporal granularity is uneven; some sources update in near real time, while others refresh periodically. Missingness patterns are systematic, reflecting business processes, customer behaviour, and technical gaps rather than random noise. These properties complicate both classical feature engineering and newer representation learning approaches. Interpretable modelling over such data must contend not only with the complexity of the predictive relationships, but also with the need to articulate them in ways that remain comprehensible despite underlying heterogeneity.

Identity resolution is a foundational challenge in Customer 360 platforms and influences interpretability indirectly [9]. Models assume that input features indeed refer to the same real-world entity. When multiple identifiers across channels and systems are linked through deterministic rules or probabilistic matching, residual ambiguity can persist. For example, one physical person may appear as several customer records, or multiple individuals may be linked to a single account. These imperfections affect model training data and, consequently, explanations derived from model outputs. If predicted behaviour is partially driven by mismatched or aggregated histories, local explanations may attribute importance to features that in fact pertain to different underlying entities [10]. Clear governance of identity resolution logic and visibility of its quality characteristics become relevant for interpreting model results over Customer 360 data.

Feature engineering in Customer 360 contexts often involves constructing aggregates over event streams, such as counts, rates, recency measures, and rolling statistics. These engineered features are typically given encoded names and managed in feature stores. While such representations can be powerful predictors, they introduce an additional abstraction layer between source events and model inputs. Interpretable modelling requires that this layer remain navigable. Stakeholders seeking to understand why a particular churn risk score is high may need to trace from the model input feature, for instance a normalised complaint frequency in a certain window, back to the underlying calls, messages, or tickets that contributed to that aggregate [11]. Without clear lineage and semantic documentation, feature-level explanations risk becoming opaque, even when the model family itself is transparent.

Data quality and bias issues in Customer 360 repositories also play a central role. Missing or inconsistent values may correlate with socio-economic variables, usage patterns, or channel access, thereby affecting the distribution of model inputs across customer segments. If these patterns remain hidden, interpretation of model outputs can be misleading. For instance, certain features may appear predictive of default risk largely because they are proxies for under-represented channels or legacy products [12]. In such cases, interpretability techniques that highlight these features may be accurate in describing the model's behaviour, yet they may fail to reveal the underlying structural data issues. Consequently, interpretability over Customer 360 data must sometimes be complemented by exploratory analyses of data coverage, lineage, and feature distributions to contextualise explanations.

Another challenge stems from the temporal dimension of Customer 360 data. Many customer-related phenomena evolve over time, with interventions, campaigns, and external events influencing behaviour. Models that ingest static snapshots at inference time may rely on lagged or aggregated variables that smooth over temporal patterns [13]. Conversely, sequential models that explicitly represent event sequences may be less transparent because they distribute predictive importance across numerous time steps. Explanation methods for temporal models attempt to attribute importance to windows, events, or recurrent states, yet their outputs can be difficult to align with business narratives. Stakeholders may find it more intuitive to reason about simple recency and frequency aggregates than about abstract temporal attention weights or latent states. Designing interpretable temporal representations that balance fidelity and simplicity is therefore an active area for Customer 360 modelling.

Finally, Customer 360 platforms are typically shared assets across multiple lines of business, meaning that features and models are reused in different contexts. A feature engineered for marketing propensity may later be employed in a service prioritisation model [14]. This reuse has implications for interpretability and governance because explanations must reflect context. The same feature might be acceptable and meaningful in one decision process but sensitive or misleading in another. Governance processes must track where features are used, with what transformations, and under which policy constraints. Interpretable models can support this by making feature contributions explicit, but the underlying meta-data and cataloguing infrastructure are equally important. Together, these considerations define the landscape within which interpretable Customer 360 models operate [15].

### 3. Interpretable Modelling Approaches for Customer 360

Interpretable models can be roughly grouped into inherently interpretable model families and post-hoc explanation techniques applied to more complex predictors. In Customer 360 settings, both categories are relevant, and practitioners frequently combine them. Inherently interpretable models include linear and logistic regression with carefully curated features, generalised additive models with potentially non-linear but shape-constrained components, decision trees and rule lists with limited depth or complexity, and scoring systems with integer-weighted factors. These models offer direct mappings from inputs to outputs via structures that can be inspected and articulated in domain language. Their simplicity can reduce the cognitive burden on stakeholders responsible for understanding and governing model behaviour.

Linear and logistic models remain widely used in customer analytics because their coefficients admit a straightforward interpretation as marginal effects under certain assumptions [16]. When features are normalised appropriately and multicollinearity is controlled, stakeholders can interpret coefficients as indicating positive or negative associations between features and target outcomes. However, Customer 360 feature spaces are often highly correlated and non-linear, which can render naive coefficient interpretations unreliable. Incorporating interaction terms and transformations can improve predictive performance but reduces transparency. To mitigate this, practitioners sometimes restrict the set of features for linear models to a small, carefully selected subset of variables with clear semantics and low redundancy, using more flexible models for auxiliary tasks or to explore non-linear patterns that might later be encoded in interpretable forms.

Generalised additive models represent a compromise between flexibility and interpretability by modelling the prediction as a sum of univariate or low-dimensional component functions of individual features or feature groups [17]. Shape constraints such as monotonicity or convexity can be imposed to align with domain knowledge or policy expectations. In Customer 360 applications, additive models allow stakeholders to examine how each feature contributes to the prediction across its range, often via partial dependence plots that are relatively easy to understand. At the same time, care must be taken in handling high-cardinality categorical features or interaction effects that are not well captured by additive structures. Feature grouping and encoding schemes need to be designed so that component functions remain interpretable, for example by aggregating rare categories or constructing composite indices that capture related behaviours.

Tree-based models occupy a prominent position because decision paths can be presented as human-readable rules [18]. Shallow trees with limited branching factors may serve as primary decision engines in certain Customer 360 use cases where simplicity is prioritised. For more complex tasks, ensembles such as gradient boosted trees and random forests often provide strong predictive performance but at the cost of reduced inherent transparency. Post-hoc explanations, such as feature importance metrics, partial dependence analyses, and instance-level attribution methods, are frequently applied to these ensembles. For example, perturbation-based methods and additive explanation frameworks assign contributions to features for individual predictions, enabling local explanations even when the ensemble structure is too large to inspect directly. Nevertheless, the interpretability of such explanations is conditional on the stability and consistency of attribution across similar instances and time periods.

Scorecards and points-based systems remain common in credit risk and can be adapted to broader Customer 360 contexts [19]. They express the prediction as a sum of discrete points assigned to ranges or categories of features, with the final score mapped to a probability or risk band. The discrete and additive structure can be more intuitive for stakeholders and customers than continuous coefficients. Designing such scorecards typically involves discretisation, binning, monotonic ordering, and manual review to ensure that the resulting system respects business constraints and legal requirements. In Customer 360 applications, scorecards can be used for churn propensity, complaint risk, or eligibility assessments where clear reasoning is needed. However, constructing and maintaining scorecards across wide feature spaces demands disciplined governance and automation support to avoid ad hoc modifications that erode validity [20].

Post-hoc model-agnostic explanation methods extend interpretability to complex models such as deep networks or large ensembles. Local surrogate modelling approximates the behaviour of a complex model around a specific instance by fitting an interpretable model on perturbed samples in its neighbourhood. Attribution methods based on additive feature contributions, path integrals, or gradient signals generate explanation vectors that sum to the model output difference relative to a baseline. In Customer 360 environments, these techniques can provide fine-grained insight into why an individual customer received a particular score or treatment recommendation. However, stakeholders must understand that such explanations describe the behaviour of the surrogate or attribution framework under specific assumptions, not necessarily causal relations in the underlying data. Communicating the limitations and stability properties of post-hoc explanations becomes part of governance [21].

For certain Customer 360 tasks, concept-based models provide an alternative interpretability mechanism. Instead of learning directly from raw or low-level features, models are structured around intermediate concepts that correspond to human-understandable constructs such as engagement, satisfaction, financial resilience, or service burden. These concepts may be defined through supervised, semi-supervised, or expert-designed aggregation functions over raw features. The final predictive model then operates on the concept space, often with a simple functional form. This arrangement allows explanations to be expressed in terms of concept contributions, which can be more meaningful to business stakeholders than individual low-level feature weights [22]. Designing stable, well-defined concepts requires collaboration between data scientists and domain experts and introduces its own governance requirements, such as concept versioning and validation.

Selecting a modelling approach for interpretability in Customer 360 applications therefore involves balancing several factors. Predictive performance, computational efficiency, and robustness to data drift must be weighed against transparency of decision logic, ease of explanation, and compatibility with stakeholder mental models. In many organisations, a layered approach emerges, where more complex models serve exploratory or advisory roles, while final decision models adhere to stronger interpretability constraints. Ensembles of interpretable models, or model cascades that fall back to simple rules under certain conditions, can also be employed. Regardless of the specific choices, interpretability is most effective when considered early in the modelling process, guiding feature engineering, model selection, and validation rather than being treated as an afterthought addressed solely through post-hoc techniques [23].

| Concept               | Description                             | Customer 360 Relevance                         |
|-----------------------|-----------------------------------------|------------------------------------------------|
| Unified Profiles      | Consolidated multi-source customer data | Enables holistic modelling across lifecycle    |
| Decision Influence    | Model-driven decisions across domains   | Supports targeting, risk scoring, intervention |
| Interpretability Need | Transparency of decision drivers        | Required for trust, governance, compliance     |

| Challenge              | Underlying Cause                         | Impact on Interpretability                      |
|------------------------|------------------------------------------|-------------------------------------------------|
| Heterogeneous Features | Mixed categorical, numeric, text signals | Increases complexity of explanation mapping     |
| Temporal Structure     | Uneven data granularity and dynamics     | Makes attribution alignment difficult for users |
| Identity Resolution    | Multi-source identifier ambiguity        | Can distort feature-level explanations          |

| Model Type        | Key Characteristics                   | Interpretability Consideration                |
|-------------------|---------------------------------------|-----------------------------------------------|
| Linear Models     | Coefficients reflect marginal effects | Sensitive to correlation and scaling issues   |
| Additive Models   | Feature-wise component functions      | More intuitive but limited in interactions    |
| Tree-Based Models | Rule-based decision paths             | Transparent when shallow, opaque in ensembles |

| Governance Area  | Requirement                 | Interpretability Link                                |
|------------------|-----------------------------|------------------------------------------------------|
| Model Monitoring | Track behaviour over time   | Explanation drift indicates underlying issues        |
| Access Control   | Manage sensitive features   | Explanation filtering needed for different audiences |
| Documentation    | Clear, consistent artefacts | Bridges technical and non-technical stakeholders     |

4. Trust, Governance, and Stakeholder Requirements

Stakeholder trust in Customer 360 models is shaped by expectations that differ across roles. Business owners typically seek assurance that models capture relevant business logic, behave consistently under anticipated scenarios, and do not inadvertently encode undesirable biases. They may evaluate models through counterfactual scenarios, asking how predictions change when certain features are varied in ways that reflect realistic interventions. Data scientists focus on statistical performance, stability, and resilience to changes in data distributions. Risk and compliance teams examine whether models meet policy and regulatory constraints, such as avoiding direct or indirect use of protected characteristics and providing sufficient reasoning for high-impact decisions [24]. External stakeholders, including customers and regulators, may require concise explanations for specific outcomes and evidence that governance processes are in place.

Governance frameworks for models over Customer 360 data typically encompass several lifecycle stages. During design and development, governance requires documentation of the model’s intended purpose, scope, and limitations, as well as the provenance and preparation of training data. Interpretability considerations include explicit statements about which features can influence outcomes and in what manner, supported by rationale linked to business context or legal obligations. Review boards or approval committees often assess whether the proposed model and feature set align with organisational policies, including restrictions on sensitive data and requirements for explainability to affected individuals [25]. Formal constraints such as monotonicity with respect to certain variables may be mandated to ensure that model behaviour adheres to intuitive expectations.

In deployment, governance extends to monitoring performance and behaviour over time. For Customer 360 models, monitoring typically tracks prediction distributions, performance metrics across segments, feature value distributions, and explanation patterns. Changes in the distribution of explanation outputs, such as shifts in which features most frequently contribute to high scores, may signal data drift, changes in upstream processes, or emerging biases. Incorporating explanation-based monitoring into governance allows stakeholders to detect and investigate such changes before they translate into problematic decision outcomes. However, this approach requires careful definition of summary statistics over explanation outputs and thresholds for triggering alerts, along with procedures for investigation and remediation [26].

Access control and data minimisation are further governance dimensions that intersect with interpretability. Customer 360 platforms often contain sensitive or regulated information, including identifiers, financial details, and potentially inferred attributes. Interpretable models expose relationships

between features and outcomes, which can increase transparency but also risk revealing sensitive associations if explanations are not properly filtered. Governance mechanisms must therefore define which features may appear in explanations and at what level of granularity. For example, internal reviewers might see detailed feature attributions, while customer-facing interfaces present higher-level explanations framed in terms of legitimate business factors [27]. This layered explanation strategy allows the organisation to maintain both transparency and data protection obligations.

An additional requirement arises from the need to reconcile interpretability with fairness and non-discrimination principles. Explanations can highlight when certain features contribute strongly to decisions, which may prompt questions about whether these features act as proxies for protected characteristics or under-served groups. Governance processes should support systematic analysis of such relationships, using statistical fairness metrics, sensitivity analyses, and scenario testing. Interpretable models can facilitate this analysis by making it easier to inspect functional forms and decision boundaries. Nevertheless, interpretability does not automatically guarantee fairness; it simply provides tools that, when embedded into governance routines, allow stakeholders to reason about fairness more concretely [28]. Responsibility for aligning models with ethical and legal standards remains with the organisation, not with the interpretability methods themselves.

Finally, trust and governance depend on communication artefacts that bridge technical and non-technical perspectives. Model documentation, or model cards, describe data sources, model structure, performance metrics, explanation methods, and usage constraints in consistent formats. In Customer 360 contexts, such artefacts may also describe data lineage from upstream systems, known data quality issues, and dependencies on other models or rules. Training and guidance for stakeholders who interact with model outputs and explanations are important, so that they understand appropriate interpretation and limitations [29]. For example, frontline staff using propensity scores in customer conversations may need practical guidance on how to incorporate explanations into dialogue without overstating certainty or revealing sensitive intermediate features. In this way, interpretability becomes a property that is jointly produced by models, platforms, governance processes, and human practices.

| Stakeholder Group   | Primary Expectation                     | Interpretability Need                     |
|---------------------|-----------------------------------------|-------------------------------------------|
| Business Owners     | Stable, logic-aligned behaviour         | Clear links between features and outcomes |
| Data Scientists     | Performance, robustness, drift handling | Reliable attribution, stable patterns     |
| Risk and Compliance | Policy adherence, non-discrimination    | Evidence of reasoning and constraints     |
| External Actors     | Justified model outputs                 | Concise, high-level explanations          |

| Governance Stage          | Core Requirement                    | Interpretability Connection            |
|---------------------------|-------------------------------------|----------------------------------------|
| Design and Development    | Purpose, scope, feature rationale   | Clarity on what influences predictions |
| Approval                  | Policy and sensitivity checks       | Ensures acceptable feature usage       |
| Deployment                | Monitoring and behavioural tracking | Explanation patterns reveal drift      |
| Post-Deployment Oversight | Auditability and documentation      | Explanations support regulatory review |

| Governance Theme        | Risk Addressed                | Explanation Implication                       |
|-------------------------|-------------------------------|-----------------------------------------------|
| Access Control          | Exposure of sensitive signals | Filter or aggregate attribution outputs       |
| Data Minimisation       | Avoiding unnecessary features | Restrict appearance of sensitive factors      |
| Fairness Requirements   | Proxy or bias detection       | Examine contribution patterns across segments |
| Communication Artefacts | Bridging technical gaps       | Provide audience-specific explanation forms   |

## 5. System Architecture and Implementation Considerations

Realising interpretable models over Customer 360 data in practice requires architectural support across data, modelling, and application layers. At the data layer, the Customer 360 platform typically offers entity resolution services, feature pipelines, and storage engines capable of handling both batch and streaming workloads. Interpretable modelling benefits when this layer exposes rich meta-data about feature definitions, provenance, transformation logic, data quality indicators, and applicable policy constraints [30]. A well-designed feature store becomes a central component, storing not only feature values but also human-readable descriptions, lineage graphs, and governance attributes such as sensitivity levels and allowed usage contexts. This meta-data can be consumed by explanation services to present feature attributions in terms that align with business understanding.

At the modelling layer, architectural choices include training infrastructure, model registry, and explainability services. Training pipelines ingest feature data, apply sampling or weighting strategies, train one or more candidate models, and evaluate their performance according to predefined metrics. For interpretable modelling, additional evaluation stages can compute measures of explanation stability, sparsity of feature attributions, and compliance with monotonic or rule-based constraints. Candidate models, along with their metrics and artefacts such as parameter sets and explanation summaries, are stored in a model registry [31]. The registry tracks versions, approvals, deployment status, and associated documentation. It can also enforce that only models that meet certain interpretability criteria, as defined by governance policies, are eligible for deployment in specific decision flows.

Explainability services often operate as shared components invoked by multiple applications. For inherently interpretable models, these services may simply retrieve and format model parameters, rules, or scorecard tables for display. For complex models, they may execute post-hoc explanation algorithms, potentially leveraging approximation or caching to meet latency budgets [32]. In Customer 360 environments, explanation requests can be frequent, especially when models underpin high-volume decisioning such as real-time offer selection. Architectural designs may therefore offload some explanation computations to pre-processing steps or batch jobs, for example by pre-computing explanations for segments of customers or for typical scenarios. Caching schemes must be reconciled with data freshness requirements and access control policies; explanations that incorporate sensitive features must not be inadvertently cached in contexts where they could be exposed more broadly than intended.

At the application layer, Customer 360 models integrate into decisioning engines, campaign management tools, service platforms, and analytics dashboards. User interfaces must be designed to present predictions and explanations in ways that match the needs and capacities of their audiences. For instance, a risk dashboard for internal analysts might display detailed feature importance vectors, partial dependence plots, and scenario simulation tools [33]. A call-centre interface might show a small number of key factors driving a recommendation, accompanied by concise textual explanations. These presentation choices have architectural implications, as explanation services must support multiple output formats and granularities. They also affect how explanation requests are triggered; some interfaces may compute explanations only upon explicit user action, while others may present them by default.

Implementation of interpretable Customer 360 models must also consider performance and resilience. Even inherently interpretable models can become complex at scale, for example when scorecards contain many bins or when rule-based systems accumulate numerous exceptions over time [34]. Rationalising such models requires periodic refactoring, supported by tools that detect redundancy, inconsistencies, and rarely used rules. Automated tests can validate that refactored models preserve key behaviours and constraints. In event-driven architectures, models may run in containers or serverless functions, and explanation services may be deployed as separate microservices. Network latency, service dependencies, and failure modes all influence the practical availability of explanations. Governance requirements may dictate that certain decisions must not be executed unless an explanation can be produced, which in turn necessitates architectural mechanisms for graceful degradation or queueing when explanation services are temporarily unavailable [35].

Security and privacy considerations intersect with interpretability at the architectural level. Customer 360 platforms already enforce access controls around raw data; explanation services must inherit and extend these controls. For example, an internal investigation team may have permission to view explanations that reference sensitive attributes, while line-of-business users see only aggregated or obfuscated versions. Implementing such policies may involve filtering explanation outputs, mapping low-level features to higher-level concepts, or redacting certain contributions. Logging and audit trails are important to record which explanations were generated, for which users, and in which contexts. These logs can support later investigations into alleged mis-use or unexpected behaviour, and they can provide evidence that governance policies around explanation access were followed [36].

In summary, system architecture for interpretable models over Customer 360 data extends beyond algorithmic choices and includes the design of feature stores, model registries, explanation services, user interfaces, and governance tooling. Implementation decisions should reflect the specific use cases and stakeholder requirements of the organisation, while maintaining a coherent approach to meta-data, access control, monitoring, and documentation. The architectural perspective underscores that interpretability is not solely a property of individual models but emerges from how models, data, and applications are integrated and managed.

## 6. Evaluation Framework and Illustrative Scenarios

Evaluating interpretable models over Customer 360 data requires metrics and procedures that capture both predictive performance and interpretability-relevant properties. Traditional performance metrics such as accuracy, area under the receiver operating characteristic curve, precision and recall, and calibration remain essential to determine whether models provide useful information [37]. However, they do not convey how understandable or trustworthy models are to stakeholders. Additional quantitative measures can characterise aspects of interpretability, such as the number of features used in a typical explanation, the concentration of attribution among a small set of features, the stability of explanations under small perturbations of inputs, and the consistency of explanations across time and segments. These measures can be computed during model development and monitored in production to detect drift.

Qualitative evaluation complements quantitative metrics by involving stakeholders in assessing the adequacy and clarity of explanations. For Customer 360 applications, structured workshops or user studies can present stakeholders with example predictions and corresponding explanations, asking them to judge whether the explanations are plausible, sufficient, and consistent with their domain understanding. Differences between stakeholder groups may surface; for instance, technical users might focus on numerical detail and statistical properties, while business users emphasise narrative coherence and practicality [38]. Such feedback can inform refinements to both models and explanation interfaces. Qualitative methods also help identify situations where technically accurate explanations nonetheless fail to support decision-making, for example when they rely on obscure feature names or refer to data fields that stakeholders do not recognise.

Illustrative scenarios provide a concrete way to reason about evaluation. Consider a churn prediction model built on Customer 360 data for a subscription service. The model uses features derived from usage logs, billing history, support interactions, and marketing campaign responses [39]. An interpretable approach might involve an additive model with monotonic constraints on certain variables, such as the relationship between consecutive missed payments and churn risk. Quantitatively, evaluation would assess discrimination and calibration on hold-out data, as well as the sparsity and stability of feature attributions. Qualitatively, explanations for high-risk customers would be examined by retention specialists, who would judge whether the highlighted drivers, such as reduced usage, unresolved complaints, or recent price increases, align with experience. Misalignments could indicate issues in feature engineering, model assumptions, or data quality.

Another scenario involves next-best-offer recommendations, where a model ranks potential offers for each customer based on predicted uplift in acceptance or value [40]. Interpretable models in this context must grapple with multiple possible actions and counterfactual reasoning. For example, a model

might estimate the incremental probability of acceptance relative to a baseline of no offer. Explanations should ideally articulate why a specific offer is preferred for a customer, referencing relevant preferences, past responses, and compatibility with product holdings. Evaluation might compare an interpretable ranking model to a black-box alternative in terms of both performance and stakeholder acceptance. If the interpretable model yields slightly lower predictive performance but is easier to understand and govern, decision-makers may still prefer it for regulatory or reputational reasons. Systematic experiments can measure business outcomes and investigate how explanation availability affects the behaviour of frontline staff and customers [41].

A third scenario concerns early warning models for complaints or regulatory breaches. Customer 360 data may include signals such as negative feedback, repeated contact attempts, or anomalies in product usage. Models predicting escalation risk can help organisations intervene proactively. Given the sensitivity of such use cases, interpretability and governance are central. Evaluation frameworks would examine whether explanations correctly highlight risk factors that justify closer attention and whether they avoid stigmatising particular demographic or behavioural segments [42]. Monitoring explanation patterns over time could reveal whether model behaviour changes as new communication channels are introduced or as policy changes alter customer interactions. Qualitative assessments with compliance officers would assess whether explanations provide sufficient detail to support documented rationales for interventions.

Across these scenarios, evaluation frameworks benefit from integrating interpretability metrics into standard model validation processes. This integration avoids treating interpretability as a separate, optional concern. For example, model selection criteria might consider a combination of performance metrics and interpretability measures, such as selecting models that achieve a target performance threshold while minimising explanation complexity or maximising stability. Governance committees reviewing models could receive evaluation summaries that include explanation-related metrics alongside traditional statistics, helping them make balanced decisions [43]. Over time, organisations can refine their interpretability evaluation practices based on accumulated experience, adjusting thresholds and procedures to match evolving expectations from regulators, customers, and internal stakeholders.

## 7. Human-in-the-Loop Decision-Making and Organisational Adoption

Human-in-the-loop decision-making is a central dimension of how interpretable models over Customer 360 data are used in practice. Even when decision flows are technically capable of full automation, many organisations choose to retain human oversight or shared control, especially for high-impact or sensitive use cases. In such settings, the role of the model is to provide quantitative assessments, rankings, or alerts that inform human judgement rather than replacing it. Interpretability then becomes not only a property of the model, but also a determinant of how effectively humans can integrate model outputs into their reasoning [44]. If explanations are too complex, unstable, or misaligned with domain concepts, human decision-makers may ignore them, rely on them mechanically, or develop inaccurate mental models of the system. Conversely, appropriately designed explanations can support calibrated trust, where users neither over-rely on nor systematically dismiss model recommendations.

The integration of interpretable models into existing organisational processes is mediated by norms, incentives, and skill profiles. Customer 360 platforms often serve multiple business units, each with distinct decision workflows and tolerance for automation. For instance, marketing teams may be comfortable with models that drive campaign targeting autonomously, as long as aggregate performance metrics remain within expected bounds, while credit risk or compliance teams may require manual review for specific thresholds or segments. Interpretable models can help these units negotiate appropriate levels of control by making trade-offs more visible [45]. A risk committee may agree that certain classes of decisions can be automated only if model logic is constrained to respect monotonic relationships with key variables and if clear, human-readable rationales are generated for audit. Such agreements are easier to reach when stakeholders can inspect and understand model structures and explanation outputs.

Training and capability building influence how effectively organisations adopt interpretable Customer 360 models. Users who interact with model outputs, such as frontline staff, analysts, and managers, benefit from guidance on how to read scores, confidence intervals, and explanations, and on how to combine them with contextual knowledge. For example, customer service agents may be told that a high churn risk score accompanied by explanations highlighting repeated unresolved issues and declining usage patterns suggests that retention efforts should focus on service remediation rather than discount offers [46]. This form of guidance connects model explanations to practical action strategies. Without such training, there is a risk that explanations are treated as opaque technical artefacts or that users infer unprincipled heuristics based on superficially salient features. Over time, organisations may develop standard interpretive frameworks, where certain patterns of feature contributions are associated with specific interventions or escalation paths.

Feedback mechanisms from human decision-makers to model developers and governance bodies are another key element of human-in-the-loop adoption. Users often encounter edge cases, novel behaviours, or contextual factors not fully captured in training data [47]. When models are interpretable, users can identify and articulate specific concerns, such as repeatedly seeing high propensity scores driven by features that appear outdated or irrelevant in recent campaigns. Structured channels, such as feedback forms linked to explanation interfaces or periodic review sessions, can collect such observations and feed them into model maintenance processes. Data scientists can then investigate whether observed patterns stem from data drift, feature leakage, or mis-specified constraints. Governance committees may use aggregated feedback to prioritise model recalibration, feature re-engineering, or policy adjustments. The interpretability of models and explanations thus reinforces a feedback loop that supports continuous improvement.

Cognitive factors shape how individuals process explanations and make decisions in the presence of model outputs [48]. There is evidence that humans favour simple, coherent narratives and may overweight a small number of prominent factors even when a decision is driven by many contributing elements. Interpretable models in Customer 360 contexts must therefore balance fidelity with cognitive accessibility. Explanations that present long lists of features with similar contribution magnitudes may be technically accurate but practically unhelpful, as users struggle to identify central drivers. Techniques that highlight a small subset of dominant contributors, map features to higher-level customer concepts, or provide summarised textual descriptions can reduce cognitive load. However, such simplifications must be managed carefully, as they may omit relevant nuances or introduce implicit prioritisation that does not fully reflect model behaviour [49]. Human factors research tailored to organisational contexts can inform the design of explanation formats that support reliable, repeatable interpretation.

The distributional effects of model-assisted decisions are also shaped by how humans use explanations. In Customer 360 applications, frontline staff may adapt their behaviour in response to model outputs, for example by directing more attention to high-risk or high-opportunity customers. If explanations emphasise certain behavioural or demographic features, staff may unconsciously associate those features with desired or undesired outcomes, potentially reinforcing biases in interactions. Organisations can mitigate such risks by framing explanations in terms of controllable, action-relevant factors, and by providing training that emphasises fairness and responsible use. For instance, explanations might focus on engagement patterns and service history rather than on attributes that could be perceived as sensitive, even if the underlying model uses a broader feature set subject to governance [50]. This does not remove the need for formal fairness analyses, but it aligns day-to-day decision practices with organisational values.

Human-in-the-loop settings often involve escalation paths and override mechanisms. Interpretable models can support principled overrides by making clear what the model suggests and why, thereby allowing decision-makers to articulate reasons for deviating from recommendations. For example, a relationship manager might override a low cross-sell recommendation score for a customer based on knowledge of a forthcoming life event, documenting that the override is motivated by information unavailable in current data feeds. Governance processes may require justification fields in systems where overrides occur, potentially pre-populated with explanation summaries that users can amend

[51]. Over time, analysis of override patterns can reveal systematic gaps in models or in Customer 360 data. Interpretable models make such analyses more informative, as overrides can be examined in light of the explanation structures that users saw at decision time.

Organisational adoption of interpretable Customer 360 models also depends on how responsibilities are allocated across teams. Data science, engineering, product, legal, compliance, and operations functions each contribute to model lifecycle management. Clear delineation of roles in relation to interpretability helps avoid gaps [52]. For instance, data scientists may be responsible for selecting model families and explanation methods, while feature engineers maintain semantic documentation and lineage links necessary for meaningful explanations. Product teams may specify user interface requirements for displaying predictions and explanations, while compliance teams define which features may appear in customer-facing rationales. Without explicit coordination, there is a risk that interpretability considerations are addressed piecemeal, with models technically capable of explanation but user interfaces or policies preventing effective use. Cross-functional working groups or model governance boards can serve as forums for aligning interpretability practices.

Change management considerations arise when organisations transition from legacy decision processes, based on rules or judgement alone, to Customer 360 models with interpretable outputs. Stakeholders accustomed to deterministic rule sets may initially perceive statistical models as less transparent, even when the latter are accompanied by detailed explanations [53]. Communication strategies can address this by demonstrating how interpretable models can reproduce and refine existing rules while providing quantified assessments of uncertainty and performance. Pilot deployments, where models are used as advisory tools before affecting actual decisions, can help build familiarity. During such pilots, discrepancies between model recommendations and current decisions can be analysed, with explanations providing insight into underlying causes. This evidence can inform adjustments to both models and policies before broader roll-out.

The maturity of data governance and analytics capabilities influences the feasibility of adopting interpretable Customer 360 models [54]. Organisations with established data catalogues, lineage tracking, and standardised documentation practices can more easily support explanation requirements, as feature semantics and provenance are already systematised. In less mature environments, implementing interpretable models may first require foundational investments in data quality, meta-data management, and access control. While it is technically possible to deploy interpretable algorithms without such infrastructure, the resulting explanations may lack necessary context, limiting their usefulness for governance and trust. Recognising interpretability as part of a broader data and analytics strategy rather than an isolated objective can help align investments and expectations.

The interaction between human decision-makers and interpretable models is dynamic over time. As models are updated, features are added or retired, and business environments change, explanations may shift in content and emphasis [55]. Users can experience such shifts as instability or inconsistency if they are not communicated and managed carefully. For example, a set of dominant features for a churn model might change following major product launches or pricing changes, leading to different patterns in explanation outputs. Regular communication about model updates, including summaries of how explanation distributions have changed and why, can support continuity of trust. Training materials and decision guidelines may need periodic revision to reflect the new explanatory landscape. Interpretable models facilitate such communication because changes in their structure and parameterisation can often be directly related to observable business or data changes [56].

An additional consideration is the potential for strategic responses by both internal and external actors to interpretable models and their explanations. Internally, teams might attempt to influence feature engineering or model configuration to favour certain business outcomes, for instance by incorporating features that are known to drive higher scores for particular segments. Externally, customers or partners might adjust their behaviour in response to perceived decision criteria, especially if explanations reveal aspects of decision logic. In some contexts, such strategic adaptation may be acceptable or even desired, as when customers are encouraged to improve behaviours that reduce risk or increase eligibility. In others, it may undermine the reliability of models [57]. Governance mechanisms should therefore assess

the extent to which explanation content could incentivise gaming and consider whether explanation granularity should vary by audience. Interpretable models allow these assessments to be grounded in explicit representations of decision logic rather than inferred from opaque systems.

Finally, the adoption of interpretable Customer 360 models is influenced by external expectations and norms in specific industries and jurisdictions. Regulators may issue guidance on explainability, fairness, and accountability that shapes organisational strategies. Industry bodies and professional associations may develop best practice frameworks for model risk management and responsible analytics. Interpretable models can align more readily with such frameworks because they expose their internal structure and can be mapped onto requirements for documentation, validation, and audit [58]. At the same time, compliance with external expectations does not guarantee that internal organisational needs are met. Stakeholders may require levels of detail or types of explanations that go beyond formal requirements, or they may have concerns that formal guidance does not address. Treating interpretability as a multi-layered concept that serves regulatory, organisational, and individual needs helps maintain a balanced approach to model adoption.

In aggregate, human-in-the-loop decision-making and organisational adoption processes underscore that interpretability in Customer 360 models is not solely a technical attribute but a component of a broader socio-technical system. Models, explanations, users, governance structures, and external constraints all interact [59]. Attention to training, feedback loops, cognitive factors, role allocation, change management, and strategic behaviour can influence whether interpretable models achieve their intended purpose of supporting informed, accountable decisions. The degree of success in these areas often determines whether Customer 360 initiatives become embedded as routine, trusted components of organisational practice or remain limited to isolated, experimental deployments.

## 8. Conclusion

Customer 360 platforms create opportunities to use integrated data for a variety of predictive and prescriptive applications, but they also introduce challenges related to transparency, accountability, and governance. Interpretable models offer one set of tools for addressing these challenges by making relationships between features and outcomes more accessible to human understanding. In the context of Customer 360 data, interpretability must contend with heterogeneous feature spaces, complex temporal dynamics, identity resolution issues, and shared infrastructure across multiple business domains. Consequently, interpretability cannot be addressed solely by choosing particular algorithms; it must be supported by data management practices, meta-data infrastructure, architectural designs, and governance processes [60].

This paper has discussed several families of interpretable models relevant to Customer 360 analytics, including linear and additive models, decision trees and scorecards, and post-hoc explanation techniques for more complex predictors. It considered how these approaches interact with the characteristics of Customer 360 data, such as event-derived features and temporal patterns, and how they can be aligned with stakeholder needs across business, technical, and oversight functions. The discussion emphasised that explanations are only useful when they map onto concepts and narratives that stakeholders recognise and can act upon, which in turn requires careful design of features, concepts, and interfaces.

Governance and trust emerged as central themes in the deployment of interpretable Customer 360 models. Governance encompasses not just regulatory compliance but also internal policies on data usage, fairness, documentation, and model monitoring [61]. Interpretable models can facilitate governance by making assumptions and decision logic more transparent, but they do not automatically ensure fair or appropriate outcomes. Evaluation frameworks that integrate performance, interpretability, and fairness metrics, supported by qualitative assessments with stakeholders, can help organisations navigate trade-offs and make informed choices about model deployment. Architectural patterns that incorporate feature stores with rich meta-data, model registries, explanation services, and carefully designed user interfaces provide technical foundations for these governance practices.

Future work in this area may explore more systematic methods for designing concept-based models that operate over Customer 360 data, better ways to quantify explanation stability and usefulness in operational settings, and stronger integration of interpretability considerations into automated model development pipelines. There is also scope for studying how different stakeholder groups interact with explanations and how these interactions influence decisions and outcomes. While no single approach will suit all organisations or use cases, the combination of interpretable modelling techniques, robust Customer 360 data management, and structured governance processes can provide a workable basis for using integrated customer data in ways that are both effective and accountable [62].

## References

- [1] M. Ahearne, T. Haumann, F. Kraus, and J. Wieseke, “It’s a matter of congruence: How interpersonal identification between sales managers and salespersons shapes sales success,” *Journal of the Academy of Marketing Science*, vol. 41, pp. 625–648, 3 2013.
- [2] K. M. Kacmar, M. C. Andrews, K. J. Harris, and B. J. Tepper, “Ethical leadership and subordinate outcomes: The mediating role of organizational politics and the moderating role of political skill,” *Journal of Business Ethics*, vol. 115, pp. 33–44, 6 2012.
- [3] E. Breugelmans, T. H. A. Bijmolt, J. Zhang, L. J. Basso, M. Dorotic, P. K. Kopalle, A. Minnema, W. J. Mijnlief, and N. V. Wunderlich, “Advancing research on loyalty programs: a future research agenda,” *Marketing Letters*, vol. 26, pp. 127–139, 6 2014.
- [4] S. H. Kukkuhalli, “Enabling customer 360 view and customer touchpoint tracking across digital and non-digital channels,” *Journal of Marketing & Supply Chain Management*, vol. 1, no. 3, 2022.
- [5] C. Tucker, “Digital data, platforms and the usual [antitrust] suspects: Network effects, switching costs, essential facility,” *Review of Industrial Organization*, vol. 54, pp. 683–694, 2 2019.
- [6] J. Bulow and J. Levin, “Matching and price competition,” *American Economic Review*, vol. 96, pp. 652–668, 5 2006.
- [7] S. Sozuer, G. S. Carpenter, P. K. Kopalle, L. McAlister, and D. R. Lehmann, “The past, present, and future of marketing strategy,” *Marketing Letters*, vol. 31, pp. 163–174, 7 2020.
- [8] M. Thoenig and T. Verdier, “A macroeconomic perspective on knowledge management,” *Journal of Economic Growth*, vol. 15, pp. 33–63, 2 2010.
- [9] T. R. Harmon-Kizer, “Let the borrower beware: Towards a framework for debiasing rollover behavior in the payday loan industry,” *Journal of Consumer Policy*, vol. 42, pp. 245–266, 2 2019.
- [10] S. Balasubramanian, R. A. Peterson, and S. L. Jarvenpaa, “Exploring the implications of m-commerce for markets and marketing,” *Journal of the Academy of Marketing Science*, vol. 30, pp. 348–361, 10 2002.
- [11] M.-H. Huang and R. T. Rust, “A strategic framework for artificial intelligence in marketing,” *Journal of the Academy of Marketing Science*, vol. 49, pp. 30–50, 11 2020.
- [12] B. So, “An implementable optimization of airfare structure,” *Journal of Revenue and Pricing Management*, vol. 16, pp. 441–465, 3 2017.
- [13] M. D. Wittman and P. Belobaba, “Dynamic availability of fare products with knowledge of customer characteristics,” *Journal of Revenue and Pricing Management*, vol. 16, pp. 201–217, 4 2016.
- [14] K. Glenk, R. J. Johnston, J. Meyerhoff, and J. Sagebiel, “Spatial dimensions of stated preference valuation in environmental and resource economics: Methods, trends and challenges,” *Environmental and Resource Economics*, vol. 75, pp. 215–242, 1 2019.
- [15] D. Grewal, P. K. Kopalle, and J. Hulland, “Addressing the greatest global challenges (un sdgs) with a marketing lens,” *Journal of the Academy of Marketing Science*, vol. 52, pp. 1263–1272, 9 2024.
- [16] J. Z. Zhang and C.-W. Chang, “Consumer dynamics: theories, methods, and emerging directions,” *Journal of the Academy of Marketing Science*, vol. 49, pp. 166–196, 3 2020.

- [17] R. Perren and R. V. Kozinets, "Lateral exchange markets: How social platforms operate in a networked economy," *Journal of Marketing*, vol. 82, pp. 20–36, 1 2018.
- [18] S. Wright and G.-X. Xie, "Perceived privacy violation: Exploring the malleability of privacy expectations," *Journal of Business Ethics*, vol. 156, pp. 123–140, 5 2017.
- [19] M. D. Johnson and F. Selnes, "Customer portfolio management: Toward a dynamic theory of exchange relationships," *Journal of Marketing*, vol. 68, pp. 1–17, 4 2004.
- [20] S. Janani, R. M. Christopher, A. N. Nikolov, and M. A. Wiles, "Marketing experience of ceos and corporate social performance," *Journal of the Academy of Marketing Science*, vol. 50, pp. 460–481, 1 2022.
- [21] W. Zhou and S. Piramuthu, "Information relevance model of customized privacy for iot," *Journal of Business Ethics*, vol. 131, pp. 19–30, 7 2014.
- [22] A. Parasuraman and D. Grewal, "The impact of technology on the quality-value-loyalty chain: A research agenda," *Journal of the Academy of Marketing Science*, vol. 28, pp. 168–174, 1 2000.
- [23] S. Brown, P. McDonagh, and C. J. Shultz, "Titanic: Consuming the myths and meanings of an ambiguous brand," *Journal of Consumer Research*, vol. 40, pp. 595–614, 12 2013.
- [24] R. Agarwal, M. Dugas, G. Gao, and P. Kannan, "Emerging technologies and analytics for a new era of value-centered marketing in healthcare," *Journal of the Academy of Marketing Science*, vol. 48, pp. 9–23, 10 2019.
- [25] Y. Luo, "Toward a cooperative view of mnc-host government relations: Building blocks and performance implications," *Journal of International Business Studies*, vol. 32, pp. 401–419, 9 2001.
- [26] L. D. Solomon, "Perspectives on human nature and their implications for business organizations," *Humanomics*, vol. 12, pp. 8–52, 1 1996.
- [27] D. Bevelander, M. J. Page, L. Pitt, and M. Parent, "On a mission : achieving distinction as a business school?," *South African Journal of Business Management*, vol. 46, pp. 29–41, 6 2015.
- [28] J. A. Arevalo and D. Aravind, "Strategic outcomes in voluntary csr: Reporting economic and reputational benefits in principles-based initiatives," *Journal of Business Ethics*, vol. 144, pp. 201–217, 9 2015.
- [29] V. D. Bush, G. M. Rose, F. W. Gilbert, and T. N. Ingram, "Managing culturally diverse buyer-seller relationships: The role of intercultural disposition and adaptive selling in developing intercultural communication competence," *Journal of the Academy of Marketing Science*, vol. 29, pp. 391–404, 10 2001.
- [30] X. Wang, X. Lin, and B. Shao, "Security and privacy protection in developing ethical ai: A mixed-methods study from a marketing employee perspective," *Journal of Business Ethics*, vol. 200, pp. 373–392, 12 2024.
- [31] S. H. Kukkuhalli, "Increasing digital sales revenue through 1:1 hyper-personalization with the use of machine learning for b2c enterprises," *Journal of Artificial Intelligence, Machine Learning and Data Science*, vol. 1, no. 1, 2023.
- [32] B. Williams-Jones and V. Ozdemir, "Challenges for corporate ethics in marketing genetic tests," *Journal of Business Ethics*, vol. 77, pp. 33–44, 2 2007.
- [33] C. Miao, M. J. Barone, S. Qian, and R. H. Humphrey, "Emotional intelligence and service quality:a meta-analysis with initial evidence on cross-cultural factors and future research directions," *Marketing Letters*, vol. 30, pp. 335–347, 8 2019.
- [34] J. Wind and A. Rangaswamy, "Customerization: The next revolution in mass customization," *Journal of Interactive Marketing*, vol. 15, pp. 13–32, 11 2001.
- [35] L. A. Fennell, "Optional price discrimination," *SSRN Electronic Journal*, 1 2022.
- [36] S. Zhang, X. Lin, X. Li, and A. Ren, "Service robots' anthropomorphism: dimensions, factors and internal relationships," *Electronic Markets*, vol. 32, pp. 277–295, 2 2022.
- [37] F. Caro, V. M. de Albeniz, and B. Apaolaza, "The value of online interactions for store execution," *SSRN Electronic Journal*, 1 2021.
- [38] D. D. Gremler, Y. V. Vaerenbergh, E. Brüggén, and K. P. Gwinner, "Understanding and managing customer relational benefits in services: a meta-analysis," *Journal of the Academy of Marketing Science*, vol. 48, pp. 565–583, 10 2019.

- [39] K. K. Wang, M. D. Wittman, and T. Fiig, "Dynamic offer creation for airline ancillaries using a markov chain choice model," *Journal of Revenue and Pricing Management*, vol. 22, pp. 103–121, 2 2023.
- [40] A. Kalra, R. Singh, V. Badrinarayanan, and A. Gupta, "How ethical leadership and ethical self-leadership enhance the effects of idiosyncratic deals on salesperson work engagement and performance," *Journal of Business Ethics*, vol. 196, pp. 169–188, 5 2024.
- [41] E. S. Oblander, S. Gupta, C. F. Mela, R. S. Winer, and D. R. Lehmann, "The past, present, and future of customer management," *Marketing Letters*, vol. 31, pp. 125–136, 6 2020.
- [42] A. Carmeli, L. E. Atwater, and A. Levi, "How leadership enhances employees' knowledge sharing: the intervening roles of relational and organizational identification," *The Journal of Technology Transfer*, vol. 36, pp. 257–274, 2 2010.
- [43] R. Panwar, E. Nybakk, E. Hansen, and J. Pinkse, "Does the business case matter? the effect of a perceived business case on small firms' social engagement," *Journal of Business Ethics*, vol. 144, pp. 597–608, 8 2015.
- [44] K. Bischoff, C. Volkmann, and D. B. Audretsch, "Stakeholder collaboration in entrepreneurship education: an analysis of the entrepreneurial ecosystems of european higher educational institutions," *The Journal of Technology Transfer*, vol. 43, pp. 20–46, 4 2017.
- [45] M. G. Walsh, L. L. Lan, M. O. Lwin, and J. D. Williams, "Marketers' boon in cyberspace: The anticybersquatting consumer protection act," *Journal of Public Policy & Marketing*, vol. 22, pp. 96–101, 4 2003.
- [46] I. Reimers and J. Waldfogel, "Digitization and pre-purchase information: The causal and welfare impacts of reviews and crowd ratings," *American Economic Review*, vol. 111, pp. 1944–1971, 6 2021.
- [47] C. Song, S. Jang, J. Wiggins, and E. L. Nowlin, "Does haste always make waste? service quantity, service quality, and incentives in speed-intensive service firms," *Service Business*, vol. 13, pp. 289–304, 8 2018.
- [48] G. A. Shinkle and A. Kriauciunas, "Institutions, size and age in transition economies: Implications for export growth," *Journal of International Business Studies*, vol. 41, pp. 267–286, 4 2009.
- [49] E. I. Altman and G. Sabato, "Effects of the new basel capital accord on bank capital requirements for smes," *Journal of Financial Services Research*, vol. 28, pp. 15–42, 11 2005.
- [50] A. Z. Röschmann, "Digital insurance brokers - old wine in new bottles? : how digital brokers create value," *Zeitschrift für die gesamte Versicherungswissenschaft*, vol. 107, pp. 273–291, 8 2018.
- [51] A. Ban, "Osborne, huw (ed.): The rise of the modernist bookshop: Books and the commerce of culture in the twentieth century," *Publishing Research Quarterly*, vol. 33, pp. 112–113, 12 2016.
- [52] P. J. Trocchia, R. Saine, and M. Luckett, "I've wanted a bmw since i was a kid: An exploratory analysis of the aspirational brand.," *Journal of Applied Business Research (JABR)*, vol. 31, pp. 331–344, 12 2014.
- [53] G. E. Metcalf and G. Norman, "Oligopoly deregulation and the taxation of commodities," *Contributions in Economic Analysis & Policy*, vol. 2, pp. 1–16, 10 2003.
- [54] U. Gretzel, C. Koo, M. Sigala, and Z. Xiang, "Special issue on smart tourism: convergence of information technologies, experiences, and theories," *Electronic Markets*, vol. 25, pp. 175–177, 7 2015.
- [55] M. Binder and A.-K. Blankenberg, "Identity and well-being in the skilled crafts and trades," *SSRN Electronic Journal*, 1 2021.
- [56] M. J. Bitner, S. W. Brown, and M. L. Meuter, "Technology infusion in service encounters," *Journal of the Academy of Marketing Science*, vol. 28, pp. 138–149, 1 2000.
- [57] B. M. Owen, "Antitrust and vertical integration in "new economy" industries with application to broadband access," *Review of Industrial Organization*, vol. 38, pp. 363–386, 5 2011.
- [58] Y. J. Park and M. M. Skoric, "Personalized ad in your google glass? wearable technology, hands-off data collection, and new policy imperative," *Journal of Business Ethics*, vol. 142, pp. 71–82, 7 2015.
- [59] S. H. Kukkuhalli, "Improving digital sales through reducing friction points in the customer digital journey using data engineering and machine learning," *International Journal of Innovative Research in Engineering & Multidisciplinary Physical Sciences*, vol. 10, no. 3, 2022.

- [60] R. T. Watson, L. Pitt, P. Berthon, and G. M. Zinkhan, “U-commerce: Expanding the universe of marketing,” *Journal of the Academy of Marketing Science*, vol. 30, pp. 333–347, 10 2002.
- [61] M. L. Meuter, A. L. Ostrom, R. I. Roundtree, and M. J. Bitner, “Self-service technologies: Understanding customer satisfaction with technology-based service encounters,” *Journal of Marketing*, vol. 64, pp. 50–64, 7 2000.
- [62] E. Hartwich, A. Rieger, J. Sedlmeir, D. Jurek, and G. Fridgen, “Machine economies,” *Electronic Markets*, vol. 33, 7 2023.