

Article

SIoT-Enabled Opportunistic Sensing Under Partial Observability: Evaluating Recruitment Policies for Human Digital Twin Context Estimation

Lorenzo Bacchiani ¹, Andrea Melis ¹, Roberta Presta ², Matteo Anedda ³, Daniele Giusto ³
and Roberto Girau ^{1,*}

¹ Department of Computer Science and Engineering, University of Bologna, 40136 Bologna, Italy; lorenzo.bacchiani2@unibo.it (L.B.); a.melis@unibo.it (A.M.)

² Centro Scienza Nuova, Università degli Studi Suor Orsola Benincasa, 80135 Napoli, Italy; roberta.presta@unisob.na.it

³ Department of Electrical and Electronic Engineering, University of Cagliari, 09123 Cagliari, Italy; matteo.anedda@unica.it (M.A.); ddgiusto@unica.it (D.G.)

* Correspondence: roberto.girau@unibo.it

Abstract

Human Digital Twins (HDTs) rely on reliable context estimation to support personalized services, adaptive decision-making, and context-aware interaction, yet local ego sensing can be noisy, intermittent, or ambiguous. This study investigates whether Social Internet of Things (SIoT)-enabled opportunistic recruitment of external sources can improve HDT context estimation under partial observability. A controlled synthetic simulator is developed in which the SIoT layer is represented as a typed object graph supporting graph-constrained candidate discovery, bounded relationship-guided discovery, and cost-aware recruitment. The evaluation compares ego-only sensing, a high-coverage opportunistic-all reference, SIoT-aware bounded recruitment, and privacy-aware SIoT recruitment across nominal conditions, degraded ego sensing, ambiguous local context, and noisy/untrusted external sources. Performance is assessed with strict joint overall context accuracy, mean variable accuracy, per-variable Macro-F1, operational cost, effective recovery cost, source-cap-matched baselines, discovery-mode comparisons, and graph-size scalability. The results show that opportunistic-all recruitment gives the highest raw OCA because it recruits many sources, whereas bounded SIoT-aware policies provide lower-cost operating points and preserve nearly the same accuracy as exhaustive SIoT discovery in the discovery-mode comparison. The findings are therefore framed as accuracy–cost–privacy–scalability trade-offs in a synthetic, reproducible testbed rather than as deployment-level claims of universal policy superiority.



Academic Editor: Ting Wang

Received: 11 May 2026

Revised: 2 July 2026

Accepted: 4 July 2026

Published: 9 July 2026

Copyright: © 2026 by the authors.

Licensee MDPI, Basel, Switzerland.

This article is an open access article distributed under the terms and conditions of the [Creative Commons Attribution \(CC BY\)](https://creativecommons.org/licenses/by/4.0/) license.

Keywords: Human Digital Twin; Social Internet of Things (SIoT); opportunistic sensing; context estimation; partial observability

1. Introduction

Human Digital Twins are increasingly framed as human-centered digital representations that support personalized services, adaptive decision-making, and context-aware interaction across IoT-enabled environments [1–4]. In these architectures, the context layer is not a peripheral component; it is the operational interface between raw observations and downstream adaptation logic. It directly supports functions such as adaptive assistance,

personalized monitoring, and context-aware service orchestration. More broadly, accurate context estimation underpins user-aware, adaptive, and accessibility-sensitive interaction in human-facing systems, including mobility-related and driver-state-aware interface settings [5,6]. Maintaining this layer requires continuous integration of heterogeneous, time-varying signals from devices with different sensing roles, reliability levels, and update dynamics. For this reason, context-layer estimation should be regarded as a foundational prerequisite for broader HDT maintenance rather than as an isolated modeling detail.

A central difficulty is partial observability: observations produced by the local reference device associated with the HDT, referred to in this work as the *ego device*, are often noisy, intermittently missing, and locally ambiguous, preventing the context state from being inferred reliably from ego sensing alone at every timestep. These conditions should be treated as structural rather than exceptional in mobile and pervasive sensing workflows [7,8]. The challenge is amplified when context is represented jointly across multiple variables: strict joint correctness requires all components of the state to be estimated correctly simultaneously, which is inherently harder than achieving acceptable performance on each variable independently. This distinction has direct methodological consequences for how context-estimation quality should be evaluated and interpreted.

In this context, Social Internet of Things (SIoT)-enabled opportunistic sensing provides a natural architectural response when local observability is insufficient. By leveraging a relationship-aware graph structure, the ego device can discover and recruit external sources beyond its local sensing boundary [9,10]. Because several SIoT relationship types are anchored to ownership, shared location, shared activity, or user-associated object interactions, the resulting candidate network is naturally aligned with the person-centered context represented by the HDT. However, external recruitment does not provide a cost-free informational gain. It introduces coupled design dimensions, including discovery scope, trustworthiness, source relevance, freshness, overhead, latency, and privacy cost, that must be managed explicitly in the design of trust-aware SIoT interactions and socially aware recruitment [11–15]. The resulting question is therefore not simply whether external sources can help, but which recruitment strategy is preferable under different observability regimes and operational constraints.

Nevertheless, existing literature provides essential building blocks but remains fragmented with respect to this specific decision problem. HDT research establishes the broader motivation and architectural scope, SIoT work develops trust, discovery, and virtual-counterpart mechanisms, and opportunistic sensing research characterizes the trade-offs associated with external-source recruitment [2,12,16,17]. However, systematic, scenario-conditioned comparisons of SIoT graph-constrained recruitment strategies for HDT-oriented context-layer estimation under partial observability remain limited. This paper addresses that gap by isolating context estimation as a bounded HDT subproblem and evaluating recruitment behavior in a controlled, fully synthetic simulator designed for comparative analysis. It therefore contributes a reproducible evaluation framework for SIoT-enabled HDT context sensing, not a general HDT architecture or a deployment-ready fusion stack.

Accordingly, this study makes five focused contributions. First, it defines an explicit SIoT graph model for HDT context estimation, including typed object relationships and relation utilities used during discovery and recruitment. Second, it introduces bounded relationship-guided discovery as the default SIoT-aware operating mode and compares it with exhaustive hop-limited SIoT discovery. Third, it evaluates ego-only sensing, opportunistic-all recruitment, SIoT-aware bounded recruitment, and privacy-aware SIoT recruitment under four controlled observability regimes. Fourth, it adds a cost-aware evaluation that separates discovery, recruitment, selection, total operational cost, effec-

tive recovery cost, and privacy exposure. Fifth, it includes source-cap-matched non-SIoT baselines and scalability experiments to separate source-selection quality from raw source coverage and graph-size effects.

The remainder of this paper is organized as follows. Section 2 reviews related work, Section 3 defines the system model and problem formulation, Section 4 presents the experimental objectives, simulator design, recruitment policies, metrics, and results, Section 5 discusses the implications, Section 6 states the limitations, and Section 7 concludes the paper.

2. Related Work

2.1. Social Internet of Things: Graph Structure, Discovery, and Trust

The Social Internet of Things (SIoT) reframes IoT interaction as a relationship-aware graph paradigm rather than as generic pairwise connectivity [9,18]. In this view, nodes do not only exchange data; they also maintain social/functional relations that structure discovery paths, interaction opportunities, and service reachability. Importantly, this graph semantics is methodologically important for context estimation because source availability is constrained by relation topology, traversal rules, and neighborhood structure, not just by radio visibility. Subsequent SIoT overviews further emphasize that graph organization and navigability are central to scalable interaction in heterogeneous deployments [10].

For recruitment-oriented sensing tasks, discovery and reachability therefore become first-order design variables. Friendship selection, relationship management, and service-discovery mechanisms directly shape the candidate source pool from which any recruitment policy can select [19–21]. This means that policy behavior depends not only on local ranking heuristics, but also on how the underlying SIoT graph exposes or withholds potential contributors across hops and relation types. As a result, evaluating recruitment policies under controlled observability regimes is most meaningful when their candidate-access assumptions are made explicit.

Trust management is the second pillar of this SIoT foundation. Trustworthiness models formalize how confidence in external entities can be computed, updated, and used to filter or weight contributed observations before fusion [11,22]. More recent SIoT trust surveys systematize this landscape, showing that trust is multidimensional and context-dependent, and that practical designs differ in evidence aggregation, trust-update dynamics, robustness against unreliable or adversarial behavior, and broader scheme-level trade-offs [12–14,23]. For context-layer estimation, this literature suggests that trust is not an ancillary security feature but a core control dimension in recruitment policy design.

2.2. Distributed and Edge-Assisted SIoT Support

A complementary line of work addresses how SIoT interaction logic can be supported operationally at scale. Collaborative edge-computing architectures exploit proximity and social relationships for distributed offloading, caching, and data processing, while broader SIoT surveys identify architectures, platforms, service discovery, network navigability, and relationship management as closely coupled dimensions of scalable SIoT operation [20,24,25]. This is relevant to recruitment because repeated discovery, trust updates, and source selection require timely access to interaction metadata and relationship context.

Distributed and edge-assisted support is therefore not only an architectural convenience; it also influences the practical cost model of recruitment. By localizing coordination and data processing, these architectures can reduce dependence on remote centralized orchestration and improve responsiveness when recruitment decisions must be repeated over time [24]. In turn, this makes it more realistic to evaluate not only accuracy outcomes but also operational dimensions such as latency and source-usage burden within relationship-aware sensing workflows.

2.3. Human-Side Digital Counterparts, Virtual User, and HDT Context Modeling

The human-side counterpart perspective extends SIoT concepts from device relations to user-centered digital representations. The Virtual User framework treats user profiles, preferences, and context as explicit computational entities in IoT ecosystems, enabling interaction logic aligned with human state rather than device state alone [16]. This perspective is directly pertinent to HDT-oriented context estimation, where the utility of sensing depends on how effectively environmental and device signals are mapped onto human-relevant context variables, including multimodal integration pipelines used for human-state modeling [26].

In parallel, HDT literature consistently identifies context representation and context inference as enabling layers for personalized and adaptive services [1–4,27,28]. However, these studies are intentionally broad, often spanning architectures, enabling technologies, and application domains. Their contribution is primarily conceptual and architectural framing, not scenario-conditioned evaluation of recruitment policies for one bounded estimation subproblem. In contrast, the present work addresses context-layer estimation as a prerequisite for HDT maintenance rather than HDT maintenance in its entirety.

This human-side perspective is further supported by application-facing studies in which interaction quality depends on how reliably user context is represented and updated. Recent work explores expert co-design workflows for haptic-enhanced automotive virtual-reality interfaces [29], reports accessibility barriers in mobility interfaces [5], and evaluates intelligent driver-state-monitoring HMI strategies [6]. More recent digital-twin research consolidates this human-centered mobility perspective. Driver Digital Twin studies emphasize multimodal state fusion, personalized modeling, and time-varying driver representations [30]. Broader mobility Digital Twin frameworks integrate human, vehicle, and traffic counterparts through cloud–edge–device architectures [31], while work on connected and automated vehicles highlights tightly coupled, adaptive, and interpretable physical–digital models [32]. A trustworthy companion-AI architecture for human-aware transition of control further shows how reliable human-state and context estimation underpins safe, adaptive interaction [33]. Although these studies do not target SIoT recruitment policies directly, they naturally extend the rationale of this subsection by showing why accurate, context-aware human-state estimation is operationally consequential in human-centered systems.

2.4. Opportunistic Sensing, Crowdsensing, and Source Recruitment Under Partial Observability

Opportunistic sensing and mobile crowdsensing literature provides the closest methodological background for recruitment-policy analysis. Classic and recent studies show that external-source participation can compensate for sparse or unreliable local sensing, but only through explicit trade-offs among information gain, communication overhead, latency, coordination complexity, and privacy exposure [8,17,34,35]. This trade-off structure aligns directly with SIoT-enabled recruitment decisions, where broader source access can improve coverage but may increase operational and privacy costs.

Context-aware mobile sensing research further underscores that partial observability is structural rather than exceptional: missingness, ambiguity, and heterogeneous source quality are recurrent conditions in real sensing pipelines [7]. Under such conditions, recruitment strategies should be judged by how they behave across different observability stressors, not only by average-case accuracy. Sparse crowdsensing work reinforces this point by showing that performance depends strongly on source selection assumptions, coverage patterns, and data sparsity regimes [8,17].

Trust-aware and socially aware recruitment approaches in mobile crowdsensing move closer to the present problem by coupling participant selection with reputation, social

relationships, task suitability, privacy, and resource constraints [15,36,37]. Complementary privacy-aware task-allocation formulations in social-sensing-based edge systems further demonstrate that participant selection must jointly balance utility and privacy risk under resource constraints [38]. This is also consistent with privacy-preserving participant-selection mechanisms proposed in mobile crowdsensing, where selection quality and privacy protection must be handled jointly [39]. Together, these works demonstrate the practical value of socially structured information for selective source use, but they do not generally provide a dedicated, scenario-conditioned comparison aimed at HDT context-layer estimation with explicit joint-versus-variable-wise reporting criteria. This leaves open a focused evaluation space at the intersection of SIoT and HDT-oriented sensing.

Related SIoT work on assigning sensing tasks to devices through a social network of objects further strengthens this bridge by explicitly framing socially informed task allocation under resource and interaction constraints [40]. For the present study, this line is methodologically relevant because it reinforces the idea that recruitment should be treated as constrained policy design over socially structured candidate sets, even when the final evaluation target is HDT context-layer estimation under partial observability.

Modern learned context and IoT-management methods address related but distinct parts of the design space. Reinforcement-learning (RL) and distributed-learning approaches typically optimize adaptive decisions such as resource allocation, task scheduling, off-loading, or sensing policies under dynamic and uncertain operating conditions [41–44]. Graph-neural-network (GNN) approaches instead learn representations over sensor, device, or service graphs, making them well suited to topology-aware IoT reasoning, graph-based service discovery, and relation-aware device representations [45,46]. Probabilistic and Bayesian approaches, including dynamic Bayesian networks, explicitly model uncertainty, dependence, and temporal evolution in multimodal or multi-sensor context fusion [47,48]. The present paper is complementary to these streams: it does not train a learned estimator, an RL policy, or a probabilistic fusion model. Instead, it keeps fusion fixed and interpretable, then evaluates how SIoT graph-constrained discovery and bounded recruitment shape accuracy, cost, privacy, and scalability trade-offs.

2.5. Positioning of This Work

Taken together, the literature forms a coherent but still partially disconnected foundation. HDT studies provide broad human-centered motivation and context-modeling rationale; SIoT studies provide graph-constrained discovery, trust mechanisms, and virtual-counterpart support; and opportunistic sensing/crowdsensing studies provide recruitment trade-offs under sparse and heterogeneous observations [2,8,10,12]. What is less commonly developed is their integration into a controlled, scenario-conditioned evaluation of recruitment policies specifically for HDT context-layer estimation under partial observability.

Beyond providing access to additional sensing sources, SIoT has a structural affinity with HDT context estimation. Its object graph is organized through relations such as ownership, co-location, co-work, social-object, and parental-object relationships, which connect devices to users, their activities, and their surrounding environments [11,16,18]. This user-anchored organization may increase the likelihood that socially or functionally related objects provide contextually relevant evidence compared with an unstructured sensor pool. In this study, this potential advantage is operationalized through graph-constrained discovery and relationship-guided recruitment; however, it is evaluated only within the controlled synthetic testbed and is not presented as an empirically demonstrated deployment advantage.

This work is positioned at that integration point. It does not propose a general HDT architecture, a learned context-fusion model, or a deployment-calibrated sensing system.

Instead, it isolates one foundational subproblem and compares ego-only sensing, a high-coverage opportunistic-all reference, SIoT-aware bounded recruitment, and privacy-aware SIoT recruitment under explicitly defined observability regimes. The contribution is therefore a focused systems-level characterization of recruitment behavior, graph-constrained discovery, source-cap-matched selection controls, and scalability-aware cost trade-offs, grounded in existing SIoT and crowdsensing principles.

3. System Model and Problem Formulation

3.1. HDT Context State Under Partial Observability

We model the HDT context at timestep t as a discrete four-variable tuple.

$$\mathbf{c}_t = (p_t, a_t, e_t, r_t), \quad (1)$$

where p_t , a_t , e_t , and r_t denote *Place*, *Activity*, *Environmental Load*, and *Resource State*, respectively. The estimator outputs an inferred tuple $\hat{\mathbf{c}}_t$ from the evidence available at timestep t . The categorical domains are synthetic simulator domains used to evaluate recruitment behavior, not a claim to exhaustively model human context. In the simulator, *Place* and *Activity* each have five categorical states, whereas *Environmental Load* and *Resource State* each have three categorical states; therefore, $|\mathcal{D}_p| = 5$, $|\mathcal{D}_a| = 5$, $|\mathcal{D}_e| = 3$, and $|\mathcal{D}_r| = 3$.

The ego device provides probabilistic, intermittently missing local observations. We denote ego confidence by $q_t \in [0, 1]$, where lower values indicate weaker local evidence under the active scenario conditions. To support adaptive recruitment diagnostics, the simulator also tracks an unresolved-variable count $u_t \in \{0, \dots, 4\}$, i.e., the number of context variables whose estimated value remains insufficiently supported by the currently available fused evidence at timestep t . Conceptually, this establishes a latent-versus-observed distinction: the underlying context tuple exists regardless of whether local sensing at that timestep is complete or reliable.

3.2. SIoT Graph and Source Attribute Model

The environment is represented as an SIoT graph $G = (V, E)$, where $V = v_0, \dots, v_n$ contains the ego node v_0 and external object/source nodes, and E contains typed SIoT relations. The simulator uses five relation types: ownership object relationship (OOR), co-location object relationship (CLOR), co-work object relationship (CWOR), social object relationship (SOR), and parental object relationship (POR). Their default relation utilities are set to 1.0, 0.8, 0.8, 0.6, and 0.5, respectively, following prior SIoT trustworthiness modeling of OOR, CLOR, CWOR, SOR, and POR relation factors [11]; here, these values are used as synthetic relation-utility priors within the controlled testbed. These relation types do not imply real user relationships in the synthetic study; they define graph semantics for discovery and scoring.

Each external source i is characterized by attributes used in policy scoring and filtering: variable relevance R_i , trust score T_i , freshness F_i , latency proxy L_i , privacy cost P_i , graph-distance/access factor A_i , relation utility U_i , and variable coverage $\Gamma_i \subseteq \{p, a, e, r\}$. These attributes provide the quality/cost descriptors that operationalize recruitment decisions. They are synthetic but fixed by scenario and simulator rules for a given run, enabling controlled policy comparison.

Candidate-source discovery is graph constrained. Under hop-limited access, the candidate set at timestep t is

$$D_t(h_t) = \{i \in V \setminus \{v_0\} : \text{dist}_G(v_0, i) \leq h_t\}, \quad (2)$$

where h_t is the active hop radius and dist_G is the shortest-path distance in the SIoT graph. Exhaustive hop-limited discovery evaluates all nodes in $D_t(h_t)$. Bounded relationship-guided discovery instead expands a relation-prioritized frontier using the OOR, CLOR, CWOR, SOR, and POR priorities, stops when candidate, quality, and variable-coverage conditions are satisfied, and applies an adaptive budget multiplier when degradation triggers are active.

3.3. Fusion and Recruitment Problem Statement

All policies share the same interpretable weighted-voting fusion mechanism; policy differences arise only from discovery and source recruitment. At each timestep, a policy selects a subset of discovered sources $S_t \subseteq D_t$, and fusion is then performed over variable-level evidence contributed by ego sensing and the recruited sources, using trust/relevance/freshness-informed weights. Selective policies rank discovered candidates using the multiplicative score

$$s_i = \frac{R_i T_i F_i U_i A_i}{(1 + \lambda_L L_i)(1 + \lambda_P P_i)}, \quad (3)$$

where λ_L and λ_P weight latency and privacy penalties. In the privacy-aware SIoT policy, the privacy penalty is strengthened within the same scoring structure. The multiplicative form is used to model conjunctive suitability: a source should be selected only when relevance, trust, freshness, relation utility, and accessibility are jointly adequate. In contrast to an additive weighted sum, this formulation prevents a very high value in one dimension from fully compensating for a near-zero value in another critical dimension. For example, a source with negligible trust or negligible relevance should contribute little to recruitment even if it is fresh or easily reachable.

The methodological problem studied in this paper is therefore: under partial observability, which recruitment policy yields the best trade-off between context-estimation quality and operational burden, including source count, discovery cost, recruitment cost, selection cost, latency proxy, and privacy exposure? This formulation is explicitly designed to isolate recruitment effects by holding fusion logic fixed across policies.

3.4. Operational Cost Model

The simulator separates discovery, recruitment, and selection costs to avoid treating all external evidence as equally cheap. Let E_t^{trav} be the set of traversed SIoT edges during discovery, D_t^{score} the candidates scored by a selective policy, and S_t the recruited sources. Discovery cost is

$$C_t^{\text{disc}} = \sum_{e \in E_t^{\text{trav}}} w_{\rho(e)} + \lambda_E |E_t^{\text{trav}}|, \quad (4)$$

where $\rho(e)$ is the relation type of edge e , $w_{\rho(e)}$ is the corresponding relation traversal cost, and λ_E penalizes traversed-edge volume. Recruitment cost is

$$C_t^{\text{rec}} = \sum_{i \in S_t} (1 + \lambda_L L_i + \lambda_P P_i), \quad (5)$$

where L_i and P_i denote latency and privacy proxies. Selection cost is

$$C_t^{\text{sel}} = \lambda_S |D_t^{\text{score}}|, \quad (6)$$

where λ_S is the candidate-scoring cost coefficient. Total operational cost is

$$C_t^{\text{tot}} = C_t^{\text{disc}} + C_t^{\text{rec}} + C_t^{\text{sel}}. \quad (7)$$

For context-transition events $\tau \in \mathcal{T}$, effective recovery cost summarizes the mean operational cost accumulated until the strict joint tuple is recovered:

$$C^{\text{eff}} = \frac{1}{|\mathcal{T}|} \sum_{\tau \in \mathcal{T}} \sum_{\Delta=0}^{T_{\text{conv}}(\tau)} C_{\tau+\Delta}^{\text{tot}}. \quad (8)$$

Here $T_{\text{conv}}(\tau)$ is defined in Equation (11). This is a synthetic operational-cost proxy, not a measured deployment energy or network-latency quantity.

3.5. Conceptual Estimation Workflow

At each timestep, the model can be read as a compact estimation chain. First, an underlying latent HDT context tuple $\mathbf{c}_t$ exists. Second, the ego device provides local evidence, which may be incomplete or uncertain under the active-observability regime. Third, the SIoT graph determines which external sources are reachable as candidates. Fourth, the recruitment policy selects which discovered candidates actually contribute evidence (the recruited set S_t). Fifth, the same weighted-voting fusion mechanism combines ego and recruited-source evidence. The output is the estimated context tuple $\hat{\mathbf{c}}_t$.

This decomposition clarifies the role of the two experimental factors. Observability scenarios affect evidence quality and availability (e.g., missingness, ambiguity, or external-source reliability), whereas recruitment policies affect which external evidence enters fusion. Because the fusion rule is shared across policies, comparative performance differences can be interpreted as recruitment effects under controlled-observability conditions rather than as artifacts of changing estimation logic.

Figure 1 summarizes this logic as a compact end-to-end flow from latent context to estimated output, separating ego/local evidence, graph-based candidate discovery, recruitment-policy admission, and shared fusion.

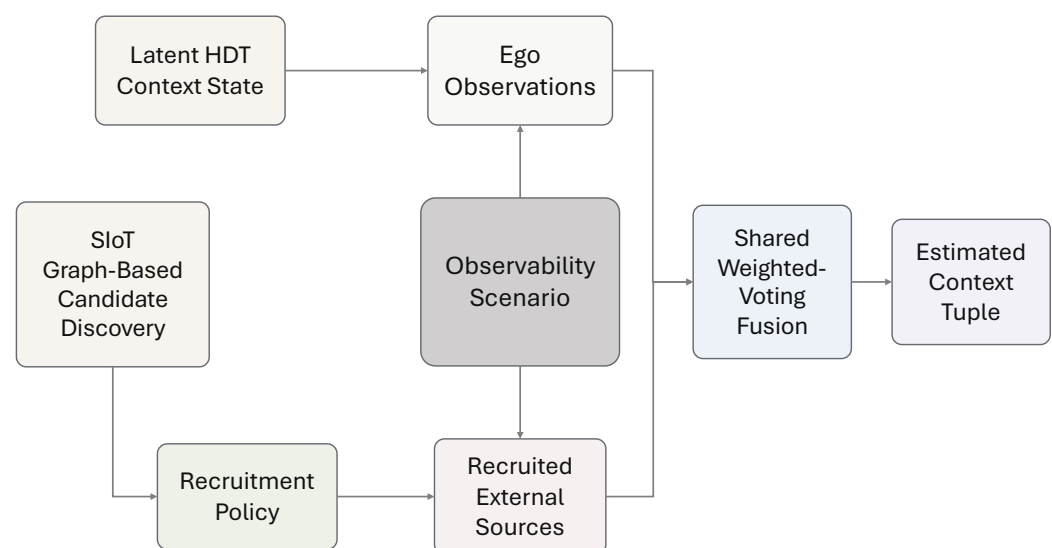


Figure 1. Conceptual workflow of HDT context estimation under SIoT-enabled recruitment. Observability conditions modulate the availability and quality of both ego and external evidence, while SIoT graph-constrained discovery defines the candidate external source set. Recruitment policies determine which external evidence is admitted to a shared weighted-voting fusion stage. Because this fusion mechanism is common across policies, comparative differences are interpreted as recruitment effects under controlled observability conditions.

Section 4 then operationalizes this conceptual workflow in a controlled simulation setting for scenario-conditioned policy comparison.

4. Experimental Evaluation

4.1. Experimental Objectives and Comparative Scope

The experimental objective is to compare recruitment policies for HDT context-layer estimation under explicitly defined partial-observability regimes. The intervention variable is recruitment behavior: all policies share the same fusion mechanism described in Section 3, while only source-discovery and source-selection decisions vary. This isolates policy-level recruitment effects from fusion-model effects.

The evaluation is intentionally comparative and controlled rather than deployment-emulative. The simulator is fully synthetic and rule-based, so the results are interpreted as policy behavior within a controlled methodological testbed, not as direct performance forecasts for real deployments. The accompanying code and processed synthetic outputs are available at: <https://github.com/RobbieGee/siot-hdt-opportunistic-sensing> (accessed on 3 July 2026).

4.2. Simulator Design and Synthetic Data Generation

The main simulator experiment uses 50 episodes of 48 timesteps for each scenario, policy, and seed, with deterministic seeds 7, 11, 13, 17, and 19. The simulator separates latent-state evolution from observation generation. At each timestep t , the latent HDT state is the context tuple $\mathbf{c}_t = (p_t, a_t, e_t, r_t)$ defined in Equation (1). Each component belongs to a finite simulator-defined categorical domain used consistently across all scenarios and recruitment policies. The domains are intentionally compact, enabling controlled and reproducible comparison of recruitment behavior within the synthetic testbed. Across timesteps, each scenario follows a simulator-defined pattern-driven latent context timeline over the 48-step episode horizon, producing alternating stable and change segments.

Internal consistency is maintained by sampling and updating the latent tuple as a single state object before generating observations. Thus, observations are noisy/incomplete views of one underlying joint context, rather than reconstructions from independently generated per-variable labels. Scenario definitions act on observability (what can be inferred from available evidence), not on the semantic definition of the latent tuple itself.

Ego observations are generated conditionally from the current latent state via scenario-conditioned reliability rules. At each timestep, ego evidence may be omitted (missingness), perturbed (noise), or made less discriminative (local ambiguity), and the resulting local evidence quality is summarized by ego-confidence score $q_t \in [0, 1]$. External observations are generated per source with source-specific quality profiles; a source contributes evidence only when it belongs to the recruited set S_t , so external evidence availability depends jointly on source quality and recruitment decisions.

Observability stressors are injected as controlled perturbations. Missingness is increased by scenario-specific omission processes, most strongly in the degraded ego scenario; noise is introduced by perturbing observed evidence; local ambiguity is induced by reducing separability of at least two context states in ego evidence; and external degradation is imposed in the noisy/untrusted external sources scenario by lowering trust, freshness, and accuracy characteristics for a subset of sources. The perturbation probabilities and trigger thresholds are fixed simulator settings shared across all policy comparisons within each scenario, so that observed differences can be attributed to recruitment behavior rather than to scenario-specific retuning. To preserve comparability across policies, all scenario parameters, latent-state domains, source-quality distributions, trigger thresholds, and selection caps are fixed before the policy comparison and kept unchanged across runs, providing stable simulator conditions for controlled and reproducible characterization. The main experimental scale, deterministic seeds, graph parameters, candidate-pool sizes, selection

caps, and discovery radii are reported to support interpretability of the comparison (see Table 1).

Table 1. Key simulator and default SIoT graph settings used in the experiments. Values summarize the default simulation configuration.

Category	Value	Interpretation
Main episodes/timesteps	50/48	Main run per seed, scenario, and policy.
Main seeds	7, 11, 13, 17, 19	Deterministic seeds used for main tables and figures.
Default graph size	61 nodes	One ego node and 60 external source/object nodes.
Mean edges	96.955	Mean typed SIoT edges in the default graph setting.
Path statistics	3.625/5.161	Mean average shortest path/effective diameter.
Ego neighborhood	degree 6.498	Mean ego degree in the default graph setting.
Candidate pools	6.498, 30.498, 53.498, 60	Mean 1-hop, 2-hop, 3-hop, and 4-hop candidate counts.
Selective caps	3/5	Default and adaptive selection caps.
Discovery radii	2/3 hops	Baseline and adaptive hop radii for selective policies.

The main experiment uses four scenarios: nominal, degraded ego sensing, ambiguous local context, and noisy/untrusted external sources. It evaluates four main policies: ego-only sensing, opportunistic-all recruitment, SIoT-aware bounded recruitment, and privacy-aware SIoT recruitment. Separate compact runs are used for the discovery-mode comparison, the source-cap-matched comparison, and scalability analysis. The main policy comparison is based on 50 episodes and five deterministic seeds, whereas the compact diagnostic comparisons use 20 episodes and three deterministic seeds to reduce computational burden while preserving comparable scenario and policy configurations. These compact runs are therefore interpreted as control and diagnostic comparisons, not as exact numerical duplicates of the main 50-episode, five-seed experiment.

Supplementary diagnostic sensitivity sweeps over hop radius, selection cap, privacy penalty, degraded-source fraction, graph density, and graph topology are also provided in the public repository to support inspection of parameter-dependent behavior; these sweeps are treated as supporting diagnostics rather than as the main evidential basis of the manuscript.

4.3. SIoT Graph and External Source Model

External-source accessibility is defined by the SIoT graph introduced in Section 3. Candidate discovery is either exhaustive hop-limited discovery, which evaluates all candidates satisfying Equation (2), or bounded relationship-guided discovery, which expands the frontier using relation priorities and stops once the bounded candidate, edge, score, and variable-coverage conditions are met. The baseline radius is $h_0 = 2$ hops; temporary expansion to $h_t = 3$ and a larger selection cap are allowed only under the adaptive mechanism.

Graph size and topology are simulator settings rather than targets of inference in this work; the study therefore treats them as controlled structural conditions for policy comparison. Operationally, the graph constrains the candidate pool before any policy-specific scoring or filtering is applied.

Each discovered source is represented by synthetic attributes used uniformly across policies: relevance (expected informativeness for the context variables), trust score (expected reliability of contributed evidence), freshness (recency of evidence), latency proxy (coordination/communication burden), privacy cost (exposure penalty), graph distance (hop-based accessibility cost), and relation utility (benefit associated with SIoT relation type/strength). These attributes are generated by the simulator under scenario rules and used as controlled inputs for policy comparison.

For fair comparison, policies are evaluated on the same latent states, graph realizations, source attributes, deterministic seeds, and scenario perturbations at a given episode/timestep/scenario. Candidate availability is then determined by each policy's specified discovery mode, so observed differences in OCA/MVA/F1 and overhead can be interpreted as the effect of policy-specific discovery and recruitment behavior rather than as changes in the underlying data-generating process.

4.4. Observability Scenarios

Four observability scenarios are used to stress different failure modes while preserving comparable experimental structure across policies.

Nominal. Conditions are balanced: ego sensing is moderately reliable, and most external sources are informative. This scenario serves as a baseline for evaluating whether selective recruitment can maintain high-quality context estimation while limiting overhead.

Degraded ego. Ego observability is severely degraded through elevated intermittency and noise (approximately 44% ego missingness), while external sources remain available. This scenario tests how well policies compensate when local sensing becomes unreliable.

Ambiguous local context. Ego signals are structurally insufficient to separate at least two context states even when data are present. Unlike pure missingness, this scenario emphasizes local non-identifiability and evaluates whether external recruitment resolves ambiguity.

Noisy untrusted external. A subset of external sources is degraded in trustworthiness, freshness, and accuracy. This scenario probes robustness when external evidence is plentiful but partially unreliable, stressing the selectivity-versus-coverage trade-off.

Together, these regimes are designed for scenario-conditioned policy comparison rather than reproduction of a specific field deployment.

4.5. Recruitment Policies

All policies use the same weighted-voting fusion mechanism; only recruitment differs.

Ego-only. External recruitment is bypassed entirely ($S_t = \emptyset$), so estimation uses only ego evidence. This defines the local-only baseline with zero recruitment overhead and zero additional privacy exposure.

Opportunistic all. Selective ranking is bypassed: all discovered candidates within the active hop radius are recruited at each timestep. This maximizes source coverage and is treated as a high-coverage, high-cost upper-bound reference rather than as a scalable deployment policy.

SIoT-aware bounded recruitment. This policy uses bounded relationship-guided discovery by default. Discovered candidates are ranked using the multiplicative score in Equation (3), based on source relevance, trust, freshness, latency-sensitive attenuation, relation utility, privacy cost, and access score. Recruitment is bounded by a per-timestep selection cap, so only a limited subset of top-ranked sources is selected. This policy also includes a trigger-based adaptive broadening mechanism: when collapse indicators are detected, specifically elevated ego missingness or the co-occurrence of low ego confidence and an elevated unresolved-variable count, the policy temporarily broadens candidate access and selection budget under bounded predefined settings shared across all runs. The mechanism remains conditional and selective; unlike opportunistic-all recruitment, it does not admit all discovered candidates unconditionally.

Privacy-aware SIoT recruitment. This policy preserves the same discovery, ranking, and selection workflow as SIoT-aware bounded recruitment, but applies stronger privacy penalization within the multiplicative ranking framework. Operationally, this pushes high-privacy-cost candidates down the ranking while retaining selective recruitment.

Source-cap-matched baselines. Two additional policies are used only in the source-cap-matched comparison. The comparison is source-cap-matched rather than total-cost-matched: the random- k , budgeted opportunistic, SIoT-aware bounded, and privacy-aware SIoT policies share the same nominal and adaptive recruitment caps, while discovery mode, realized recruited-source count, and total operational cost remain policy-dependent. The random- k baseline selects a random subset of discovered candidates under this cap. The budgeted opportunistic baseline selects a bounded subset without SIoT-aware relationship scoring. These baselines test whether SIoT-aware ranking improves source selection under comparable recruitment-cap constraints, rather than simply benefiting from a larger recruited-source set.

In the results tables and figures, policy names are abbreviated as follows: EGO = ego-only; OPP-ALL = opportunistic-all; SIoT = SIoT-aware bounded recruitment; P-SIoT = privacy-aware SIoT recruitment; R- k = random- k ; B-OPP = budgeted opportunistic; SIoT-EXH = SIoT-aware exhaustive discovery; SIoT-BND = SIoT-aware bounded discovery; P-SIoT-EXH = privacy-aware SIoT exhaustive discovery; and P-SIoT-BND = privacy-aware SIoT bounded discovery.

4.6. Metrics and Statistical Reporting

Let $K = 4$ denote the number of context variables and T the number of timesteps.

Overall context accuracy (OCA) is the strict joint metric:

$$\text{OCA} = \frac{1}{T} \sum_{t=1}^T \mathbb{I}[\hat{\mathbf{c}}_t = \mathbf{c}_t], \quad (9)$$

which counts a timestep as correct only when all four variables are simultaneously correct.

Mean variable accuracy (MVA) is

$$\text{MVA} = \frac{1}{KT} \sum_{t=1}^T \sum_{k=1}^K \mathbb{I}[\hat{c}_{t,k} = c_{t,k}], \quad (10)$$

which averages variable-wise correctness and is therefore structurally less strict than OCA. Given the simulator domain cardinalities, a uniform random guess over the four context variables would have expected joint accuracy $\text{OCA}_{\text{rand}} = 1/(5 \cdot 5 \cdot 3 \cdot 3) = 0.0044$ and expected mean variable accuracy $\text{MVA}_{\text{rand}} = \frac{1}{4}(1/5 + 1/5 + 1/3 + 1/3) = 0.2667$. These values are reported only as domain-size reference baselines; the simulator latent-state generator itself is pattern-driven and non-uniform.

We additionally report per-variable Macro-F1, convergence time, and operational overhead indicators (recruited-source count, latency proxy, privacy exposure). In this study, convergence time is computed relative to true context-state transition events and reported as the number of timesteps until strict overall context accuracy first becomes 1.0 after the transition. This first-hit recovery measure is used comparatively across policies and should be read as a dynamic-response comparison within the synthetic setup, not as network or deployment latency.

Formally, if τ denotes a true context-transition timestep, convergence time is defined as

$$T_{\text{conv}}(\tau) = \min\{\Delta \geq 0 : \hat{\mathbf{c}}_{\tau+\Delta} = \mathbf{c}_{\tau+\Delta}\}, \quad (11)$$

i.e., the first post-transition timestep at which the full joint context tuple is recovered. For effective-recovery-cost computation, transitions that are not recovered before the end of the episode are treated as censored at the final timestep, and operational costs are accumulated from the transition until either full-tuple recovery or this censoring point. The same censoring rule is applied across all policies and scenarios.

The missingness-oriented analysis (performance versus ego missing rate) is used to characterize sensitivity to local-observability loss and to compare how quickly policy performance separates as ego evidence becomes sparse.

Uncertainty is summarized with bootstrap 95% confidence intervals computed from seed-episode aggregate units with 1000 bootstrap resamples. These intervals are used as comparative uncertainty summaries for this synthetic testbed; they are not used as stand-alone inferential proof of deployment-level superiority.

To isolate the contribution of the trigger-based adaptive broadening mechanism, we additionally ran a targeted ablation within the SIoT-aware and privacy-aware SIoT policies. The same policies were evaluated with adaptive broadening enabled and disabled, suppressing trigger activation while leaving source scoring, fusion, seeds, episode count, timestep horizon, scenario definitions, and bounded-discovery rules unchanged.

4.7. Experimental Results

4.7.1. Main Joint Accuracy by Scenario and Policy

Table 2 and Figure 2 summarize the main experiment. Raw OCA is highest for opportunistic-all in all four scenarios: 0.805 in nominal, 0.774 in degraded ego, 0.743 in ambiguous local context, and 0.303 in noisy/untrusted external sources. This is expected because the policy recruits almost all discovered sources from the graph-constrained candidate pool. Accordingly, opportunistic-all is interpreted as a high-coverage, high-cost accuracy ceiling rather than as a scalable deployment target. By contrast, the bounded SIoT-aware policies intentionally restrict candidate admission, obtaining lower raw OCA but using far fewer sources and substantially lower operational cost, as shown in Table 3.

Table 2. Joint overall context accuracy (OCA) by recruitment policy and scenario for the main experiment. Values are means with bootstrap 95% confidence intervals computed over seed-episode aggregate units. OCA is a strict joint metric that counts a timestep as correct only when all four context variables are simultaneously correct.

Policy	Nominal	Degraded	Ambiguous	Noisy Ext.
EGO	0.140 [0.134, 0.146]	0.021 [0.018, 0.023]	0.052 [0.048, 0.056]	0.083 [0.078, 0.088]
OPP-ALL	0.805 [0.794, 0.816]	0.774 [0.763, 0.785]	0.743 [0.729, 0.756]	0.303 [0.295, 0.311]
SIoT	0.219 [0.211, 0.228]	0.154 [0.146, 0.162]	0.136 [0.130, 0.142]	0.098 [0.092, 0.104]
P-SIoT	0.203 [0.195, 0.210]	0.137 [0.129, 0.144]	0.122 [0.116, 0.128]	0.095 [0.089, 0.101]

The confidence intervals show that the gap between opportunistic-all and the bounded SIoT-aware policies is much larger than the bootstrap uncertainty range in all scenarios. By contrast, the differences between SIoT-aware and privacy-aware SIoT are comparatively small and should be interpreted as descriptive operating-point differences rather than as strong evidence of dominance.

These comparisons should be interpreted with the metric definition in mind. OCA is a strict joint criterion that requires simultaneous correctness of all four context variables at each timestep; therefore, OCA values are structurally lower than variable-wise metrics. For example, in the nominal scenario, SIoT-aware bounded recruitment has OCA = 0.219 and MVA = 0.676, whereas opportunistic-all has OCA = 0.805 and MVA = 0.948. OCA, MVA, and per-variable Macro-F1 are complementary and should not be used interchangeably.

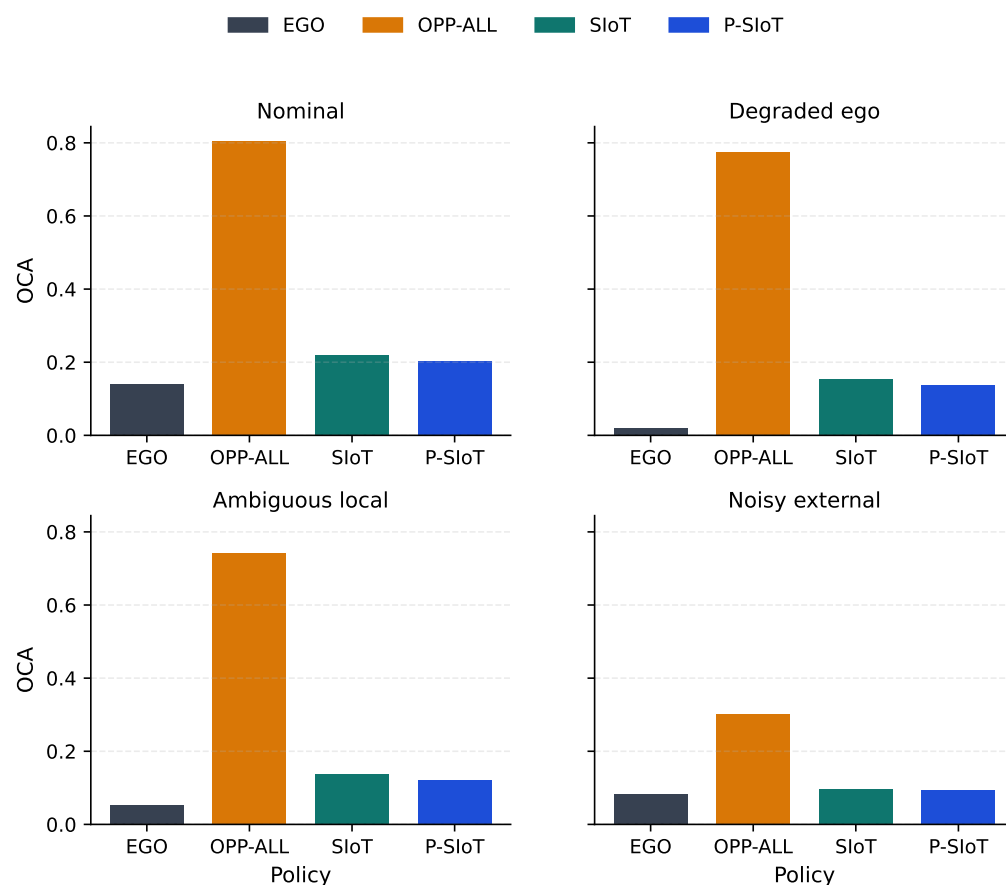


Figure 2. Joint overall context accuracy (OCA) by scenario and recruitment strategy for the main experiment. Policy abbreviations: EGO = ego-only; OPP-ALL = opportunistic-all; SIoT = SIoT-aware bounded recruitment; P-SIoT = privacy-aware SIoT. The high raw OCA of OPP-ALL reflects high source coverage and should be interpreted together with operational cost.

Table 3. Scenario-resolved operational-overhead indicators for the main experiment. Costs are synthetic simulator proxies and support comparative interpretation only.

Scenario	Policy	Src.	Disc.	Rec.	Sel.	Total	Priv.
Nominal	EGO	0.000	0.000	0.000	0.000	0.000	0.000
Nominal	OPP-ALL	53.456	156.645	84.409	0.000	241.054	8.491
Nominal	SIoT	3.208	51.327	4.654	8.142	56.795	1.044
Nominal	P-SIoT	2.982	51.327	4.309	8.142	56.450	0.967
Degraded	EGO	0.000	0.000	0.000	0.000	0.000	0.000
Degraded	OPP-ALL	53.552	157.552	84.858	0.000	242.410	8.507
Degraded	SIoT	4.132	52.928	5.887	8.257	59.640	1.103
Degraded	P-SIoT	3.888	52.928	5.532	8.257	59.286	1.051
Ambiguous	EGO	0.000	0.000	0.000	0.000	0.000	0.000
Ambiguous	OPP-ALL	53.532	156.553	84.637	0.000	241.190	8.487
Ambiguous	SIoT	3.399	51.547	4.926	8.119	57.285	1.062
Ambiguous	P-SIoT	3.119	51.547	4.497	8.119	56.855	0.970
Noisy ext.	EGO	0.000	0.000	0.000	0.000	0.000	0.000
Noisy ext.	OPP-ALL	53.452	156.615	88.656	0.000	245.271	8.457
Noisy ext.	SIoT	2.088	52.120	3.105	8.194	56.045	0.718
Noisy ext.	P-SIoT	1.475	52.120	2.182	8.194	55.122	0.504

4.7.2. Operational Cost and Effective Recovery Cost

The operational trade-off is reported in Table 3. Opportunistic-all recruits roughly 53.5 sources per timestep in the default 60-source graph, with total operational cost above 241 in all scenarios. By contrast, SIoT-aware bounded policies recruit about 1.5–4.1 sources and keep total operational cost near 55–60. In the noisy/untrusted external sources scenario, privacy-aware SIoT recruits 1.475 sources on average and has privacy cost 0.504, compared with 53.452 sources and privacy cost 8.457 for opportunistic-all. The raw-OCA advantage of opportunistic-all is therefore primarily a coverage effect, about 53 recruited sources per timestep versus about 1.5–4.1 for bounded SIoT-aware policies in the default graph, and the same near-exhaustive recruitment behavior explains why it is useful as an accuracy ceiling but becomes impractical in the scalability analysis.

For readability, Table 3 reports point estimates for overhead indicators; bootstrap confidence intervals for the overhead components are included in the processed repository outputs. The main overhead differences are large in magnitude, with opportunistic-all operating at total costs above 241 in all scenarios, whereas bounded SIoT-aware policies remain near 55–60.

Effective recovery cost combines recovery dynamics with operational cost after true context transitions. Figure 3 shows that lower raw source use does not automatically minimize recovery cost: bounded policies spend less per timestep, but their longer recovery windows can increase effective recovery cost in degraded and ambiguous regimes. In these cases, opportunistic-all can be more favorable under the specific recovery-cost metric despite its much higher continuous operational burden. Bounded SIoT recruitment should therefore be interpreted as a scalable low-overhead operating point, not as a universally recovery-optimal policy. This reinforces the need to interpret accuracy and cost jointly rather than treating either metric alone as decisive.

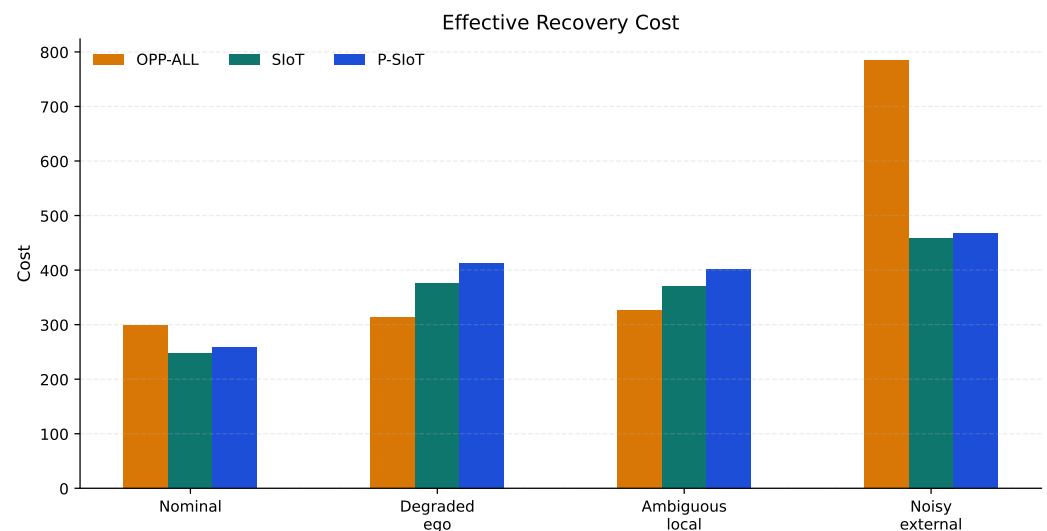


Figure 3. Effective recovery cost by policy and scenario. Values summarize synthetic operational cost accumulated during recovery after context transitions; they are not measured deployment latency or energy costs. EGO is omitted because its recruitment-related operational cost is identically zero by construction.

4.7.3. Source-Cap-Matched and Discovery-Mode Comparisons

Figure 4 reports the source-cap-matched comparison and provides a fairer test of SIoT-aware selection against non-SIoT-aware selection under comparable source-use constraints. This comparison is source-cap-matched rather than total-cost-matched: all policies share the same nominal and adaptive recruitment caps, while discovery mode, realized

recruitment count, and total operational cost remain policy-dependent. In the nominal scenario, SIoT-aware bounded recruitment reaches $OCA = 0.222$ with total cost 56.497, compared with $OCA = 0.190$ and cost 95.015 for random- k , and $OCA = 0.204$ and cost 94.789 for budgeted opportunistic recruitment. In degraded ego, SIoT-aware recruitment reaches $OCA = 0.154$ and OCA -per-total-cost 0.00260, compared with 0.126 and 0.00088 for budgeted opportunistic recruitment. In the noisy/untrusted external sources case, non-SIoT-bounded baselines obtain slightly higher raw OCA , but at roughly twice the total cost; SIoT-aware policies keep the better OCA -per-total-cost ratios. Accordingly, OCA per total cost should be interpreted as an empirical cost-effectiveness ratio under the implemented cost model, not as performance under equal total expenditure.

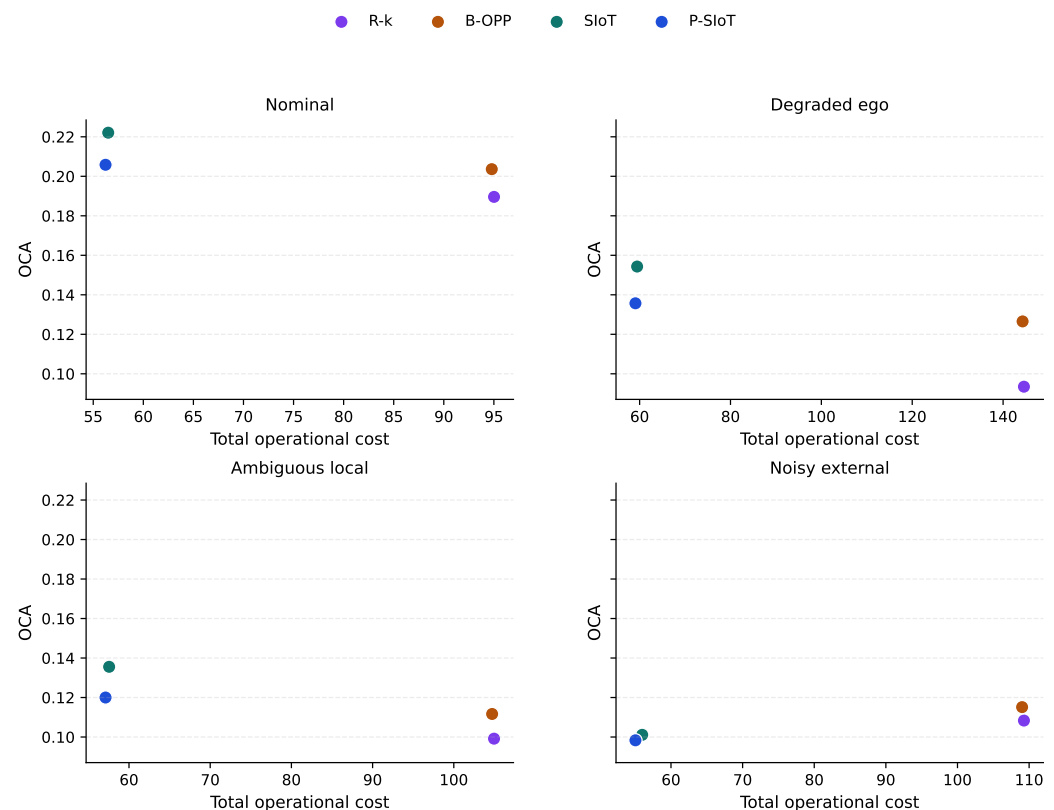


Figure 4. Source-cap-matched OCA and cost comparison. The comparison evaluates SIoT-aware selection against non-SIoT-aware bounded baselines under comparable recruitment caps; total operational cost remains policy-dependent because discovery and scoring costs differ across policies.

Figure 5 reports the discovery-mode comparison, separating SIoT-aware ranking from the cost of discovering candidates. Bounded discovery preserves almost the same OCA as exhaustive SIoT discovery while substantially reducing the candidate pool and total cost. For SIoT-aware recruitment in the nominal compact run, exhaustive discovery gives $OCA = 0.236$ and cost = 94.959, whereas bounded discovery gives $OCA = 0.236$ and cost = 56.369. In degraded ego, the corresponding values are $OCA = 0.149$ and cost = 144.926 for exhaustive discovery, versus $OCA = 0.148$ and cost = 59.483 for bounded discovery. This supports bounded relationship-guided discovery as the scalable SIoT operating mode.

4.7.4. Pareto and Scalability Analysis

Figure 6 provides the Pareto view of OCA versus total operational cost and further clarifies why raw OCA alone is insufficient. Opportunistic-all lies on the high-accuracy/high-cost part of the frontier, while bounded SIoT policies provide lower-cost operating points with substantially higher OCA per recruited source. This does not mean that SIoT-aware

policies maximize raw accuracy; rather, they change the feasible trade-off when operational cost and privacy exposure are part of the evaluation.

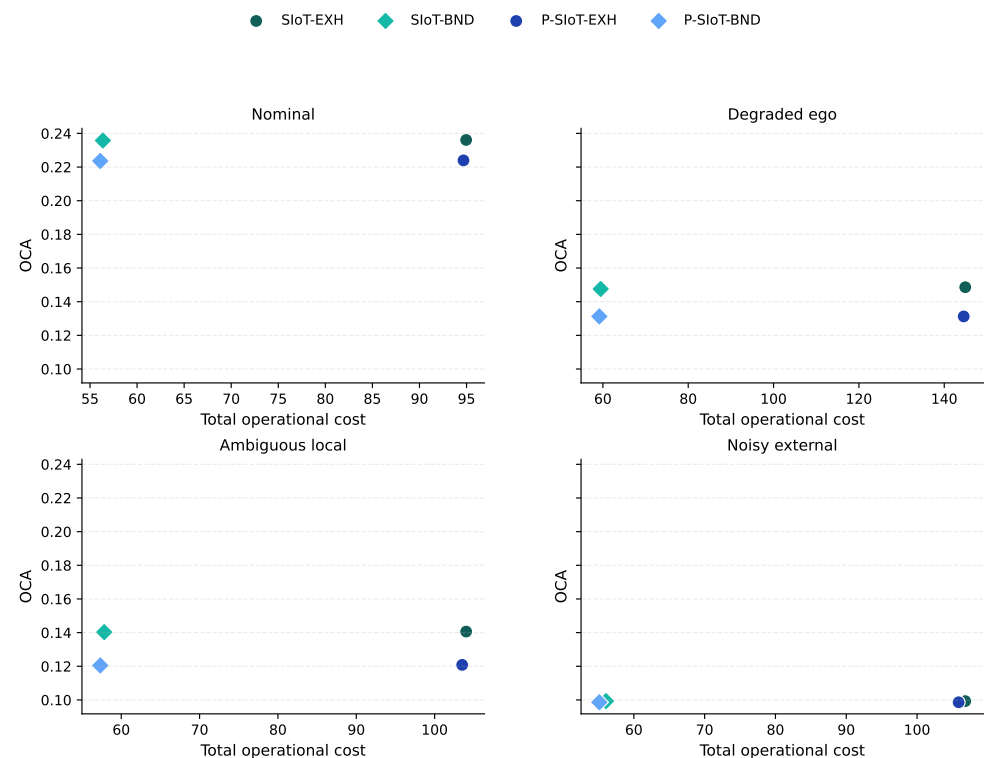


Figure 5. Exhaustive versus bounded SIoT discovery. The compact comparison uses a different run configuration and sample size from Table 2; it should therefore be interpreted as a discovery-mode comparison, not as a direct duplicate of the main table.

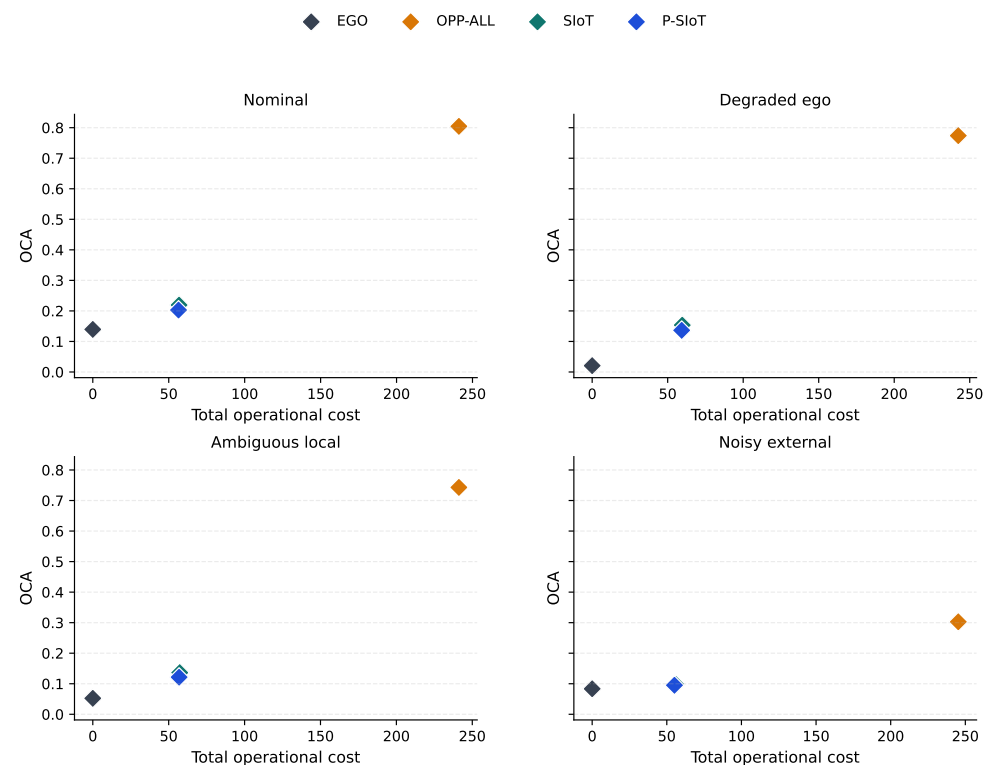


Figure 6. Pareto view of OCA versus total operational cost for the main policy set. The plot distinguishes the high-coverage accuracy ceiling from lower-cost bounded SIoT operating points.

Figure 7 reports the scalability analysis and shows the strongest separation between exhaustive opportunistic recruitment and bounded SIoT discovery. At graph size 800, opportunistic-all recruits roughly 606 sources and reaches total operational cost above 10,000 in all scenarios. Bounded SIoT policies keep the discovered candidate count near eight and the recruited-source count near two to four, with total operational cost around 79–88 depending on scenario and policy. Thus, opportunistic-all becomes impractical as graph size grows, even when its raw OCA remains high, whereas bounded SIoT discovery keeps discovery, scoring, and recruitment bounded.

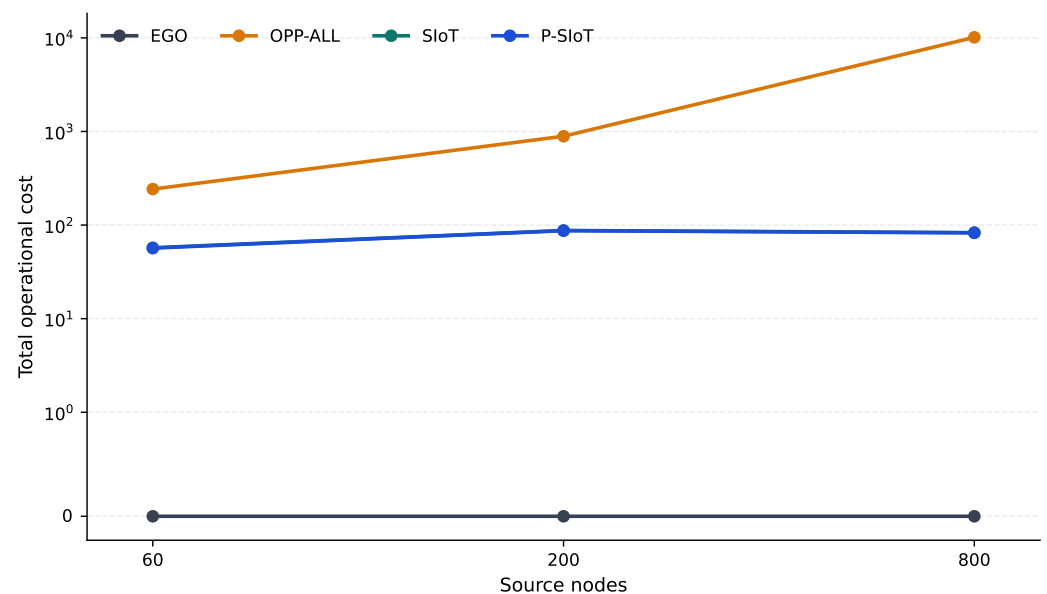


Figure 7. Scalability of total operational cost versus graph size. The y-axis uses a logarithmic scale to show both the high-cost opportunistic-all reference and the bounded SIoT policies. Exhaustive opportunistic recruitment grows rapidly with graph size, whereas bounded SIoT discovery keeps candidate scoring and recruitment bounded.

4.7.5. Per-Variable Error Analysis

Figure 8 reports per-variable Macro-F1 scores and shows that the four context variables are not equally estimable from the available source pool, and that the hardest variable changes by scenario and policy. In the nominal and noisy/untrusted external-source scenarios, the selective SIoT-aware policies show their lowest or near-lowest Macro-F1 on *Resource State* and *Environmental Load*, whereas opportunistic-all remains consistently stronger because it recruits many more sources. In an ambiguous local context, *Place* becomes particularly difficult for the SIoT-aware policies, with Macro-F1 values of 0.481 and 0.461 for SIoT-aware and privacy-aware SIoT, respectively; the corresponding *Place* misclassification rates are 0.486 and 0.507. These patterns explain why strict joint OCA can remain low even when variable-wise metrics are more moderate.

Although *Environmental Load* sensing is widely represented in the generated SIoT graphs, bounded recruitment often yields comparatively little redundant evidence for this variable. This behavior is not caused by a shortage of *Environmental Load*-capable sources, but by the interaction of source-level ranking, bounded candidate admission, scenario-dependent evidence quality, and the absence of an explicit per-variable redundancy term in recruitment. The *Environmental Load* Macro-F1 should therefore be interpreted as a consequence of selective recruitment and scenario-dependent evidence quality rather than as evidence that the generated graph lacks capable sources.

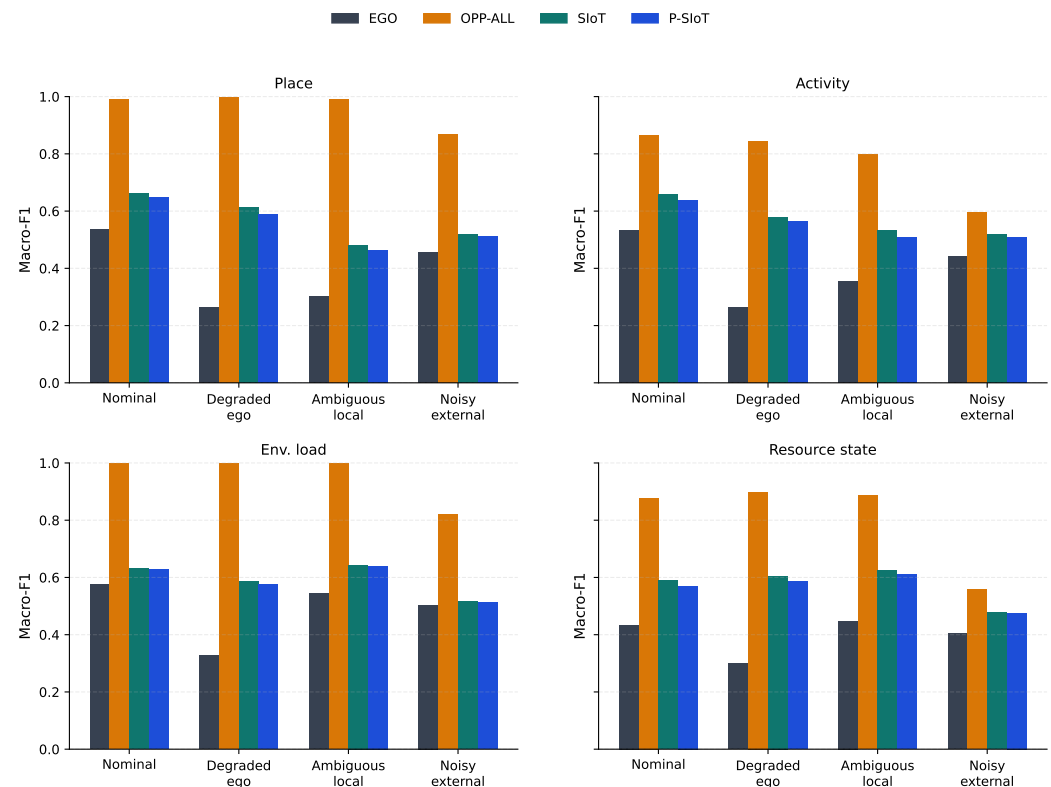


Figure 8. Per-variable Macro-F1 scores for the four HDT context variables by policy and scenario. The spread across variables and scenarios indicates that estimability depends on both the specific context variable and the active observability regime.

4.7.6. Adaptive Mechanism Diagnostics

The adaptive broadening ablation shows scenario-dependent but limited benefits. For SIoT-aware recruitment, enabling adaptive broadening raises OCA from 0.084 to 0.154 in degraded ego, from 0.105 to 0.136 in ambiguous local context, from 0.197 to 0.219 in nominal, and from 0.095 to 0.098 in noisy/untrusted external sources. The gain is therefore largest when ego sensing collapses and modest when external unreliability dominates. Adaptive activation rates are 0.158 in nominal, 0.706 in degraded ego, 0.268 in ambiguous local context, and 0.312 in noisy/untrusted external sources (see Table 4).

Table 4. Adaptive broadening ablation. Disabled/enabled columns report OCA means with bootstrap 95% confidence intervals computed over seed-episode aggregate units; source count and activation rate refer to the adaptive-enabled condition.

Scenario	Policy	Disabled	Enabled	Sources	Activation
Nominal	SIoT	0.197 [0.189, 0.205]	0.219 [0.211, 0.228]	3.208	0.158
Nominal	P-SIoT	0.182 [0.175, 0.189]	0.203 [0.196, 0.210]	2.982	0.158
Degraded	SIoT	0.084 [0.078, 0.090]	0.154 [0.145, 0.161]	4.132	0.706
Degraded	P-SIoT	0.075 [0.069, 0.080]	0.137 [0.129, 0.144]	3.888	0.706
Ambiguous	SIoT	0.105 [0.100, 0.111]	0.136 [0.131, 0.143]	3.399	0.268
Ambiguous	P-SIoT	0.095 [0.090, 0.100]	0.122 [0.115, 0.127]	3.119	0.268

Table 4. Cont.

Scenario	Policy	Disabled	Enabled	Sources	Activation
Noisy ext.	SIoT	0.095 [0.090, 0.101]	0.098 [0.093, 0.104]	2.088	0.312
Noisy ext.	P-SIoTT	0.090 [0.084, 0.095]	0.095 [0.089, 0.101]	1.475	0.312

The confidence intervals indicate that adaptive broadening produces its clearest gains in the degraded ego and ambiguous local-context scenarios, whereas the noisy/untrusted external-source case shows only a small enabled-versus-disabled difference with overlapping intervals.

These diagnostics support a bounded interpretation of adaptive broadening. It helps most when local ego sensing is degraded, but it does not close the raw-OCA gap to opportunistic-all. Its role is therefore to mitigate collapse within a bounded SIoT operating point, not to make selective recruitment equivalent to all-source recruitment.

5. Discussion

5.1. Scenario-Conditioned Policy Behavior

Within the boundaries of this controlled synthetic evaluation, policy performance is strongly scenario-dependent once accuracy, cost, privacy, and scalability are considered together. Raw OCA is highest for opportunistic-all, but that result follows from recruiting almost every discovered source and should be interpreted as an upper-bound reference. The more informative comparison for deployable SIoT recruitment is the trade-off between bounded SIoT discovery and non-SIoT bounded baselines under comparable recruitment-cap constraints.

The trigger-based adaptive broadening mechanism embedded in the SIoT-aware policies behaves consistently with its intended degradation-response role, but its benefit is bounded. It improves OCA most clearly under degraded ego sensing and more modestly in the other scenarios. The mechanism should therefore be interpreted as a collapse-mitigation heuristic within a bounded SIoT operating point, not as evidence of uniformly optimal adaptation.

5.2. Coverage–Selectivity Trade-Off

The central design tension is between coverage and selectivity. Opportunistic-all recruitment demonstrates how much raw accuracy can be obtained by using almost all available sources, but it also exposes why such a policy is not the practical target in a growing SIoT graph. Bounded SIoT-aware recruitment uses relation semantics to constrain discovery and rank candidates, thereby trading some raw OCA for substantially lower source use, privacy exposure, and total operational cost.

This interpretation is consistent with the simulated dynamics of trust, freshness, graph reachability, and missingness, but it remains scenario-conditional. The present results support comparative reasoning about recruitment strategies in a synthetic testbed; they do not by themselves justify deployment-level superiority claims.

5.3. Joint vs. Variable-Wise Accuracy and Design Implications

A key methodological implication is the structural difference between strict joint OCA and variable-wise metrics (MVA and Macro-F1). Because OCA requires all variables to be simultaneously correct, it is systematically lower and reflects a different reliability criterion. This is especially important for selective policies: moderate variable-wise performance can still yield low joint tuple accuracy when errors are distributed across different variables.

Reporting only variable-wise performance would therefore overstate effective context-state correctness for HDT services that require full joint consistency.

Within this controlled setting, the evidence supports scenario-conditional guidance rather than universal prescriptions: opportunistic-all establishes the high-cost accuracy ceiling; SIoT-aware bounded discovery provides a lower-cost operating point; source-cap-matched comparisons are needed to compare selection quality fairly; and adaptive broadening remains a heuristic direction requiring stronger validation. For downstream HDT services, the practical implication is not that one recruitment policy is universally best, but that HDT context layers should expose accuracy, cost, privacy, and scalability as joint design variables.

6. Limitations

These findings should therefore be interpreted under several explicit methodological constraints.

Fully synthetic environment. All observations, source behaviors, graph relations, trust values, freshness effects, latency proxies, privacy costs, and context labels are generated by simulator rules. No real sensor traces, human-subject data, user studies, or field deployments are included. Accordingly, conclusions about comparative behavior in this controlled environment are not directly generalizable to real HDT deployments.

Stylized source and trust models. Source quality, trust dynamics, typed relations, missingness, and privacy/latency costs are parameterized abstractions rather than empirically calibrated distributions. Alternative parameter settings could change policy ordering or alter the observed coverage–selectivity crossover.

Controlled scale and synthetic graph generation. The main experiment uses 50 episodes, 48 timesteps, and five deterministic seeds, and scalability is evaluated up to 800 external source nodes. Nevertheless, graph topology and source attributes remain simulator-generated. Bootstrap intervals provide uncertainty summaries for the synthetic testbed, but the study remains exploratory/comparative rather than statistically definitive for real deployments.

Fixed fusion mechanism. All policies share the same interpretable weighted-voting fusion model. This isolates recruitment effects, but conclusions may not transfer unchanged to more expressive fusion approaches.

Heuristic adaptive mechanism. The trigger-based adaptive broadening mechanism embedded in the SIoT-aware selective policy is intentionally interpretable and rule-based. It is not globally tuned or optimized, and the current results should not be interpreted as evidence of optimal adaptation.

No deployment latency, energy, or human-in-the-loop validation. The latency and cost variables are simulator proxies. The study does not measure network latency, energy use, user acceptability, human-in-the-loop decision quality, or downstream HDT service outcomes.

Future work should calibrate source reliability, missingness, graph topology, latency, and privacy-cost models with real mobile, wearable, environmental, or deployed SIoT/HDT datasets. It should also compare the present interpretable recruitment baseline with learned RL policies, graph-neural-network-based source selection, probabilistic graphical context-fusion models, and human-in-the-loop HDT validation protocols. Benchmarking these alternatives under matched observability, recruitment-cap, and cost assumptions would clarify when interpretable SIoT recruitment is sufficient and when learned or probabilistic fusion models are warranted.

7. Conclusions

Overall, this paper presented a controlled synthetic comparison of SIoT-enabled recruitment strategies for HDT context-layer estimation under partial observability. The evaluation makes the SIoT graph layer explicit, uses bounded relationship-guided discovery, separates discovery/recruitment/selection costs, and includes source-cap-matched and scalability-aware comparisons.

In this evaluation, opportunistic-all achieves the highest raw OCA because it recruits many sources and is therefore best interpreted as a high-cost accuracy ceiling. SIoT-aware bounded policies do not maximize raw OCA; instead, they provide lower-cost operating points, improve cost-effectiveness in source-cap-matched comparisons, and preserve nearly the same OCA as exhaustive SIoT discovery while strongly reducing discovery and total operational cost. The adaptive broadening ablation further shows that trigger-based expansion is most beneficial under degraded ego sensing, where it substantially improves OCA while preserving bounded recruitment; however, it remains a collapse-mitigation heuristic rather than an optimal adaptation mechanism. Scalability results show that exhaustive opportunistic recruitment becomes impractical as graph size grows, whereas bounded SIoT discovery keeps candidate scoring and recruitment bounded. The results also reaffirm that strict joint OCA and variable-wise metrics should be reported together for an accurate interpretation of context-estimation quality. The per-variable analysis further suggests that future SIoT-aware recruitment policies should consider explicit variable-level coverage or redundancy terms, since source-level ranking alone may leave some context variables less supported even when capable sources exist in the graph.

Future work should prioritize empirical calibration with real sensing datasets, deployment-level latency and energy measurements, human-in-the-loop HDT validation, and comparisons with RL, GNN, and probabilistic context-fusion policies. The public simulator and processed synthetic outputs are provided to support reproducibility and to make these extensions easier to benchmark.

Author Contributions: Conceptualization, R.G. and L.B.; methodology, R.G., A.M. and L.B.; software, L.B.; validation, L.B., A.M. and R.G.; formal analysis, L.B. and R.G.; investigation, L.B., A.M. and R.G.; writing—original draft preparation, R.G. and L.B.; writing—review and editing, R.G., A.M., R.P., M.A. and D.G.; supervision, R.G.; project administration, R.G. All authors have read and agreed to the published version of the manuscript.

Funding: This research received no external funding.

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Data Availability Statement: The data and code supporting the findings of this study are publicly available in the accompanying GitHub repository: <https://github.com/RobbieGee/siot-hdt-opportunistic-sensing> (accessed on 3 July 2026). The repository includes the synthetic simulator, processed synthetic outputs used to generate the tables and figures reported in the manuscript, meta-data, and scripts for reproducing the reported plots and summary statistics. No personal, sensitive, proprietary, or real-world human-subject data were used in this study.

Conflicts of Interest: The authors declare no conflicts of interest.

References

1. Minerva, R.; Lee, G.M.; Crespi, N. Digital Twin in the IoT Context: A Survey on Technical Features, Scenarios, and Architectural Models. *Proc. IEEE* **2020**, *108*, 1785–1824. [CrossRef]
2. Lin, Y.; Chen, L.; Ali, A.; Nugent, C.; Cleland, I.; Li, R.; Gao, D.; Wang, H.; Wang, Y.; Ning, H. Human Digital Twin: A Survey. *J. Cloud Comput.* **2024**, *13*, 131. [CrossRef]

3. Asad, U.; Khan, M.; Khalid, A.; Lughmani, W.A. Human-Centric Digital Twins in Industry: A Comprehensive Review of Enabling Technologies and Implementation Strategies. *Sensors* **2023**, *23*, 3938. [\[CrossRef\]](#) [\[PubMed\]](#)
4. Krupas, M.; Kajati, E.; Liu, C.; Zolotova, I. Towards a Human-Centric Digital Twin for Human–Machine Collaboration: A Review on Enabling Technologies and Methods. *Sensors* **2024**, *24*, 2232. [\[CrossRef\]](#) [\[PubMed\]](#)
5. De Simone, F.; Casillo, M.; Presta, R. When Technology Becomes a Barrier: The Accessibility Issue in Mobility Interfaces. In *Proceedings of the International Conference on Human-Computer Interaction*; Lecture Notes in Computer Science; Springer: Berlin/Heidelberg, Germany, 2025; Volume 16336, pp. 229–246. [\[CrossRef\]](#)
6. Presta, R.; De Simone, F.; Tancredi, C.; Chiesa, S. Nudging the Safe Zone: Design and Assessment of HMI Strategies Based on Intelligent Driver State Monitoring Systems. In *Proceedings of the International Conference on Human-Computer Interaction*; Lecture Notes in Computer Science; Springer Nature: Cham, Switzerland, 2023; Volume 14048, pp. 166–185. [\[CrossRef\]](#)
7. Yurur, O.; Liu, C.H.; Sheng, Z.; Leung, V.C.M.; Moreno, W.; Leung, K.K. Context-Awareness for Mobile Sensing: A Survey and Future Directions. *IEEE Commun. Surv. Tutor.* **2016**, *18*, 68–93. [\[CrossRef\]](#)
8. Liu, Y.; Kong, L.; Chen, G. Data-Oriented Mobile Crowdsensing: A Comprehensive Survey. *IEEE Commun. Surv. Tutor.* **2019**, *21*, 2849–2885. [\[CrossRef\]](#)
9. Atzori, L.; Iera, A.; Morabito, G.; Nitti, M. The Social Internet of Things (SIoT)—When Social Networks Meet the Internet of Things: Concept, Architecture and Network Characterization. *Comput. Netw.* **2012**, *56*, 3594–3608. [\[CrossRef\]](#)
10. Malekshahi Rad, M.; Rahmani, A.; Sahafi, A.; Qader, N. Social Internet of Things: Vision, Challenges, and Trends. *Hum. Centric Comput. Inf. Sci.* **2020**, *10*, 52. [\[CrossRef\]](#)
11. Nitti, M.; Girau, R.; Atzori, L. Trustworthiness Management in the Social Internet of Things. *IEEE Trans. Knowl. Data Eng.* **2014**, *26*, 1253–1266. [\[CrossRef\]](#)
12. Alam, S.; Zardari, S.; Noor, S.; Ahmed, S.; Mouratidis, H. Trust Management in Social Internet of Things (SIoT): A Survey. *IEEE Access* **2022**, *10*, 108924–108954. [\[CrossRef\]](#)
13. Chahal, R.K.; Kumar, N.; Batra, S. Trust Management in Social Internet of Things: A Taxonomy, Open Issues, and Challenges. *Comput. Commun.* **2020**, *150*, 13–46. [\[CrossRef\]](#)
14. Khan, W.Z.; Arshad, Q.; Hakak, S.; Khan, M.; Rehman, S.U. Trust Management in Social Internet of Things: Architectures, Recent Advancements, and Future Challenges. *IEEE Internet Things J.* **2020**, *8*, 7768–7788. [\[CrossRef\]](#)
15. Fu, Y.; Lu, S.; Chen, J.; Liu, X.; Huang, B. Socially-Aware and Privacy-Preserving Multi-Objective Worker Recruitment in Mobile Crowd Sensing. *Peer-to-Peer Netw. Appl.* **2024**, *17*, 1001–1019. [\[CrossRef\]](#)
16. Girau, R.; Cossu, R.; Farina, M.; Pilloni, V.; Atzori, L. Virtual User in the IoT: Definition, Technologies and Experiments. *Sensors* **2019**, *19*, 4489. [\[CrossRef\]](#) [\[PubMed\]](#)
17. Zhao, S.; Qi, G.; He, T.; Chen, J.; Liu, Z.; Wei, K. A Survey of Sparse Mobile Crowdsensing: Developments and Opportunities. *IEEE Open J. Comput. Soc.* **2022**, *3*, 73–85. [\[CrossRef\]](#)
18. Atzori, L.; Iera, A.; Morabito, G. SIoT: Giving a Social Structure to the Internet of Things. *IEEE Commun. Lett.* **2011**, *15*, 1193–1195. [\[CrossRef\]](#)
19. Farhadi, B.; Rahmani, A.M.; Asghari, P.; Hosseinzadeh, M. Friendship Selection and Management in Social Internet of Things: A Systematic Review. *Comput. Netw.* **2021**, *201*, 108568. [\[CrossRef\]](#)
20. Amin, F.; Majeed, A.; Mateen, A.; Abbasi, R.; Hwang, S.O. A Systematic Survey on the Recent Advancements in the Social Internet of Things. *IEEE Access* **2022**, *10*, 63867–63884. [\[CrossRef\]](#)
21. Khelloufi, A.; Ning, H.; Dhelim, S.; Qiu, T.; Ma, J.; Huang, R.; Atzori, L. A Social-Relationships-Based Service Recommendation System for SIoT Devices. *IEEE Internet Things J.* **2020**, *8*, 1859–1870. [\[CrossRef\]](#)
22. Chen, R.; Bao, F.; Guo, J. Trust-Based Service Management for Social Internet of Things Systems. *IEEE Trans. Dependable Secur. Comput.* **2015**, *13*, 684–696. [\[CrossRef\]](#)
23. Sagar, S.; Mahmood, A.; Sheng, Q.Z.; Zhang, W.E.; Zhang, Y.; Pabani, J.K. Understanding the Trustworthiness Management in the Social Internet of Things: A Survey. *Comput. Netw.* **2024**, *251*, 110611. [\[CrossRef\]](#)
24. Dong, P.; Ge, J.; Wang, X.; Guo, S. Collaborative Edge Computing for Social Internet of Things: Applications, Solutions, and Challenges. *IEEE Trans. Comput. Soc. Syst.* **2021**, *9*, 291–301. [\[CrossRef\]](#)
25. Roopa, M.S.; Pattar, S.; Buyya, R.; Venugopal, K.R.; Iyengar, S.S.; Patnaik, L.M. Social Internet of Things (SIoT): Foundations, Thrust Areas, Systematic Review and Future Directions. *Comput. Commun.* **2019**, *139*, 32–57. [\[CrossRef\]](#)
26. Zhong, R.; Hu, B.; Feng, Y.; Zheng, H.; Hong, Z.; Lou, S.; Tan, J. Construction of Human Digital Twin Model Based on Multimodal Data and Its Application in Locomotion Mode Identification. *Chin. J. Mech. Eng.* **2023**, *36*, 126. [\[CrossRef\]](#)
27. Miller, M.E.; Spatz, E. A Unified View of a Human Digital Twin. *Hum.-Intell. Syst. Integr.* **2022**, *4*, 23–33. [\[CrossRef\]](#)
28. Gaffinet, B.; Al Haj Ali, J.; Naudet, Y.; Panetto, H. Human Digital Twins: A Systematic Literature Review and Concept Disambiguation for Industry 5.0. *Comput. Ind.* **2025**, *166*, 104230. [\[CrossRef\]](#)

29. Presta, R.; Tancredi, C.; Mancuso, L.; Caccavale, D.; Riccio, F.M. Co-Designing with Experts: Exploring Scenarios for Haptic-Enhanced Virtual Reality in Automotive HMI Design. In *Proceedings of the International Conference on Human-Computer Interaction; Lecture Notes in Computer Science; Springer Nature: Cham, Switzerland, 2025; Volume 15817, pp. 76–95. [CrossRef]*
30. Hu, Z.; Lou, S.; Xing, Y.; Wang, X.; Cao, D.; Lv, C. Review and Perspectives on Driver Digital Twin and Its Enabling Technologies for Intelligent Vehicles. *IEEE Trans. Intell. Veh.* **2022**, *7*, 417–440. [\[CrossRef\]](#)
31. Wang, Z.; Gupta, R.; Han, K.; Wang, H.; Ganlath, A.; Ammar, N.; Tiwari, P. Mobility Digital Twin: Concept, Architecture, Case Study, and Future Challenges. *IEEE Internet Things J.* **2022**, *9*, 17452–17467. [\[CrossRef\]](#)
32. Schwarz, C.; Wang, Z. The Role of Digital Twins in Connected and Automated Vehicles. *IEEE Intell. Transp. Syst. Mag.* **2022**, *14*, 41–51. [\[CrossRef\]](#)
33. Presta, R.; De Simone, F.; Bacchiani, L.; Girau, R. Trustworthy Companion AI for Human-Aware Transition of Control: Motivation, Architecture, and Research Roadmap. *Technologies* **2026**, *14*, 386. [\[CrossRef\]](#)
34. Ganti, R.K.; Ye, F.; Lei, H. Mobile Crowdsensing: Current State and Future Challenges. *IEEE Commun. Mag.* **2011**, *49*, 32–39. [\[CrossRef\]](#)
35. Capponi, A.; Fiandrino, C.; Kantarci, B.; Foschini, L.; Kliazovich, D.; Bouvry, P. A Survey on Mobile Crowdsensing Systems: Challenges, Solutions, and Opportunities. *IEEE Commun. Surv. Tutor.* **2019**, *21*, 2419–2465. [\[CrossRef\]](#)
36. Amintoosi, H.; Kanhere, S.S.; Allahbakhsh, M. Trust-Based Privacy-Aware Participant Selection in Social Participatory Sensing. *J. Inf. Secur. Appl.* **2015**, *20*, 11–25. [\[CrossRef\]](#)
37. Wang, L.; Yang, D.; Yu, Z.; Han, Q.; Wang, E.; Zhou, K.; Guo, B. Acceptance-Aware Mobile Crowdsourcing Worker Recruitment in Social Networks. *IEEE Trans. Mob. Comput.* **2021**, *22*, 634–646. [\[CrossRef\]](#)
38. Zhang, D.; Ma, Y.; Hu, X.S.; Wang, D. Toward Privacy-Aware Task Allocation in Social Sensing-Based Edge Computing Systems. *IEEE Internet Things J.* **2020**, *7*, 11384–11400. [\[CrossRef\]](#)
39. Li, T.; Jung, T.; Li, H.; Cao, L.; Wang, W.; Li, X.Y.; Wang, Y. Scalable Privacy-Preserving Participant Selection in Mobile Crowd Sensing. In *Proceedings of the Annual IEEE International Conference on Pervasive Computing and Communications*; IEEE: New York, NY, USA, 2017; pp. 59–68. [\[CrossRef\]](#)
40. Atzori, L.; Girau, R.; Pilloni, V.; Uras, M. Assignment of Sensing Tasks to IoT Devices: Exploitation of a Social Network of Objects. *IEEE Internet Things J.* **2019**, *6*, 2679–2692. [\[CrossRef\]](#)
41. Rosenberger, J.; Urlaub, M.; Rauterberg, F.; Lutz, T.; Selig, A.; Bühren, M.; Schramm, D. Deep Reinforcement Learning Multi-Agent System for Resource Allocation in Industrial Internet of Things. *Sensors* **2022**, *22*, 4099. [\[CrossRef\]](#) [\[PubMed\]](#)
42. Sivamayil, K.; Rajasekar, E.; Aljafari, B.; Nikolovski, S.; Vairavasundaram, S.; Vairavasundaram, I. A Systematic Study on Reinforcement Learning Based Applications. *Energies* **2023**, *16*, 1512. [\[CrossRef\]](#)
43. Mali, S.; Zeng, F.; Adhikari, D.; Ullah, I.; Al-Khasawneh, M.A.; Alfarraj, O.; Alblehai, F. Federated Reinforcement Learning-Based Dynamic Resource Allocation and Task Scheduling in Edge for IoT Applications. *Sensors* **2025**, *25*, 2197. [\[CrossRef\]](#) [\[PubMed\]](#)
44. Le, M.; Huynh-The, T.; Do-Duy, T.; Vu, T.H.; Hwang, W.J.; Pham, Q.V. Applications of Distributed Machine Learning for the Internet-of-Things: A Comprehensive Survey. *IEEE Commun. Surv. Tutor.* **2023**, *27*, 1053–1100. [\[CrossRef\]](#)
45. Dong, G.; Tang, M.; Wang, Z.; Gao, J.; Guo, S.; Cai, L.; Gutierrez, R.; Campbell, B.; Barnes, L.E.; Boukhechba, M. Graph Neural Networks in IoT: A Survey. *ACM Trans. Sens. Netw.* **2022**, *19*, 1–50. [\[CrossRef\]](#)
46. Hamrouni, A.; Khanfor, A.; Ghazzai, H.; Massoud, Y. Context-Aware Service Discovery: Graph Techniques for IoT Network Learning and Socially Connected Objects. *IEEE Access* **2022**, *10*, 107330–107345. [\[CrossRef\]](#)
47. Ullah, I.; Kim, J.B.; Han, Y.H. Compound Context-Aware Bayesian Inference Scheme for Smart IoT Environment. *Sensors* **2022**, *22*, 3022. [\[CrossRef\]](#) [\[PubMed\]](#)
48. Han, V.Y.; Wang, T.; Cho, H.; Todi, K.; Fernandes, A.S.; Levi, A.; Zhang, Z.; Grossman, T.; Ion, A.; Jonker, T.R. A Dynamic Bayesian Network Based Framework for Multimodal Context-Aware Interactions. In *Proceedings of the International Conference on Intelligent User Interfaces; IUI '25; ACM: New York, NY, USA, 2025; pp. 54–69. [CrossRef]*

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.