

TAPSPY: Semantic Usage Analysis and Triangulated Risk Assessment of Session Replay Services in Mobile Apps

Jiawei Guo

University at Buffalo, SUNY
Buffalo, NY, United States
jiaweigu@buffalo.edu

Zongze Li

University at Buffalo, SUNY
Buffalo, NY, United States
zongzeli@buffalo.edu

Haipeng Cai*

University at Buffalo, SUNY
Buffalo, NY, United States
haipengc@buffalo.edu

Abstract

Session Replay (SR) is a powerful tool that reconstructs a user’s in-app session from fine-grained UI states and interaction events collected over time. Although SR is legitimately used by app developers, it may introduce significant security and privacy risks by recording sensitive on-screen interactions. While these risks have been studied on the web, their prevalence and nature in the mobile ecosystem remain largely unexplored. In this paper, we present the first large-scale, systematic dissection of SR in Android apps. To enable the study, we developed TAPSPY, a novel auditing framework that synergistically combines program analysis with schema-guided, evidence-grounded large language model reasoning to perform a deep semantic audit of SR usage.

Using TAPSPY, we conducted a measurement study on 47,855 high-impact Android apps. Our findings reveal a systemic failure in the mobile ecosystem’s approach to SR security and privacy. Transparency is the first casualty: the vast majority of apps using SR fail to disclose this data collection in their Data Safety sections, creating a significant gap between their declared and actual privacy practices. We uncover a stark asymmetry in developer effort, where developers meticulously enrich session data with user identifiers and custom events for analytics but overwhelmingly rely on coarse, default masking for privacy. Consequently, risks rarely occur in isolation but instead form dangerous bundles, such as failure in redacting sensitive UI content (i.e., masking), logging sensitive information, linking it to a user’s identity, and failing to disclose any of it. Our work provides the first large-scale empirical study of SR risks of these widespread issues in mobile apps and offers actionable insights for developers, SR vendors, and platform operators to mitigate the risks associated with this powerful technology. In line with responsible disclosure, we reported high-severity risks to the developers of affected apps.

CCS Concepts

• Security and privacy → Software and application security.

Keywords

Session replay, mobile apps, security and privacy risks, session replay usage specification (SRUS), triangulated auditing

*Haipeng Cai is the corresponding author.



This work is licensed under a Creative Commons Attribution International 4.0 License.

CCS ’26, The Hague, Netherlands

© 2026 Copyright held by the owner/author(s).

ACM ISBN 979-8-4007-2871-6/2026/11

<https://doi.org/10.1145/3830454.3832570>

ACM Reference Format:

Jiawei Guo, Zongze Li, and Haipeng Cai. 2026. TAPSPY: Semantic Usage Analysis and Triangulated Risk Assessment of Session Replay Services in Mobile Apps. In *Proceedings of the 2026 ACM SIGSAC Conference on Computer and Communications Security (CCS ’26)*, November 15–19, 2026, The Hague, Netherlands. ACM, New York, NY, USA, 15 pages. <https://doi.org/10.1145/3830454.3832570>

1 Introduction

Session Replay (SR) is a powerful technology that helps app developers reconstruct how users interact with an app by recording fine-grained UI states and interaction events over time in modern application (app) development, providing legitimate merits in debugging and user experience optimization. Notably, SR is *fundamentally different* from traditional mobile analytics and logging services. First, while conventional analytics collect coarse-grained telemetry (e.g., predefined events), SR services reconstruct entire user sessions with rich context, capturing detailed UI state and user inputs as they occur. Second, SR systems expose developer-configurable redaction semantics (e.g., redacting UI content by masking controls at global or element granularity), where misconfigurations can silently disable default protections. Third, the privacy impact is further amplified when replay data is associated with persistent identifiers (e.g., user accounts or profile attributes), enabling recorded interactions to be attributed to specific individuals. Fourth, SR is frequently framed as “analytics” or “debugging”, despite recording content and interactions that users typically perceive as private. This *uniqueness* of SR also makes it *riskier than general third-party services*.

In fact, risks induced by flawed SR service usage have already been found in high-profile real-world apps [1, 2]. In 2019, for example, several popular iOS apps were found to be recording entire user screens through SR services without proper disclosure [47]. Some of these apps employed the Glassbox SR service [6], resulting in sensitive data (e.g., payment information and private messages) being inadvertently captured and transmitted. Following media reports, Apple intervened and required developers to remove the offending code or face App Store removal, citing violations of policies mandating explicit user consent and visible indicators for recording [46]. In some cases, these SR-induced risks even triggered legal action [9].

However, extant relevant research has mainly focused on, hence being limited to, *SR usage in web platforms* [21, 22, 39, 52]. Apart from the platform coverage difference, the dominant methodology is runtime- or traffic-centric rather than code-based: prior web studies largely infer SR behavior from dynamic observations (e.g., monitoring script actions and outbound requests, detecting exfiltration, or instrumenting form interactions). This is because the end-to-end SR usage logic is encoded in large, minified first-party bundles

and opaque third-party scripts, and may further depend on server-side decisions that are not directly auditable from the client code alone [21, 39]. As a result, these studies typically do not (and often cannot) systematically examine the first-party integration logic that governs when SR is enabled, what data is unmasked or enriched, and under what guards these behaviors are triggered—precisely the logic that determines SR’s effective security and privacy impact.

This leaves a *significant gap* in understanding the semantic behaviors and associated specific security and privacy risks of *SR usage in mobile apps*—mobile SR studies to date focus on limited features of SR in only a specific category of apps. This gap is particularly concerning given the intimate nature of mobile device usage and the potential for capturing sensitive user interactions that may not occur in web browsing contexts.

As it stands, there is no at-scale, systematized study that (i) *detects SR usage* in mobile apps, (ii) *analyzes SR semantics* governing data-capture and masking at code level, and (iii) *assesses real-world risks* (including, but not limited to, leakage and transparency inconsistencies) induced by SR when the usage is flawed.

In this paper, we aim to bridge this critical gap, presenting the first large-scale, systematic dissection of SR services in mobile apps to characterize their usage and assess their risks. Given the scale of the mobile ecosystem, manual analysis is clearly infeasible, necessitating an automated approach. However, achieving such automation is non-trivial and faces several fundamental challenges.

First (*Challenge 1*), there is a large semantic gap between the low-level code patterns of SR usage and the high-level, abstract security and privacy principles and policies. An effective, automated analysis must be able to infer the functional intent of code (e.g., “this code unmask a login UI”) and connect it to a reference standard (e.g., “do not record credentials”). Moreover, in line with their uniqueness, SR behaviors, which may have security and privacy implications or consequences, are composed of (a) the relevant set of SR APIs the app invokes, (b) the control-flow structure and data-flow context of (i.e., program dependencies among) those API calls, and (c) constraints or conditions (e.g., lifecycle callbacks, conditional branches, global configurations) that govern when and where those APIs are invoked. These SR behaviors, referred to as *compositional SR semantics*, model and determine what SR services actually record and how they annotate or link session recordings. Understanding these semantics from low-level code, which is not a property of any single SR call in isolation, nor point-to-point reachability [16] as in source-sink or taint flow [23], is inherently difficult.

Second (*Challenge 2*), the mobile software ecosystem is marked by immense heterogeneity: e.g., developers use SR APIs in varied and unpredictable ways, and the security and privacy documentation of both apps and SR services are often irregular, incomplete, or ambiguous. This wide variety makes it difficult to apply a single, static set of rules for risk assessment against various SR service vendors and apps. Then, the compositional nature of SR semantics exacerbates this heterogeneity-induced challenge. Different SR service vendors expose distinct API surfaces and defaults (e.g., masking models, consent interfaces, identity primitives), leading to diverse composition patterns in real apps. As a result, SR auditing must both recover composition-level usage context and normalize vendor-specific APIs into comparable semantic concepts.

The advent of large language models (LLMs) seems to bring hope for cracking the above challenges. Yet LLMs have their own inherent limitations (*Challenge 3*), including susceptibility to hallucination and practical constraints on context window size, making them struggle to offer reliable results and prohibitive for analyzing entire apps. Moreover, despite their general potential for code understanding, LLMs may not have contextual reasoning capabilities specific to compositional SR semantics and factual knowledge necessary for SR risk analysis (which involves various information sources such as app disclosures or vendor requirements).

Due to these challenges and uniqueness of SR services and their risks, existing relevant solutions fall short when applied to assess SR risks. To address such challenges, our work is guided by several key insights. First, we leverage LLMs’ reasoning capabilities to interpret the high-level functional semantics of low-level code and link them to natural-language security and privacy concepts defined in our framework, mitigating *Challenge 1*. Second, we introduce a structured approach, using program analysis to first abstract an app’s complex codebase into a concise, code-level representation, referred to as the *Session Replay Usage Specification (SRUS)*. We then develop a global schema for SRUS to normalize LLM’s output as SRUS’s natural-language representation. Meanwhile, we model SR vendors’ policies (e.g., terms of use) and apps’ disclosure/documentation (e.g., data-safety sections) as external references in a principled manner, while applying the same global schema to unify those references but refined specifically to each SR service and individual app. These together overcome *Challenge 2*, while further addressing *Challenge 1*. The code-level abstraction via SRUS dramatically reduces the code analysis scope for the LLMs, mitigating their context window constraints; we further introduce a triangulation-based approach to risk auditing, grounding LLM-based assessment with verifiable evidence based on a multi-faceted reference model, which counters model hallucination to overcome *Challenge 3*.

Based on these insights, we developed TAPSPY, a novel, automated technique that analyzes compositional SR semantics hence assesses security and privacy risks of SR services as used in mobile apps. Using TAPSPY, we then conducted a large-scale study of 47,855 real-world apps to provide the first empirical view of SR-induced risks in the mobile ecosystem. Our evaluation shows that TAPSPY is highly effective, achieving 89.47% precision and 94.44% recall in identifying the risks, at low costs of <3 minutes and <\$0.1 per app.

Our study reveals a systemic failure in the mobile ecosystem’s approach to SR security and privacy. While transparency violations are widespread, we show that SR risks in mobile apps are driven less by the mere presence of recording and more by how developers *configure and enrich* SR sessions in first-party code. In particular, risky deployments frequently combine (i) weak or overridden masking, (ii) identity binding, and (iii) verbose enrichment (custom events or properties), creating replay artifacts that are both information-rich and readily attributable to a specific user.

Among others, our noteworthy findings are summarized below.

- We identify **551** Android apps in our dataset that integrate SR services, spanning 14 unique SR service vendors.
- Among these apps, **312 (56.6%)** exhibit at least one risk, spanning *all* app functionality categories, and *all* risks in our taxonomy

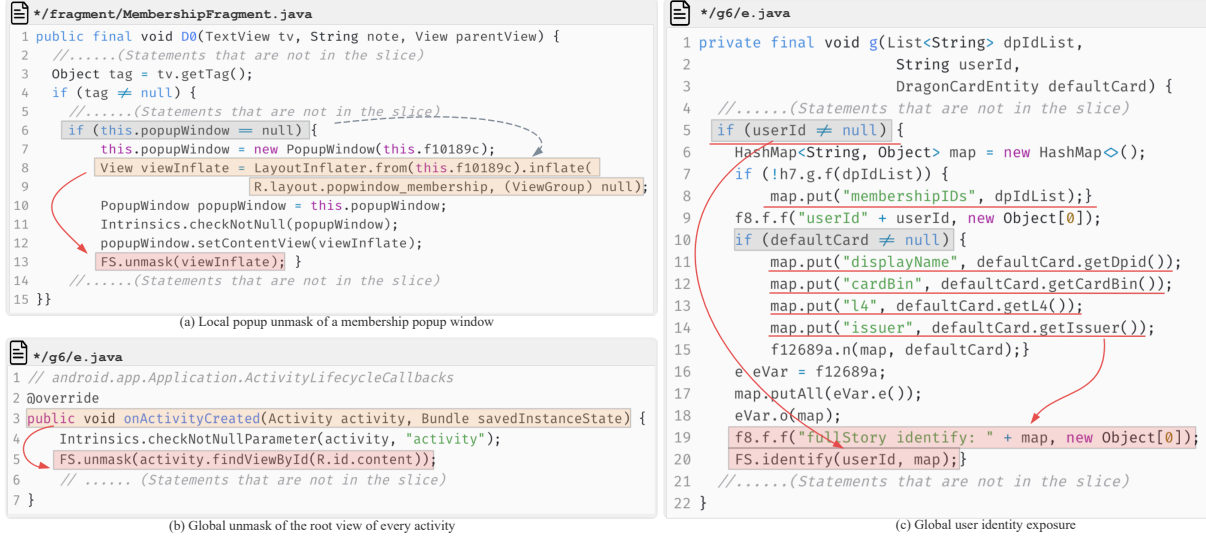


Figure 1: Examples of privacy-risking uses of the SR service by FullStory in a real-world app. (a) A membership popup window is explicitly unmasked, allowing its contents to be recorded. (b) User identity and payment-related attributes (user ID, membership IDs, card BIN, last four digits, issuer) are bundled into a map and sent to FullStory via `FS.identify()`. (c) A global lifecycle hook unmask the root view of every activity, effectively disabling masking protections app-wide and exposing all UI contents.

appear in the wild. Most critically, we observe SR-specific implementation failures that can directly expose sensitive on-screen content and bind it to user identities:

- (1) *(Un)Masking Misconfiguration* includes insufficient masking, overly permissive rules (e.g., wildcard-style unmasking), and even explicit unmasking of sensitive UI;
 - (2) *Logging of Sensitive Info* is largely driven by developers logging PII through enrichment APIs; and
 - (3) *PII Leakage in Identification* spans multiple identifier types, including email, user IDs, phone numbers, and names, enabling straightforward re-identification of replay sessions.
- For a small subset of apps, we instrument SR API entry points at runtime and have validated that the risky semantics detected statically do occur during the apps’ executions. Further captured outbound traffic and observable end-to-end SR upload or transmission workflow show concrete evidence that SR artifacts are exported off-device to SR vendor infrastructure.
 - To evaluate the real-world impact of our findings, we initiated a responsible disclosure process focusing on 64 apps that exhibited both technical vulnerabilities (e.g., masking misconfigurations) and a failure to disclose such practices in their privacy policies. To date, developers of 16 unique applications, spanning the Business, Social, Entertainment, Shopping, and Maps & Navigation categories, have formally confirmed our findings. These developer-validated cases substantiate the four primary risk categories identified by our tool: masking misconfigurations, identification leakage, enrichment logging, and consent bypass.

In summary, our work makes the following contributions:

- A novel auditing framework: we design and implement TAPSPY, a new methodology that synergistically combines program analysis with schema-guided, evidence-grounded LLM reasoning to perform a deep semantic audit of SR usage in mobile apps (§3).

- A large-scale measurement: we present results of the first large-scale study of SR services in Android, quantifying their prevalence and revealing widespread risky implementation patterns and policy inconsistencies in real-world apps (§5).
- Actionable insights: we provide insights for multiple stakeholders, including developers, SR vendors, and platform operators, to help mitigate the security and privacy risks associated with this powerful and popular technology (§6).

2 Background and Motivation

2.1 Session Replay (SR)

Session replay refers to the practice of reconstructing a user’s experience within an app—capturing interactions like taps, scrolls, view changes, and other UI events—to visualize how users navigate and interact with the app. It is widely valued in product development, UX research, and customer support as it provides qualitative context to quantitative analytics, enabling teams to understand why users behave in certain ways, not just what they do.

On mobile platforms, SR services typically do not rely on video-like screen capture. Instead, they operate by capturing snapshots of the app’s view hierarchies—structures of UI elements (e.g., views, text boxes, images, input fields) that represent the app’s state—and by logging interaction events over time. These snapshots, often represented as “frames,” are transmitted (frequently as diffs) to the vendor’s backend, where the UI is reconstructed for playback, recreating a user’s journey with greater efficiency and control. These legitimate uses do not make SR inherently unsafe; rather, the risk depends on how developers configure SR in first-party code, as SR services often expose developer controls over what is captured and how sessions are contextualized. These controls include masking or unmasking UI content at different scopes (e.g., element, screen, or global), as well as enrichment and identification APIs that attach business events or user attributes to a replay.

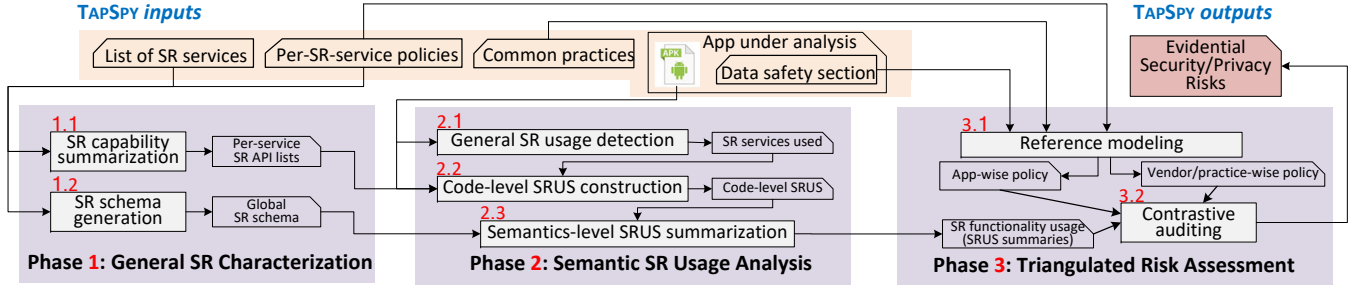


Figure 2: Overview of our TAPSPY framework, including its inputs, three key working phases, and outputs. Starting from the app under analysis together with SR documentation, policies, common practices, and the app’s Data Safety disclosures, TapSpy derives a unified SR schema (Phase 1), recovers app-specific SR usage semantics (Phase 2), and contrasts the recovered behavior against disclosure- and policy-based references to identify evidence-backed security and privacy risks (Phase 3).

2.2 Motivation

Despite its justifiable merits and prevalence, SR comes with security and privacy risks, which are often subtle yet severe. To illustrate, we consider a high-profile popular travel rewards app with over a million downloads and a stellar 4.8-star rating from over 60k users. It integrates a leading SR service, FullStory [8], and in doing so, implements several privacy-risking practices, as shown in Figure 1.

Figure 1(a) depicts a seemingly benign feature: a membership popup window. However, the code explicitly disables the default privacy protections for this UI element. After inflating the popup view (viewInflate at Line 10), the developer passes it to the `FS.unmask()` API (Line 15). This single call instructs the SR service to record the full UI content of the popup, which may include the user’s name, membership status, or other personal details, defeating the SR service’s default redaction.

While unmasking a single element is concerning, Figure 1(b) reveals a far more systemic issue. The app registers a global `ActivityLifecycleCallbacks` (Line 1) that triggers every time a new screen (Activity) is created. Within this callback, the code finds the root view of the new screen (Line 4) and calls `FS.unmask()` on it (Line 5). The consequence is profound: every single screen in the app is unmasked, effectively nullifying the SR service’s privacy safeguards app-wide and exposing all on-screen content to recording.

Since most SR services can enrich the captured sessions with more information beyond capturing the user interface, SR can be used to exfiltrate sensitive data directly. In Figure 1(c), the app gathers user and payment-related attributes, including a user ID, membership IDs, and payment card details like the Card BIN and issuer (Lines 3, 12-15). This data is then bundled into a Map object, and as the data-flow arrows indicate, passed to `FS.identify()` (Line 23). This explicitly sends a detailed profile of the user and their payment instrument to a third-party analytics service. This kind of instrumentation may be intended to support richer analytics or replay segmentation. Notably, even if some values are visually masked in the UI, the risk here remains because the data is passed directly to the SR service through an identification API rather than captured from the rendered screen. Such behaviors are particularly concerning given that the app’s Data Safety section claims “No data shared with third parties” and does not declare any data as being collected for “Analytics” purposes, putting its implementation in contradiction with its public disclosures.

The effective semantics in this example is the *composition* of three SR behaviors, each enabled by a specific API call under a

specific guard or context to conduct (i) local unmasking, (ii) global unmasking, and (iii) identity or profile linkage. Together, these calls yield a composed semantics where the SR (a) records UI content that would otherwise be redacted (now including all screens due to the lifecycle-guarded global unmask), and (b) links those recordings to an identified user enriched with sensitive attributes—amplifying privacy risk beyond what any single call would suggest on its own.

Manually auditing thousands of apps for these nuanced issues is intractable. This motivates the need for an automated framework capable of operating at the scale of the app ecosystem. Next, we present our technical design to address these very problems.

3 TAPSPY Design

We start with an overview, followed by details of each component.

3.1 Overview

Figure 2 gives a high-level overview of TAPSPY highlighting its overall design. In addition to (1) the Android app (APK) under analysis, **TAPSPY Inputs** also include (2) a *list of SR services* (names) targeted and their official API documentation and security and privacy policies (*per-SR-service policies*); (3) common security, privacy, and compliance requirements (*common practices*)—individual SR service vendors do not always include these common policies (e.g., “session recordings should not lead to PII leaks”); and (4) data collection and sharing disclosure of the given app (*data safety section*)—the Google Play Data Safety section of the app’s listing.

With these inputs, in **Phase 1**, TAPSPY builds a general understanding of SR capabilities from the target SR services’ documentation, yielding a comprehensive *list of SR APIs* for each service. It then summarizes all security and privacy aspects relevant to SR capabilities, resulting in a *global schema* that defines a unified vocabulary of SR features and their associated privacy risks. The goal of this phase is to lay a *unified common foundation* for the consistent semantic analysis in **Phase 2 across heterogeneous SR services and their diverse usages in various apps**.

Taking the global SR schema and per-service API lists, in **Phase 2**, TAPSPY detects and characterizes actual SR usage in the given app. It first detects the presence of any SR libraries in this app. If the app package (APK) does not even include any SR library, TAPSPY may end here. Otherwise, it proceeds by constructing the Session Replay Usage Specification (SRUS) of the app, a code-level abstract representation that isolates the minimal, relevant program

logic as a subgraph of each API caller’s program dependence graph (PDG) [18], which includes each API callsite and its intraprocedural data- and control-flow context. The rationale is that such contextualized SR API calls collectively capture how the app uses SR functionalities of each included library. To enable high-level (semantic) understanding of SR usage, the SRUS is then converted to its natural-language form by an LLM via graph-based reasoning guided and output-regulated by the global schema. The goal is to allow for reasoning about risks induced by the SR usage at the same (i.e., natural-language) level (hence more accurately) in **Phase 3**.

Finally, in **Phase 3**, TapSPy audits the semantic SRUS summary for security and privacy risks and policy violations through a triangulation-based cross-referencing strategy. First, it constructs a reference model based on two data sources: the app-wise policy of the app (i.e., its Data Safety declarations) and the SR service vendor- and practice-wise policy (i.e., the vendor’s terms of use or service and common security and privacy practices). Then, relevant violations are detected by triangulating the two data sources together with the app’s SR functionality usage (i.e., natural-language/semantic SRUS summary)—any inconsistencies found by cross-referencing among the three entities signal risks and violations, with the contrasted references serving as evidence. Such audit reports of evidential risks constitute the main **TapSPy Outputs**.

3.2 General SR Characterization (Phase 1)

To build a foundational, vendor-agnostic understanding of SR capabilities, Phase 1 works in two steps below.

SR capability summarization (1.1). To understand the capabilities of SR, we construct a comprehensive corpus of public-facing APIs for each SR service. This is a prerequisite for understanding how app developers integrate and utilize these services. The primary challenge is that most SR services are commercial, closed-source products and often lack exhaustive, machine-readable documentation like Javadoc. Instead, their functionality is typically presented through tutorial-style developer guides.

We begin by manually curating an initial set of APIs from each SR service’s official developer documentation. This is because these documents highlight the APIs that vendors intend for developers to use, effectively isolating the public API surface from internal implementation details. However, this documentation is often incomplete from a privacy auditing perspective. For instance, a tutorial might demonstrate how to tag a user session by showcasing a single API call (e.g., `setUserId()`), while omitting other semantically related methods within the same class (e.g., `setEmail()`, `setAddress()`, or `setName()`) that could also handle sensitive information such as PII.

To address this incompleteness, we employ a class-based expansion technique. For each API method curated from the documentation, we first identify its parent class. We then programmatically inspect this class to extract all other public methods it declares. This process systematically expands our initial seed set to include functionally adjacent APIs that are exposed to the developer but may not be explicitly mentioned in the tutorials.

SR schema generation (1.2). Next, to enable a consistent and scalable analysis across different apps and SR libraries, we abstract the discovered APIs and features into a global schema. This schema establishes a unified vocabulary of SR functionalities and their associated risks, serving as the canonical knowledge base for our entire

auditing framework. To construct it, we first parse the developer documentation for each SR service and extract all distinct features. After manually reviewing these results for accuracy, we unionize SR features across all services to form a comprehensive catalog. This set was then filtered to retain only those features observable through static code analysis. Finally, we enriched the schema by manually analyzing and annotating the security and privacy risks of each feature, such as its capacity to handle PII or capture user input. The resulting global schema maps code-observable features to their potential risks, providing a robust foundation for our automated audit, as outlined in Table 1.

3.3 Semantic Analysis of SR Usage (Phase 2)

Using the SR API lists from Phase 1, TapSPy then aims to analyze if and how the given app interacts with any of the targeted SR services. This is achieved by the three steps described below.

General SR usage detection (2.1). We first detect if the given app bundles a library with SR service. Thus, we extract all class definitions from the app APK and match against service-unique package namespaces. A positive match indicates the presence of an SR service usage in the app. Apparently, this process requires meaningful semantic information (e.g., method names, class structures) in the app. If the app is not obfuscated, the detection of explicit SR presence is efficient and accurate; otherwise, the information may be lost. Thus, we take a two-pronged approach, first checking explicit presence, followed by detecting implicit (obfuscated) usage and merging the results. The rationale is that an app may both explicitly and implicitly use SR services.

For the implicit presence checking, we leverage state-of-the-art (SOTA) obfuscation-resilient Android third party library (TPL) detection [48], obtaining a mapping between meaningful identifiers (e.g., class and method names) and their obfuscated counterparts. Apps may also dynamically load an SR service, followed by accessing SR APIs through reflection. Thus, we proceed with best-effort handling dynamic and reflective code loading by scanning uses of class-loading interfaces and, where feasible, resolving their target class names or bundled code paths; when these resolve to known SR namespaces or attributable code artifacts, we mark SR presence and carry the corresponding SR API callsites to the next step.

Code-level SRUS construction (2.2). With the curated API corpus, this step aims to characterize how these APIs are invoked. To that end, we introduce SRUS. For each user-defined method M that calls SR APIs, we formally define its SRUS as a subgraph of M ’s Program Dependence Graph (PDG). The rationale is that SRUS can capture the complete and minimal program logic that influences the app’s interaction with the SR service.

The construction begins by identifying all user-defined methods that contain callsites to our target APIs. We iterate through every method in the app packages and match invoke instructions to our API corpus. When identifiers are obfuscated (e.g., renamed), we use the TPL detector’s (same one as used in Step 2.1) internal mapping records to recover canonical class and method names and re-match. For reflective and dynamic invocations, we leverage the dynamically loaded SR library names resolved in Step 2.1 and further utilize a SOTA reflection resolver to identify the reflective call targets.

To focus exclusively on the app developer’s integration choices, we exclude calls originating from the SR library’s own packages or

the standard Android framework. Using method D0 in Figure 1(a) for example, once we found there is a library API call (Line 15) and the containing class name does not start with the SR library's package name, i.e., `com.fullstory` in this example, we label this method as a potential user-defined method. It is worth noting that we observed that for some general analytics service vendors, SR is an optional, distinct feature rather than the default behavior. For instance, Sentry [3], a widely used crash reporting and performance monitoring library, just recently introduced a stable SR feature. An app might integrate Sentry for its core features without ever enabling SR. To handle these cases, we manually identified the specific APIs required to **enable** SR functionality for each such SR service vendor. During our analysis, we only consider an app to be using SR if it explicitly invokes one of these **enablement** APIs. This ensures our study accurately targets apps that have actively opted into SR, not just including a library that offers it as a feature.

For each such user-defined method M , the SRUS subgraph is constructed by first performing a two-level union of backward slices to identify all relevant statements, and then *projecting* the method's full PDG to the union slice. Formally, the SRUS is constructed as:

$$SRUS(M) = PDG(M) \left[\bigcup_{c \in C_M} \bigcup_{a \in A_c} \text{Slice}(M, c, a) \right]$$

Where:

- $PDG(M)$ is the program dependence graph of method M .
- C_M is the set of all SR API callsites within M .
- A_c is the set of arguments for a given SR API callsite c .
- $\text{Slice}(M, c, a)$ is the static backward slice tracing the data and control dependencies that influence the value of an argument a at a callsite c .
- $G[V_{sub}]$ is an operator that yields the subgraph of a graph G that is induced by the vertex subset V_{sub} .

This definition ensures that the resulting SRUS is a concrete subgraph containing only the nodes (statements) and edges (dependencies) relevant to SR API interactions. Moreover, this code-level representation of SRUS corresponds to our notion of compositional SR semantics—as the PDG projection captures control flow structure, data flow dependencies, and constraints and conditions which together compose the SR semantics. To prepare the SRUS for downstream semantic analysis, we enrich this subgraph by numbering each statement and explicitly annotating the dependencies, for example, as $s_i \rightarrow s_j$. The resulting collection of these annotated PDG subgraphs constitutes the final SRUS for an app.

To illustrate, still consider Figure 1(a), where the API call `FS.unmask(viewInflate)` at line 15 is identified as the seed in the previous phase. The backward slice from this seed collects all statements that define or propagate its operands or guard its execution. Concretely, the slice retains the predicate input Line 5 (`tag = tv.getTag()`), the enclosing guards Line 6 (`if (tag != null)`) and Line 8 (`if (this.popupWindow == null)`), the creation of the inflated view Line 10–11 and its propagation into the popup Line 14, the receiver chain through Line 9, 12, and 13, and the seed itself 15. Projecting the method's PDG onto these nodes yields a compact subgraph whose edges include control dependencies (6→15, 8→15), data dependencies along the argument chain (10→14→15), data dependencies along the receiver chain (9→12→13→14), and the flow into the control predicate (5→6). This SRUS thus captures what object is

unmasked, how it flows to the API, which receiver instance is involved, and under which conditions the action occurs. If there are more seeds for this method, we repeat the same slicing procedure for each additional SR-API call and merge the resulting nodes and edges; their union constitutes the complete SRUS for the method.

It is important to note that this slicing and projection is performed on recompiled, Java-like source code rather than on a lower-level intermediate representation like *Jimple*. While *Jimple* is well-suited for traditional data-flow analysis, its three-address code format often introduces semantic redundancy by breaking down single high-level expressions into numerous simpler statements. This verbosity increases the analysis burden for both human auditors and the LLM. By operating on a more concise, source-level representation, we significantly reduce the number of input tokens required for the LLM, making our large-scale semantic audit more efficient and cost-effective.

Semantics-level SRUS summarization (2.3). With an app's SRUS constructed, the next step is to automatically generate its natural-language summary (i.e., natural-language representation of SRUS). This summary aims to answer what SR features an app uses and how it implements them. To achieve this at scale, we employ an LLM-based process guided by the schema defined previously.

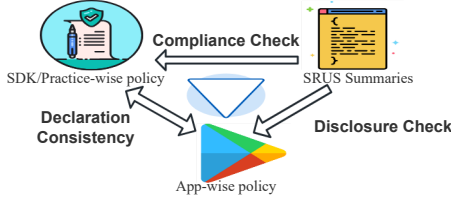
The process begins with schema localization. With the global schema from the preceding step, we create a localized, service-specific schema by filtering for only those features relevant to the specific SR service present in the app. This crucial step focuses the LLM's analysis exclusively on the features that are actually available in the service being used. For example, in Figure 1(a), the app used *FullStory*, so the vendor-specific schema retained only the features supported by *FullStory*.

To generate the summary, we prompt the LLM to perform graph-based reasoning [26] on the SRUS. Instead of providing only the raw source code, we provide the entire SRUS structure: the numbered statements and the explicit data and control dependency edges (e.g., $s_i \rightarrow s_j$). The LLM is specifically instructed to utilize this graph to trace how data flows into API call arguments and to understand the conditions under which these APIs are invoked. This leverages the rich dependency information from the PDG to ground the LLM's interpretation in the underlying program logic, leading to more faithful and accurate summaries.

For each feature in the localized schema, the LLM generates a structured output containing: (1) Usage: a boolean determination of whether the app utilizes the feature; (2) Description: a natural language summary, derived from its graph-based reasoning, of how the feature is implemented; and (3) Evidence: the concrete arguments and key statement excerpts from the SRUS that support its conclusion. Using the same example in Figure 1, for Figures 1(a) and (b), the LLM will detect the feature presence of *Element Level* and *Screen Level* redaction control and output *true* for two features, with a natural language description—for example: "*The app invokes FS.unmask() in two distinct ways: (i) it conditionally un.masks a membership popup view when a specific UI or feature condition holds; and (ii) it registers ActivityLifecycleCallbacks and, during lifecycle events, repeatedly accesses each activity's root view and applies FS.unmask(), effectively performing systematic unmasking across screens.*" and the corresponding SRUS statement excerpts as evidence. The LLM prompt template used is provided in our [artifact](#).

Table 1: Risk Taxonomy for SR Practices

Risk Category	Risk	Description
Disclosure Transparency (RC1)	Omitted Data Practices (R1)	The app's code behavior shows it collects or shares a specific type of data (e.g., location) via SR, but this practice is never mentioned in its Data Safety section.
Behavioral Implementation (RC2)	(Un)Masking Misconfiguration (R2)	The developer fails to mask enough sensitive information in a UI page or deliberately un.masks sensitive information.
	Logging of Sensitive Info (R3)	The developer misuses a Data Enrichment API (e.g., logEvent, setCustomEvent, or log) to pass sensitive, unencrypted information.
	PII Leakage in Identification (R4)	The app writes PII, credentials, or sensitive information when identifying the user's identity associated with the session.
	Erroneous Privacy Configuration (R5)	The developer calls the privacy APIs, but in the wrong order.
Consent & User Control (RC3)	Disregard User Choice (R6)	The user has explicitly denied consent, but the app or SR service proceeds to collect or share data anyway.
	Hard-Coded Consent (R7)	The developer hard-codes a "true" or "opt-in" value when calling an SR's privacy API, completely ignoring any preference the user might have selected in the UI.

**Figure 3: Reference modeling and triangulated risk auditing.**

3.4 Triangulated Risk Assessment (Phase 3)

Leveraging the natural-language summaries from the preceding phase, the final step in our methodology is to audit these behaviors for potential privacy risks and policy violations. Our approach is a triangulated audit (Figure 3) that positions the app's observed behavior, the app's own disclosures, vendor and common privacy mandates [31, 34, 35, 43, 53, 57] as three vertices of a reference triangle. By examining the relationships across each edge, TAPSPY systematically checks for mismatches and omissions. The final output of this process is a detailed privacy audit report for each app, identifying discrepancies between its actual implementation of SR and the policies it is expected to follow.

Reference modeling (3.1). To ground the LLM's reasoning, we build a triangulation model from three distinct sources. The rationale is two-fold. First, by providing the LLM with explicit reference policies, we mitigate the risk of hallucination and ensure its judgments are based on documented ground truths. Second, it allows us to account for developer transparency; as we detail in our threat model, a potentially risky data practice may not constitute a violation if it is clearly disclosed to users. These sources, shown as three vertices of Figure 3, represent distinct kinds of evidence:

- **SRUS Summaries (SR functionality usage):** The semantic summary of how the app actually uses SR features, as derived from the SRUS constructed in Phase 2.
- **App-wise policy:** The developer's self-declared data practices from the Google Play Data Safety section, specifically data claimed as collected or shared for "Analytics" purposes.
- **Vendor/practice-wise policy:** A unified model synthesized from SR service vendor terms of service and widely acknowledged privacy principles (e.g., GDPR, CCPA), capturing contractual requirements and common best practices.

Together these three sources form our reference model against which contrastive checks are performed. This triangulated view

mitigates LLM hallucination by grounding analysis in verifiable inputs, and it accounts for both developer-facing obligations (vendor mandates) and user-facing transparency (Data Safety declarations).

Contrastive auditing (3.2). We leverage LLMs for auditing, which takes the above three kinds of evidence as its inputs.

Given these inputs, the LLM performs a structured contrastive analysis. Rather than assessing the code in isolation, it evaluates whether the observed behaviors are consistent with the expectations instantiated in the reference model. In practice, this involves checking each edge of Figure 3:

- **Disclosure Check (SRUS ↔ App-wise policy):** Does the app's SR usage involve data not disclosed in its Data Safety section?
- **Compliance Check (SRUS ↔ Vendor/practice-wise policy):** Does the app's SR usage violate vendor-imposed rules or widely acknowledged privacy practices?
- **Declaration Consistency (App-wise policy ↔ vendor/practice-wise policy):** Shown in the model for completeness, this edge represents the conceptual relationship between what developers disclose and what vendors or privacy principles would expect to be disclosed. TAPSPY currently does not directly audit this edge, but it highlights potential for systematic under-reporting.

The output of contrastive auditing is a detailed audit report for the app. Each report enumerates the specific risks identified, including the violated policy dimension, a natural-language justification of why the behavior is risky, and the supporting SRUS evidence (API calls and arguments) that triggered the finding. This ensures that results are not only reproducible but also interpretable, grounding high-level assessment in concrete code-level facts. The LLM prompt template used is given in our [artifact](#). Moreover, to enrich the audit report, we developed a taxonomy of SR risks as per the global SR schema. We first synthesize privacy guidelines from official vendor documentation, established issues from prior literature, and standard security best practices, before instantiating the taxonomy for specific SR services; we then use this taxonomy to associate each risk item in the report with an SR risk category.

Table 1 presents the *risk taxonomy* for session replay (SR) practices, which systematically organizes privacy and security concerns into three Risk Categories (RC1–RC3) and seven distinct risks (R1–R7). Each risk highlights a concrete way in which SR implementations can misalign with transparency expectations, data protection principles, or user consent requirements.

- **Disclosure Transparency (RC1)** focuses on mismatches between declared and actual data practices. For example, an app

may use SR service to capture sensitive data (e.g., location) without ever disclosing such practices in its Data Safety section (R1).

- **Behavioral Implementation Risks (RC2)** capture problematic developer practices when configuring or invoking SR APIs. These include insufficient or deliberately inverted masking (R2), logging of sensitive information through enrichment APIs (R3), leakage of personally identifiable information (PII) during session identification (R4), and misordered or incorrect invocation of privacy controls (R5).
- **Consent and User Control Risks (RC3)** address violations of user expectations around consent. Apps may disregard explicit denial of consent (R6), or worse, hard-code consent to always "true" regardless of user input (R7).

This taxonomy provides a systematic lens for identifying how SR practices can create privacy and compliance risks, complementing our proposed global SR schema (§3.2).

To illustrate, we revisit the example in Figure 1, following the output from Step 2.2. Using the NL summaries that indicate unmasking of (i) a membership popup and (ii) activity root views via lifecycle callbacks, TAPSPY matches these behaviors to the reference model's notion of redaction control and default masking. Because both summaries describe actions that disable or override redaction (albeit at different scopes), the audit flags them as *(Un)Masking Misconfiguration (R2)* with scope-specific evidence: the report distinguishes the localized popup unmask in (a) from the systematic, lifecycle-driven root-view unmask in (b), while grounding both in the SRUS slices that expose the corresponding targets and guards.

These examples demonstrate the granularity of TAPSPY's outputs: each risk is labeled with its taxonomy category, grounded in evidence from the SRUS (natural-language representation), and accompanied by a natural-language explanation of the privacy implication. This allows the results to serve both as a quantitative dataset (in our large-scale study) and as actionable evidence for developers and auditors.

Technique generalizability. The core analysis engine (SRUS graph extraction and guided-LLM reasoning via graph-based prompting) is decoupled from the pluggable SR-specific API schema. By swapping this schema, our method can be retargeted to perform a similar semantic audit on other mobile app infrastructure (e.g., analytics, ads, privacy-configurable SDKs). Our Phase-3 triangulated audit (code semantics $\leftrightarrow$ vendor guidance $\leftrightarrow$ app disclosure) is also fully general, providing a new way to find deep code-documentation mismatches in any Android app ecosystem.

4 Implementation and Evaluation

4.1 Implementation of TAPSPY

We implemented TAPSPY for Android apps. To support consistent analysis across SR services, we first constructed a global SR feature schema by manually curating SR capabilities from vendor documentation, unionizing them across services, and retaining only features observable in app code. We then manually derived the risk taxonomy by synthesizing privacy guidelines from official vendor documentation, established issues from prior literature, and standard security best practices. For the LLM-involved steps, we used Gemini-2.5-Pro [7].

For static analysis, we use Soot [44] for SR library presence detection and API call-site identification. We use LibScan [48] to

recover meaningful library identifiers. To construct code-level SRUS, we decompile APKs with Jadx [41], resolve reflective calls with DroidRA [42], and use Joern [51] for dependence analysis and backward slicing.

4.2 Evaluation of TAPSPY

In this section, we evaluate the performance of TAPSPY in terms of both effectiveness and efficiency. In addition to the end-to-end violation detection results, we explicitly evaluate the two LLM-involved steps in our pipeline: (i) generating a natural-language (NL) representation of the SR usage from SRUS, and (ii) performing risk assessment grounded by our triangulated reference model.

4.2.1 Setup. Reliability of the global schema and risk taxonomy. To assess the global SR schema and the SR risk taxonomy curated through collective author decisions (i.e., which features and risks belong in the canonical vocabulary), we quantify inter-author agreement on inclusion and exclusion decisions. Concretely, we first form a closed candidate pool by taking the union of items proposed by documentation parsing and manual inspection; then multiple authors independently vote whether each candidate should be included. We compute Fleiss's Kappa on these pre-adjudication votes to measure agreement beyond chance, and resolve disagreements through discussion to produce the final schema (taxonomy).

To assess the effectiveness of TAPSPY, we established a ground truth by manually inspecting a statistically significant subset of our dataset. From a population of 551 unique apps confirmed to use SR services, we randomly selected 84 app-service pairs. This sample size was determined using a 90% confidence level, a 5% margin of error, and an estimated population proportion of 10%.

Ground truth for SRUS NL representation. For each selected pair, annotators inspected the SR-relevant code paths and the SRUS slice, then produced a reference annotation of the key SR semantics that should appear in an accurate NL representation (e.g., whether SR is enabled or disabled, masking/unmasking scope, identity linkage, and activation guards such as consent/lifecycle/region logic). We then evaluated whether the generated NL representation from each approach correctly captured these annotated semantics, reporting precision/recall/F1. We also report the underlying ground-truth reliability using inter-annotator agreement, denoted by Fleiss Kappa (FK) score.

Ground truth for risk assessment. For the same 84 pairs, annotators determined the presence and type(s) of SR-related privacy/security risks by inspecting the corresponding code paths and vendor documentation where needed. This ground truth serves as the basis for measuring precision/recall/F1 of the final risk assessment stage, and we again report FK to quantify labeling reliability.

Baseline. To quantify the performance contribution of SRUS and the reference model (beyond the LLM itself), we implement a standalone LLM baseline that takes as input the same SR-invoked API callsites but without access to SRUS or our triangulated reference model. The baseline uses a fixed prompt to (i) summarize the SR behavior and (ii) judge whether the callsite indicates any privacy/security risk. We compute its precision/recall/F1 against the same ground truth.

To assess efficiency, we measured execution time, peak memory usage, and LLM API costs. We additionally break down runtime across the three primary phases of TAPSPY to identify bottlenecks.

4.2.2 Results. For the two curated knowledge bases, SR feature schema and risk taxonomy, on independent include or exclude votes over the closed candidate pools, we obtain a Kappa of 0.83 for the SR schema and 0.91 for the risk taxonomy. These scores indicate high agreement beyond chance, suggesting that both the schema (what SR functionalities are code-observable and should be included) and the taxonomy (what risk types should be represented) are reliable foundations for the downstream auditing stages.

Table 2 reports results for both LLM-involved subtasks. For **SRUS NL representation**, TapSPY achieves high precision and recall ($F1=0.941$), while the baseline performs substantially worse ($F1=0.436$). This indicates that graph-guided prompting grounded in SRUS is critical for producing faithful, semantically complete descriptions of SR behavior, rather than relying on the LLM to infer missing context from isolated callsites.

For the **final risk assessment**, TapSPY again substantially outperforms the baseline ($F1=0.919$ vs. 0.250). This performance gap supports our central claim that TapSPY’s accuracy comes from combining program-analysis-grounded SRUS with a triangulated reference model, which enables composition-aware reasoning about SR semantics (e.g., masking scope, identity linkage, and activation logic) that is difficult to recover from callsites alone. Finally, the FK values in Table 2 show that the manual ground truth for both subtasks is reliable, reducing the likelihood that the observed performance differences are artifacts of annotation noise.

In Table 3 the average execution time of TapSPY is less than 3 minutes per app with a reasonable memory load. The LLM API cost and token cost are also minimal, which are key advantages of our approach to be scalable in practice.

In Table 4 the most resource-intensive phase is Phase 2, which is responsible for the code-level construction and semantics-level summarization of the SRUS. The high resource usage in Phase 2 is expected, as it involves detailed static analysis to build a comprehensive dependency graph of the app’s codebase. This process is crucial for providing the granular, grounded context needed to prevent LLM hallucination and ensure the accuracy of the final risk detection. Our analysis demonstrates that while the system requires significant computational resources during the SRUS construction phase, the overall cost of the approach is low, making it a viable solution for large-scale security audits.

5 Measurement Study

Using TapSPY, we conducted a large-scale measurement study, aiming to answer four research questions (RQs).

- **RQ1:** What is the prevalence and feature landscape of SR services in Android apps?
- **RQ2:** What are the common feature usage patterns of SR in Android apps?
- **RQ3:** To what extent do real-world SR implementations introduce privacy risks?
- **RQ4:** What do in-depth case studies reveal about real-world privacy risks introduced by SR?

5.1 Methodology

Dataset. To build our initial list of SR service providers, we consulted public software review platforms, starting with G2’s “Best

Session Replay Software” list [4]. For each candidate SR service vendor, we manually verified from its official website whether it supported the Android platform, as many are web-exclusive. This process yielded a set of potential SR services for our study. However, we encountered a practical limitation: several enterprise-focused SR service vendors (e.g., Glassbox [6], Quantum Metric [5]) do not provide public developer documentation or accessible AAR/JAR files. Access to their technical details typically requires direct contact with a sales team. As this prevented us from systematically mapping their API surface, we excluded these services from our analysis. For app collection, we scraped the Google Play Store to collect high-impact apps, defined as those with one million or more downloads. This resulted in a dataset of 47,855 apps.

Metrics. In our large-scale measurement study, we use the following metrics: (1) Prevalence and Usage: We measure the adoption of SR service libraries/enablers of SR support in the library and the usage of specific SR features (as defined in our schema) as a percentage of the apps in our dataset. (2) Co-occurrence Patterns: To understand how features and risks are combined, we analyze their co-occurrence. We use the Jaccard similarity for pairs of items (SR features or risks) and Support (the proportion of apps containing a specific itemset) for identifying common combinations of three or more items. (3) Feature-Risk Association: To identify feature sets that are predictive of privacy risks, we use standard association rule mining metrics: Support, Confidence (the conditional probability of a risk occurring given a feature set), and Lift (the ratio of the observed confidence to the expected confidence).

Procedure. First, to answer RQ1, we began by characterizing the capabilities of 14 prominent SR services against our global feature schema. We then executed Phase 1 of TapSPY on our dataset of 47,855 apps to detect the presence of these services. We aggregated results to determine the market share of each service and analyzed their distribution across different Google Play Store categories.

Next, to address RQ2, we processed the apps identified in the previous step through the complete TapSPY pipeline. The semantic summaries generated in Phase 2 allowed us to identify which of the 16 features each app implemented. We then aggregated this data to calculate overall feature usage rates and performed the co-occurrence analysis to identify common feature bundles. We also categorized the specific arguments passed to APIs of features in FC2 and FC4 to understand developer implementation choices.

For RQ3, we used the final output from TapSPY’s triangulated risk audit (Phase 3). This provided specific privacy risks present in each app. We calculated the prevalence of each risk across the dataset and performed a co-occurrence analysis to find which risks tend to appear together. Finally, we conducted an association analysis to identify the feature combinations most strongly correlated with specific privacy risks. To answer RQ4, we manually selected illustrative examples from the apps analyzed for in-depth analysis.

5.2 General SR Usage (RQ1)

To understand the capabilities and diversity of the SR ecosystem, we first analyze the feature sets of 14 prominent SR services. Figure 4 shows that SR adoption is highly uneven across services and app categories. Each bubble represents a vendor–category pair, where size indicates how many apps in that category use the SR service and color indicates the relative scale of installs. For example, the

Table 2: Effectiveness of TAPSPY on two LLM-involved subtasks: (i) SRUS natural-language (NL) representation generation, and (ii) final risk assessment

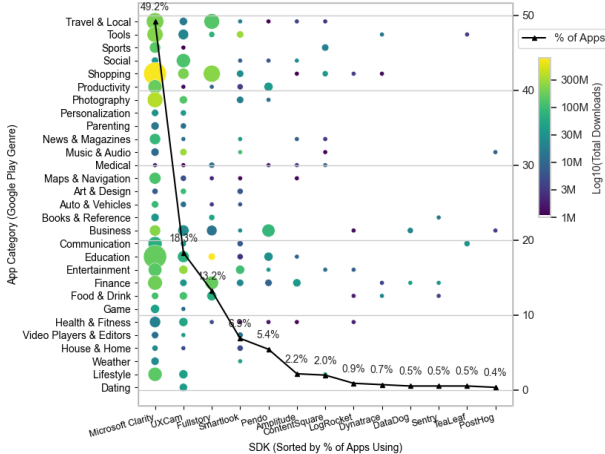
Approach	Graph-guided prompting for SRUS NL representation				Risk assessment grounded by triangulated reference model			
	Ground Truth's Fleiss Kappa	P	R	F1	Ground Truth's Fleiss Kappa	P	R	F1
TAPSPY	0.78	0.952	0.931	0.941	0.92	0.895	0.944	0.919
Baseline (LLM)	–	0.454	0.412	0.436	–	0.273	0.231	0.250

Table 3: Efficiency of TAPSPY

Metric	Min	Max	Mean
Execution Time (s)	41.51	245.39	166.69
Peak Memory (MB)	2987.60	8730.09	3879.28
LLM API Tokens Used	3257	6802	4743
LLM API Cost (\$)	0.03	0.07	0.06

Table 4: Average Execution Time per Phase of TAPSPY

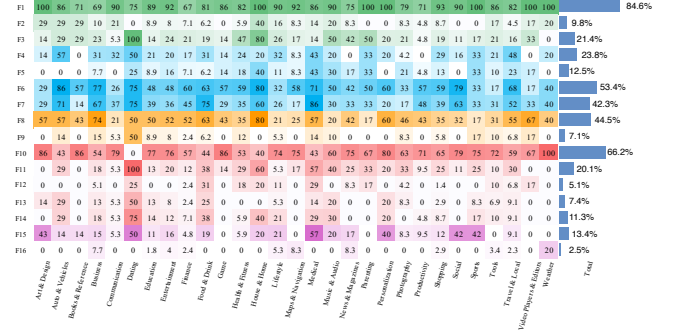
Phase	Sub-Phase	Time (s)
Phase 1	SR presence detection	6.14
Phase 2	Code-level SRUS	97.66
	Semantics-level SRUS	40.25
Phase 3	Contrastive detection	22.64

**Figure 4: Distribution of SR service presence in studied apps.**

Shopping–Clarity bubble is large and bright, meaning that many Shopping apps integrate Clarity and together account for substantial installs. These examples show that SR adoption can produce either broad exposure via many mid-sized apps or concentrated exposure through a few blockbuster apps. Weather or Dating categories show faint or no bubbles, indicating minimal use.

Our results also show that the SR adoption is concentrated in a handful of providers: Microsoft Clarity alone appears in nearly half of all apps, with UXCam and FullStory trailing far behind. Together with the bubble matrix, this shows that SR exposure is shaped both by a few services that dominate the ecosystem and by a few high-impact category–vendor pairings.

The category distribution is similarly skewed. Shopping, Education, and Travel & Local are the most prominent rows, each showing dense activity across multiple services. Finance also stands out: although its row is less dense, its bubbles are consistently large,

**Figure 5: Distribution of SR Feature usage in the studied apps.**

implying that SR-using apps in this domain reach especially broad audiences despite smaller adoption.

This distinction between adoption count and audience reach is critical. The prevalence rate of 1.2% (551/47,855 apps) should be contextualized by the nature of these tools: unlike ubiquitous free analytics services, SR services are often paid, enterprise-grade products. Consequently, their impact is better captured by user exposure than raw app counts. The apps in our dataset using SR average over 15 million downloads each, creating a cumulative exposure footprint exceeding 8 billion downloads. Furthermore, SR adoption spans 29 distinct categories, penetrating sensitive domains such as Finance, Medical, and Business. Thus, while ecosystem-level prevalence is modest, the combination of high-volume user exposure and broad deployment in sensitive domains shows that SR has substantial real-world reach.

Concentration in a handful of SR services, and concentration of popularity in specific categories, imply that misconfigurations or weak defaults in just a few services can propagate widely, particularly in high-stakes app categories such as Finance.

5.3 Semantic SR analysis (RQ2)

In this RQ, we show how SR features are distributed in the apps analyzed. Figure 5 shows that some SR features are nearly universal, while others are rare. Initialization Setup (F1) dominates, reflecting that most apps explicitly configure SR on startup. Default Masking (F10) and User Identification (F6) are also widespread, showing frequent reliance on built-in privacy controls and common linkage of sessions to user profiles.

Across categories, the chart reveals systematic gaps. Some genres never employ particular features—for instance, F10 is absent in Dating. These omissions suggest that developers in some domains place less value on such controls, either because the perceived sensitivity is lower or the workflows are simpler.

Notably, F1 is not present in every app. This does not mean SR is running without initialization; rather, setup may occur in ways

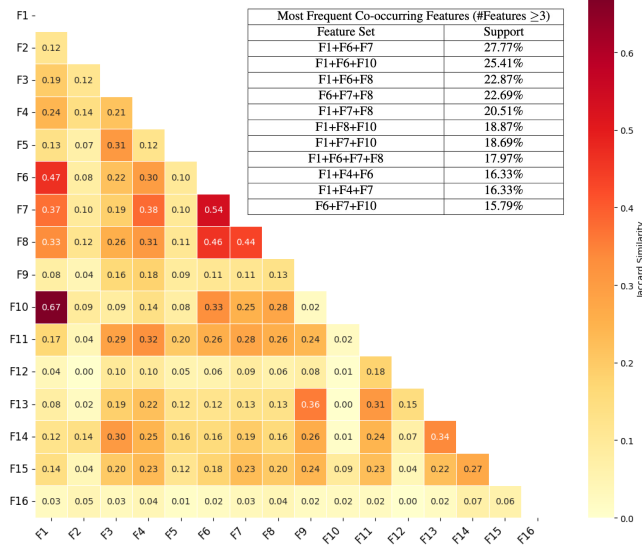


Figure 6: Co-occurrence of SR Features in the studied apps.

not visible to our code analysis, such as through configuration files outside the decompiled package or remote configuration from a vendor dashboard. In such cases, no explicit initialization call appears in the app code even though the service is active.

Developers often use mandatory or default features, while optional controls (e.g., consent and fine-grained masking) remain uncommon.

Figure 6 presents pairwise co-occurrence together with the most frequent ≥ 3 -feature sets. The column for F1 is especially informative: its brightest cells align with F10, F6, Custom Event Tracking (F7), and Session Control (F8), indicating that once SR is initialized, developers commonly enable identification, log business events, and control session boundaries while leaving default masking in place. The table confirms this “business-context bundle”: the top-ranked triad is F1+F6+F7, with support covering over a quarter of apps, showing that identification and event logging are frequently deployed together after initialization rather than in isolation.

The higher-order co-occurrence table clarifies patterns that a pairwise Jaccard view cannot capture. Several top triads include F1 and at least one of F6, F7, F8, demonstrating that value-generating features are activated in bundles. For example, F1+F6+F10 and F1+F6+F8 indicate that, alongside initialization, developers either keep default masking or add session controls; F6+F7+F8 shows that identification, event logging, and lifecycle control often co-occur even without F1 in the triple.

Equally important is what does not cluster. Privacy add-ons that require extra engineering—granular masking at component/screen scope (F12/F13) and explicit consent (F16)—do not appear among the most common triads, despite being potent mitigations. Their absence in the top higher-order patterns, together with their weak pairwise associations in the heatmap, suggests that developers largely rely on defaults (F10) and seldom layer fine-grained protections or consent logic on top of identification and tracking.

SR deployments are built around tightly coupled “analytics bundles” (initialize → identify users → track custom events → manage sessions), while optional privacy controls remain peripheral.



Figure 7: Distribution of SR Risks in the studied apps.

5.4 SR Risk Assessment (RQ3)

In this section, we show how SR risks are distributed in analyzed apps. Out of 551 apps that have been found with SR usages, 312 apps (56.6%) are assessed to have at least one type of risks, covering all app functionality categories. This conditional risk is important for interpreting the significance of our findings. While only 0.65% of all apps in the full dataset (312/47,855) exhibit at least one SR-related risk, among SR-using apps the proportion rises to 56.6%. Combined with the popularity of these apps and their presence in sensitive domains, this indicates that the problem is not merely that SR exists, but that risky SR usage is common once SR is adopted. Additionally, all type of risks in the taxonomy are found in the apps analyzed. Figure 7 shows that Omitted Data Practices (R1) are by far the most prevalent risk, often exceeding half of apps in many categories and reaching near-universal presence in others.

By contrast, Unmasking Misconfiguration (R2) and Logging of Sensitive Info (R3) are more scattered. R2 appears in only a handful of categories such as Business, Finance, or Shopping, suggesting that most apps do not attempt to override masking defaults but some do so in ways that introduce leakage. R3 shows a similarly uneven profile: while absent in most genres, it spikes in categories like Medical and Shopping, where custom event tracking is more used and sensitive details are more likely to be logged.

PII Leakage in Identification (R4) shows modest but widespread presence, with rates in the 5–20% range across categories like Lifestyle and Social. This suggests that while not universal, a meaningful minority of apps explicitly attach identifiers such as emails to SR sessions. Erroneous Privacy Configuration (R5) is rarer still, surfacing only sporadically in Business, Entertainment, and Travel & Local, reflecting configuration mistakes rather than deliberate design. The remaining risks are close to absent, such as Disregarding User Choice (R6) and Hard-Coded Consent (R7).

SR usage risks are prevalent. Transparency failures are systemic and dominate the landscape, while enrichment-driven risks appear more selectively in categories where developers actively track user actions or identifiers. Second, consent-related violations are rare, because they generally do not implement any explicit consent logic at all, leaving little opportunity to misconfigure it.

Figure 8 highlights how privacy risks tend to co-occur within the same app. The heatmap shows that R3 frequently overlaps with both R1 and R4, with Jaccard similarity exceeding 0.1 in each case. This indicates that once developers start enriching sessions with sensitive fields, they are often simultaneously failing to disclose these practices or attaching identifiers, amplifying privacy risk.

The higher-order table further confirms this combination. The most frequent triad is R1+R3+R4—particularly concerning since it both increases exposure and undermines transparency. Other recurring sets, such as R1+R3+R2/R7, show that omitted practices almost always accompany enrichment-driven risks. By contrast, risks like R5 or R6 rarely co-occur with others, appearing only in

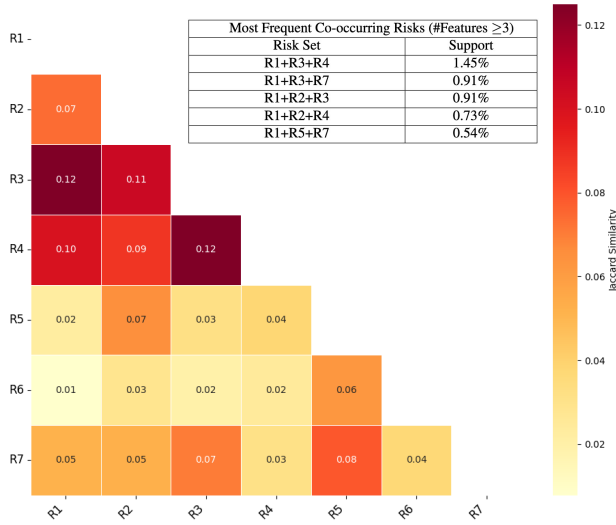


Figure 8: Co-occurrence of SR Risks in the app.

```

1 public final class SplashActivity extends AppCompatActivity {
2   @Override protected void onCreate(Bundle savedInstanceState) {
3     DynatraceSessionReplay.setConfiguration(Dynatrace.applyUserPrivacyOptions(
4       UserPrivacyOptions.builder()
5         .withDataCollectionLevel(DataCollectionLevel.USER_BEHAVIOR)
6         .withCrashReportingOptedIn(true)
7         .withCrashReplayOptedIn(true)
8         .build());
9     Configuration.INSTANCE.builder().withMaskingConfiguration(
10       new MaskingConfiguration.Custom().removeAllMaskedViews().build());
11   ...
12 public final class LoginCheckActivity extends PushNotificationsActivity {
13   public final void redirect(String cpfLogado) {
14     Application application = getApplication();
15     Usuario user = new LoginDao(application).getUser();
16     Dynatrace.identifyUser(user.getCpf());
17   ...
18 public void onclick(Selecionaveiculo CompradorAdapter.ServiceModel obj) {
19   DTXAction dTXActionEnterAction = Dynatrace.enterAction("CRV");
20   String numeroCrvVeiculo = obj.getVeiculo().getNumCrvVeiculo();
21   dTXActionEnterAction.reportValue("CRV", numeroCrvVeiculo);

```

Figure 9: Severe privacy risks via SR in a government app.

weaker combinations. This reflects that such misconfigurations are isolated mistakes rather than systematic patterns.

The SR ecosystem is not just with random, independent risks. The most damaging patterns arise from bundles of enrichment and identification features that systematically co-occur with nondisclosure.

5.5 Case Studies (RQ4)

This section presents in-depth analyses of a small set of real-world apps identified by TAPSPY. We organize the case studies into two parts. First, we show how our static analysis surfaces SR usage and risk-relevant implementation details from app code. Second, we validate a subset of these findings via dynamic instrumentation and network observation, showing that SR APIs are invoked with concrete runtime values and that SR session artifacts are subsequently exported off-device to SR vendor infrastructure.

5.5.1 Static Code Demonstration. We first present static, code-level evidence for representative apps to illustrate the spectrum of SR implementation practices. For each app, we trace the integration of SR service and highlight the specific SR APIs invoked with the surrounding application logic.

```

1 public class ProfileFragment extends BaseFragment implements f57,an3,nr6 {
2   @Override public void f1(String str) {
3     if (TextUtils.isEmpty(str)) {
4       Q5(this.C.J, getString(R.string.add_birthday),
5         R.color.tlp_banner_text, new View.OnClickListener() { ... });
6       FS.unmask(this.C.J);
7     } else {
8       FS.mask(this.C.J);
9       this.C.J.setText(str);

```

Figure 10: A privacy-preserving SR usage.

```

1 public final void startSessionIfEnabledMain (Context context) {
2   ...
3   if (ProfileUtil.INSTANCE.isCountryUS())
4     66 ProfileUtil.INSTANCE.isPhysician()) {
5     Contentsquare.setDefaultMasking(false);
6     Contentsquare.unMask((Class<?>) EditText.class);
7     Contentsquare.start(context);
8     Contentsquare.optIn(context);
9     Contentsquare.send("Launch screen");
10    Settings.singleton(context).saveSetting(
11      com.medscape.android.Constants.IS_CONTENTSQUARE_SDK_ENABLED, true);

```

Figure 11: Targeted, high-risk data collection logic via SR.

A Cascade of Risks in a Government App. The first case is an official digital wallet from a major Brazilian city, with 10M+ downloads and used to manage licenses and registrations. Given this high-trust context, the privacy failures are striking. TAPSPY uncovered a chain of high-severity risks (Figure 9). The app hard-codes consent (Line 7), forcing session recording without user choice (R7). It then calls `removeAllMaskedViews()`, a global unmask that disables all redaction. Next, it leaks national identifiers by passing a citizen’s CPF number (the Brazilian equivalent of a Social Security Number) to `identifyUser()` (Line 16), leak PII via identification (R4). Finally, it logs sensitive data by reporting each user’s Vehicle Registration Certificate number (`numeroCrvVeiculo`, Line 21) as a custom event.

Exemplary Privacy-Aware Implementation. In contrast, we also observed privacy-preserving SR use. The Taco Bell app [10] illustrates how SR can support UX insights while safeguarding data (Figure 10). In its `ProfileFragment`, the birthday field is conditionally handled. When empty (Line 3), the app calls `FS.unmask()` (Line 6), allowing capture of the placeholder state for usability analysis. When populated, it first applies `FS.mask()` (Line 8) before inserting the user’s birthdate (Line 9), ensuring PII is never recorded. Unlike the global unmasking in the government wallet, this design reflects “privacy by design”: granular, context-aware masking embedded directly into app logic. It demonstrates that SR’s analytic value and user privacy can be reconciled, offering a practical model for responsible development.

Targeted Unmasking in a Medical Reference App. A popular medical reference app for millions of healthcare professionals demonstrates a deliberate, conditional SR strategy (Figure 11). SR activates only if the user is a US-based physician (Lines 3–4), suggesting a calculated attempt to apply different standards across regions. For this group, the app disables default masking via `setDefaultMasking(false)` (Line 5), then calls `unMask(EditText.class)` (Line 6), exposing all text input fields. It also hard-codes consent with `optIn(context)` (Line 8), preventing refusal.

This illustrates a more complex privacy challenge. It does not stem from developer negligence, but a deliberate strategy to monitor a specific user base. The conditional logic, while potentially designed to navigate legal compliance, results in a two-tiered privacy system where one group of users is subjected to intense monitoring

Table 5: App cases exercised with verified risks found in SRUS

Cases	Run-time Stack Trace	Risk Description	Risk
1	... → ...fragments.onViewCreated → com.fullstory.FS.addClass(getView() , FS.UNMASK_CLASS);	Unmasks bottom-sheet view, exposing on-screen text/inputs; undisclosed.	R1,R2
2	... → ...MainViewModel\$1\$1.emit → ...com.uxcam.UXCam.setUserProperty(user_email, dummy@email.com)	Links user email via setUserProperty (identity binding); undisclosed.	R1,R4
3	... → ...main.more.MoreFragment.initView → ...main.more.g.onClick → ... → ...UX-Cam.logEvent("Signout", map)	Logs "Signout" with user-profile PII key-value map, enriching the session; undisclosed.	R1,R3
4	... → ...application.App.onCreate → ...com.fullstory.FS.consent(true)	Hard-codes consent(true) at startup; undisclosed.	R1,R7

without their explicit consent. This example highlights how SR can be used for targeted surveillance and raises ethical questions about the practice of segmenting users for different levels of privacy.

5.5.2 Dynamic Behavior Validation. To strengthen the evidence of SR risks identified by TAPSPY, we perform dynamic validation on a small subset of apps by executing them and instrumenting SR API entry points at runtime. Using Frida-based hooks [37], we record the arguments supplied to SR APIs revealed by TAPSPY when they are invoked during interaction, and we compare these observed values against the risk-relevant arguments reported by TAPSPY. We also capture the app’s outbound network traffic to confirm that SR artifacts are uploaded to the SR vendor’s infrastructure.

Table 5 shows part of the apps we exercised. We redacted app package names as the studied apps are closed-source, proprietary products. In Case 1, the app de-redacts a bottom-sheet during onViewCreated by calling ..., exposing on-screen text/inputs. Case 2 links identity by setting UXCam.setUserProperty(...), binding our input dummy email address to the session. Case 3 enriches the session timeline by logging a "Signout" event via UXCam.logEvent(...), where *map* denotes a key-value dictionary of user-profile PII. Case 4 pre-enables recording at app startup by invoking FS.consent(true).

Network traffic analysis. To further corroborate that SR data leaves the device, we captured outbound traffic from the Case 2 app during execution and observed the canonical UXCam upload workflow. The app first issues a verification request to `https://verify-{{region_code}}.uxcam.com/v4/verify`, after which the server returns a session identifier, capture settings (including `videoRecording=true`), and pre-signed upload parameters that authorize the client to upload two session artifacts: a `video.zip` and a `data.zip`. Immediately afterward, the device performs multipart uploads to UXCam-controlled storage endpoints (`https://prod-video-zip-{{region_code}}-1.uxcam-storage.com/` and `https://prod-data-zip-{{region_code}}-1.uxcam-storage.com/`), both returning HTTP 200 OK. The uploaded files are encrypted at the application layer (e.g., `video.mp4.aes` and `data.json.gz.aes`), so we do not inspect plaintext contents; nevertheless, the `verify` → `upload` sequence provides end-to-end evidence that SR session artifacts (screen recording and associated interaction/state data) are exported off-device to the SR vendor infrastructure.

Dynamic validation confirms that risk-relevant SR behaviors do occur at runtime. Network captures further show that these replay artifacts are subsequently uploaded off-device to SR vendor endpoints.

6 Discussion

6.1 Implications of Results

Our findings show actionable implications for developers, SR vendors, and platform operators, requiring coordinated measures to mitigate privacy risks of SR usage.

For Application Developers. The asymmetry between aggressive data enrichment and coarse, default masking indicates that apps often maximize analytics while investing little in privacy hygiene. To reduce leakage risk, developers could adopt a *deny-by-default* masking strategy: mask all UI elements at startup and explicitly unmask only the non-sensitive components needed for analysis, rather than relying on SDK defaults. Developers might also routinely audit the data sent through enrichment APIs to prevent PII from being logged as custom properties. Finally, the prevalence of R1 suggests that Data Safety disclosures may be maintained as living documents and reviewed as part of the release cycle to ensure SR collection is transparently reported to users [12, 29].

For SR Vendors. Although vendors often advertise SR as “private by default,” their APIs can make it easy to disable protections globally. For example, broad unmasking primitives (e.g., unmasking an activity’s root view) effectively provide a one-line way to nullify masking, the primary safeguard. Vendors can improve safety by adopting safer vendor defaults and discouraging risky patterns: deprecate or harden global unmasking in favor of APIs that require explicit, fine-grained unmasking of specific elements, introducing intentional friction against wholesale privacy disablement. Vendors could also build in privacy nudges and detection: emit high-priority debug warnings when sensitive UI (e.g., password fields) is unmasked or when identify/enrichment arguments match common PII patterns, and augment dashboards with backend monitoring that flags newly observed sensitive fields in uploaded streams.

For Platform Operators. The widespread failure to disclose SR in Data Safety labels shows that self-reporting alone is unreliable. Platforms should therefore incorporate SR detection into app review: scan submissions for SR integrations and cross-check them against declared data practices; if SR is present but not disclosed, require correction before approval. To strengthen informed consent, a potential direction for further study is to explore whether platform-supported indicators for active SR recording, analogous to camera or microphone indicators, can improve real-time user awareness beyond policy text alone. However, such mechanisms would require careful design and evaluation, since prior usable-privacy research shows that icons and indicators do not necessarily convey privacy choices effectively and may risk habituation [27]. Moreover, deploying such support would likely require substantial platform-level adoption, which can be challenging to realize in practice. Apple’s 2019 intervention [46] was a necessary step, but it did not solve the broader problem of under-disclosure.

More broadly, these mitigations are amenable to tool-supported enforcement: automated analyses such as TAPSPY can help developers catch risky SR configurations before release, help vendors identify risky integration patterns in deployed apps, and help platforms prioritize apps for disclosure review.

6.2 Limitations

Impact of obfuscation and implicit usage. To ensure broad coverage, TAPSPY incorporates best-effort defenses against common malicious evasion/benign protection techniques. However, limitations remain regarding advanced obfuscation and dynamic behaviors [17]. Techniques such as control-flow flattening or aggressive string encryption may still hinder Phase 2, resulting in incomplete or overly complex SRUS slices that can degrade the fidelity of the LLM’s semantic reasoning. Moreover, while we resolve reflection and dynamic code loading involving constants, targets derived from runtime-external sources (e.g., network configuration payloads or complex user inputs) remain invisible to our static pipeline. Consequently, apps relying on fully dynamic remote code loading to deploy SR logic may still evade detection. Accordingly, our measured prevalence should be interpreted as estimates over the subset of SR ecosystems that are observable through analyzable SDK artifacts and recoverable app code; they may therefore underapproximate the true prevalence of SR deployment and SR-related risks in the broader Android ecosystem.

Geographical differences. Although our case studies include apps from different countries and demonstrate that SR-related privacy risks arise in diverse geographic contexts, it is not sufficient to support claims about geographic disparity or the impact of local regulations. A rigorous treatment would require reliably attributing each app (and its SR configuration) to developer location and user-region targeting [25], and then controlling for confounders such as app category, market segment, and vendor choice. We leave a large-scale, geographically analysis to future work.

Threats to validity. Our analysis targets native Android apps; services that support integrated via cross-platform frameworks (e.g., React Native, Flutter) could not be analyzed, limiting generalizability beyond the native ecosystem. TAPSPY also examines only compiled APKs, missing “out-of-band” configurations such as vendor dashboard settings or Gradle build scripts. These external mechanisms may under-report certain SR features or risks.

For SR services which we rule out due to lack of public docs and packages, we cannot find corresponding package or class names, therefore it is hard to know their exact distribution in our dataset or real world, which might impact our default list of SR services currently collected and considered in our measurement study.

7 Related Work

We discuss several lines of prior work closely related to our analysis of SR services in mobile apps.

Analysis of Android Third-Party Libraries. Work on detecting third-party libraries (TPLs) in Android apps is extensive. LibRadar clusters and fingerprints packages at scale [33], while LibScout matches versioned signatures [13]. Later systems improve obfuscation resilience: LibID [54] remains accurate under renaming, shrinking, and control-flow randomization, and LibScan [48] uses cost-effective fingerprints with over-approximated class correspondences. Beyond detection, studies examine risks from outdated or vulnerable TPLs, showing pervasive outdatedness and challenges of updating mobile dependencies [19, 45].

Risks Due to Third-Party Libraries. Research on SDK/TPL risks spans multiple threads. Network- and app-level studies reveal

widespread third-party tracking and policy gaps [11, 12, 14, 32, 38, 40, 50, 55–57]. More broadly, default SDK choices strongly influence privacy outcomes before customization [30].

Session Replay (SR) Studies. SR research has focused on the web. Embedded SR scripts have been shown to exfiltrate keystrokes, unsubmitted form data, and credentials [21, 22]. Large-scale studies reveal pre-submit exfiltration of emails and passwords [39], SR deployment on 3.6% of hospital sites leaking sensitive patient data [52], and frequent nondisclosure of SR use in privacy policies [28].

Our work connects to a growing body of mobile app privacy research showing persistent gaps between declared and actual data practices, as well as recurring failures of notice and consent mechanisms. As prior work shows that privacy failures often stem from misalignment among disclosures [11, 12, 14, 49, 50, 57], implemented behavior [20, 53, 55], and meaningful user choice [34–36]. Our findings place SR within this broader pattern, but through a distinct instrumentation layer where the resulting exposure can be especially consequential. In this sense, SR does not create an entirely new class of privacy failure; rather, it introduces a new technical vector through which known problems—policy–practice inconsistency, weak transparency, and ineffective consent—can be amplified: they may manifest in especially rich and sensitive form as SR captures fine-grained UI state, user interactions, and identity-linked contextual data.

Finally, instead of addressing malware behaviors [15, 24], our threat model assumes benign developers and benign apps, with risks arising from unintended SR usage, misconfiguration, or disclosure inconsistencies rather than purposeful malicious intent.

8 Conclusion

We presented the first large-scale study of mobile Session Replay (SR) usages. To enable this study, we also developed TAPSPY, a novel framework that synergizes program analysis with large language models (LLMs) to realize semantic usage analysis of SR services and evidence-guided risk assessment in mobile apps. Using TAPSPY, we audited over 47,000 Android apps and uncovered systemic privacy failures: most apps omit disclosures, developers enrich data for analytics while neglecting privacy, and risks cluster into dangerous bundles of sensitive logging, identification, and nondisclosure. These findings call for coordinated action from developers, vendors, and platforms to enforce transparency and safeguard users.

Acknowledgment

We thank reviewers for their constructive comments. This work was supported by the National Science Foundation (NSF) under Grant CCF-2505223 and Office of Naval Research (ONR) under Grant N000142512252.

References

- [1] 2025. App Analysis: Air Canada. <https://theappanalyst.com/aircanada.html>.
- [2] 2025. App Analytics SDKs Could Expose Sensitive Data. <https://www.security.com/expert-perspectives/app-analytics-sdks-could-expose-sensitive-data>
- [3] 2025. Application Performance Monitoring & Error Tracking Software. <https://sentry.io/welcome/>.
- [4] 2025. Best Session Replay Software: User Reviews from August 2025. <https://www.g2.com/categories/session-replay>.
- [5] 2025. Digital Analytics Platform. <https://www.quantummetric.com>.
- [6] 2025. Digital Customer Experience Analytics. <https://www.glassbox.com/>.
- [7] 2025. Gemini 2.5 Pro. <https://deepmind.google/models/gemini/pro/>.

- [8] 2025. Intro to Fullstory for Mobile Apps. <https://help.fullstory.com/hc/en-us/articles/4414642861463-Intro-to-Fullstory-for-Mobile-Apps>.
- [9] 2025. Papa John's Sued for 'wiretap' Spying on Website Visitors. https://www.theregister.com/2022/10/06/papa_johns_spying_lawsuit/.
- [10] 2025. Taco Bell Fast Food & Delivery - Apps on Google Play. https://play.google.com/store/apps/details?id=com.tacobell.ordering&hl=en_US.
- [11] Benjamin Andow, Samin Yaseer Mahmud, Justin Whitaker, William Enck, Bradley Reaves, Kapil Singh, and Serge Egelman. 2020. Actions speak louder than words: {Entity-Sensitive} privacy policy and data flow analysis with {PoliCheck}. In *USENIX Security*. 985–1002.
- [12] Ioannis Arkalakis, Michalis Diamantaris, Serafeim Moustakas, Sotiris Ioannidis, Jason Polakis, and Panagiotis Iliia. 2024. Abandon All Hope Ye Who Enter Here: A Dynamic, Longitudinal Investigation of Android's Data Safety Section. In *USENIX Security*. 5645–5662.
- [13] Michael Backes, Sven Bugiel, and Erik Derr. 2016. Reliable third-party library detection in android and its security applications. In *CCS*. 356–367.
- [14] Duc Bui, Yuan Yao, Kang G Shin, Jong-Min Choi, and Junbum Shin. 2021. Consistency analysis of data-usage purposes in mobile apps. In *CCS*. 2824–2843.
- [15] Haipeng Cai, Xiaoqin Fu, and Abdelwahab Hamou-Lhadj. 2020. A study of runtime behavioral evolution of benign versus malicious apps in android. *IST* 122 (2020), 106291.
- [16] Haipeng Cai and John Jenkins. 2018. Leveraging historical versions of Android apps for efficient and precise taint analysis. In *MSR*. 265–269.
- [17] Haipeng Cai and Barbara Ryder. 2020. A longitudinal study of application structure and behaviors in android. *TSE* 47, 12 (2020), 2934–2955.
- [18] Haipeng Cai and Raul Santelices. 2015. Abstracting program dependencies using the method dependence graph. In *QRS*. 49–58.
- [19] Erik Derr, Sven Bugiel, Sascha Fahl, Yasemin Acar, and Michael Backes. 2017. Keep me updated: An empirical study of third-party library updatability on android. In *CCS*. 2187–2200.
- [20] Xiaolin Du, Zheming Yang, Jiapeng Lin, Yinzi Cao, and Min Yang. 2024. Withdrawing is believing? detecting inconsistencies between withdrawal choices and third-party data collections in mobile apps. In *S&P*. 735–751.
- [21] Steven Englehardt, Gunes Acar, and Arvind Narayanan. 2017. No boundaries: Exfiltration of personal data by session-replay scripts. <https://blog.citp.princeton.edu/2017/11/15/no-boundaries-exfiltration-of-personal-data-by-session-replay-scripts/>. *Freedom to tinker* 15 (2017).
- [22] Steve Englehardt, Gunes Acar, and Arvind Narayanan. 2018. No boundaries for credentials: New password leaks to Mixpanel and Session Replay Companies. <https://blog.citp.princeton.edu/2018/02/26/no-boundaries-for-credentials-password-leaks-to-mixpanel-and-session-replay-companies/>.
- [23] Jiawei Guo and Haipeng Cai. 2026. EvoTaint: Incremental static taint analysis of evolving Android apps. *TOSEM* 35, 4 (2026), 1–46.
- [24] Jiawei Guo, Xiaoqin Fu, Li Li, Tao Zhang, Mattia Fazzini, and Haipeng Cai. 2025. Characterizing installation-and run-time compatibility issues in Android benign apps and malware. *TOSEM* 35, 1 (2025), 1–44.
- [25] Jiawei Guo, Yu Nong, Zhiqiang Lin, and Haipeng Cai. 2025. Code Speaks Louder: Exploring Security and Privacy Relevant Regional Variations in Mobile Applications. In *S&P*. 3952–3970.
- [26] Jiawei Guo, Haoran Yang, and Haipeng Cai. 2025. VerLog: Enhancing release note generation for Android apps using large language models. *PACMSE* 2, ISSTA (2025), 1910–1932.
- [27] Hana Habib, Yixin Zou, Yaxing Yao, Alessandro Acquisti, Lorrie Cranor, Joel Reidenberg, Norman Sadeh, and Florian Schaub. 2021. Toggles, dollar signs, and triangles: How to (in) effectively convey privacy choices with icons and link texts. In *CHI*. 1–25.
- [28] Daichi Kajima and Hiroaki Kikuchi. 2023. Failure of Privacy Policy for Session Replay Services Used for Monitor Your Keystroke. In *NBIS*. Springer, 123–129.
- [29] Rishabh Khandelwal, Asmit Nayak, Paul Chung, and Kassem Fawaz. 2024. Unpacking privacy labels: A measurement and developer perspective on google's data safety section. In *USENIX Security*. 2831–2848.
- [30] Simon Koch, Manuel Karl, Robin Kirchner, Malte Wessels, Anne Paschke, and Martin Johns. 2025. The Impact of Default Mobile SDK Usage on Privacy and Data Protection. *PETS* 2025 (2025), 808–823. Issue 1.
- [31] Shuai Li, Zheming Yang, Nan Hua, Peng Liu, Xiaohan Zhang, Guangliang Yang, and Min Yang. 2022. Collect Responsibly But Deliver Arbitrarily?: A Study on Cross-User Privacy Leakage in Mobile Apps. In *CCS*. 1887–1900.
- [32] Xing Liu, Jiqiang Liu, Sencun Zhu, Wei Wang, and Xiangliang Zhang. 2019. Privacy risk analysis and mitigation of analytics libraries in the android ecosystem. *TMC* 19, 5 (2019), 1184–1199.
- [33] Ziang Ma, Haoyu Wang, Yao Guo, and Xiangqun Chen. 2016. Libadar: Fast and accurate detection of third-party libraries in android apps. In *ICSE Companion Proceedings*. 653–656.
- [34] Trung Tin Nguyen, Michael Backes, Ninja Marnau, and Ben Stock. 2021. Share first, ask later (or never?) studying violations of {GDPR's} explicit consent in android apps. In *USENIX Security*. 3667–3684.
- [35] Trung Tin Nguyen, Michael Backes, and Ben Stock. 2022. Freely Given Consent?: Studying Consent Notice of Third-Party Tracking and Its Violations of GDPR in Android Apps. In *CCS*. 2369–2383.
- [36] Midas Nouwens, Janus Bager Kristensen, Kristjan Maalt, and Rolf Bagge. 2025. A Cross-Country Analysis of GDPR Cookie Banners and Flexible Methods For Scraping Them. In *CHI*. 1–28.
- [37] Ole André Vadla Ravnås. 2025. Frida: A world-class dynamic instrumentation toolkit. <https://frida.re/>.
- [38] Abbas Razaghpanah, Rishabh Nithyanand, Narseo Vallina-Rodriguez, Srikanth Sundaresan, Mark Allman, and Christian Kreibich Phillipa Gill. 2018. Apps, trackers, privacy, and regulators. In *NDSS*. 1–15.
- [39] Asuman Senol, Gunes Acar, Mathias Humbert, and Frederik Zuiderveen Borgesius. 2022. Leaky forms: A study of email and password exfiltration before form submission. In *USENIX Security*. 1813–1830.
- [40] Yun Shen, Pierre-Antoine Vervier, and Gianluca Stringhini. 2021. Understanding Worldwide Private Information Collection on Android. *arXiv:2102.12869* [cs].
- [41] skylot. 2025. Skylot/Jadx. <https://github.com/skylot/jadx>.
- [42] Xiaoyu Sun, Li Li, Tegawendé F Bissyandé, Jacques Klein, Damien Octeau, and John Grundy. 2021. Taming reflection: An essential step toward whole-program analysis of android apps. *TOSEM* 30, 3 (2021), 1–36.
- [43] Zeya Tan and Wei Song. 2023. PTPDroid: Detecting Violated User Privacy Disclosures to Third-Parties of Android Apps. In *ICSE*. 473–485.
- [44] Raja Vallée-Rai, Phong Co, Etienne Gagnon, Laurie Hendren, Patrick Lam, and Vijay Sundaresan. 2010. Soot: A Java bytecode optimization framework. In *CASCON First Decade High Impact Papers*. 214–224.
- [45] Ying Wang, Bihuan Chen, Kaifeng Huang, Bowen Shi, Congying Xu, Xin Peng, Yijian Wu, and Yang Liu. 2020. An empirical study of usages, updates and risks of third-party libraries in java projects. In *ICSE*. 35–45.
- [46] Zack Whittaker. 2019. Apple Tells App Developers to Disclose or Remove Screen Recording Code. <https://techcrunch.com/2019/02/07/apple-glassbox-apps/>.
- [47] Zack Whittaker. 2019. Many Popular iPhone Apps Secretly Record Your Screen without Asking. <https://techcrunch.com/2019/02/06/iphone-session-replay-screenshots/>.
- [48] Yafei Wu, Cong Sun, Dongrui Zeng, Gang Tan, Siqi Ma, and Peicheng Wang. 2023. {LibScan}: Towards more precise {Third-Party} library identification for Android applications. In *USENIX Security*. 3385–3402.
- [49] Yue Xiao, Zhengyi Li, Yue Qin, Xiaolong Bai, Jiale Guan, Xiaojing Liao, and Luyi Xing. 2023. Lalaine: Measuring and characterizing {Non-Compliance} of apple privacy labels. In *USENIX Security*. 1091–1108.
- [50] Yue Xiao, Chaoqi Zhang, Yue Qin, Fares Fahad S Alharbi, Luyi Xing, and Xiaojing Liao. 2024. Measuring Compliance Implications of Third-party Libraries' Privacy Label Disclosure Guidelines. In *CCS*. 1641–1655.
- [51] Fabian Yamaguchi. 2022. A platform for robust analysis of C/C++ code. <https://joern.readthedocs.io/en/latest/installation.html>.
- [52] Xiufen Yu, Nayanamana Samarasinghe, Mohammad Mannan, and Amr Youssef. 2022. Got sick and tracked: privacy analysis of hospital websites. In *European Symposium on Security and Privacy Workshops*. 278–286.
- [53] Chang Yue, Kai Chen, Zhixiu Guo, Jun Dai, Xiaoyan Sun, and Yi Yang. 2025. What's Done Is Not What's Claimed: Detecting and Interpreting Inconsistencies in App Behaviors. In *NDSS*. 1–18.
- [54] Jiexin Zhang, Alastair R Beresford, and Stephan A Kollmann. 2019. Libid: reliable identification of obfuscated third-party android libraries. In *ISSTA*. 55–65.
- [55] Xueling Zhang, Xiaoyin Wang, Rocky Slavin, Travis Breaux, and Jianwei Niu. 2020. How does misconfiguration of analytic services compromise mobile privacy? In *ICSE*. 1572–1583.
- [56] Yifan Zhang, Zhaojie Hu, Xueqiang Wang, Yuhui Hong, Yuhong Nan, XiaoFeng Wang, Jiatao Cheng, and Luyi Xing. 2024. Navigating the Privacy Compliance Maze: Understanding Risks with {Privacy-Configurable} Mobile {SDKs}. In *USENIX Security*. 6543–6560.
- [57] Kaifa Zhao, Xian Zhan, Le Yu, Shiyao Zhou, Hao Zhou, Xiapu Luo, Haoyu Wang, and Yeping Liu. 2023. Demystifying Privacy Policy of Third-Party Libraries in Mobile Apps. In *ICSE*. 1583–1595.

Ethics Considerations

We analyzed only publicly available APKs and did not access user data. We anonymize affected commercial apps and omit exploit-enabling details. Before publication, we disclosed high-severity findings to affected developers; 16 confirmed our findings so far, and several updated their privacy policies or Data Safety disclosures.

Open Science

Our artifact for this work is publicly accessible at <https://doi.org/10.5281/zenodo.20600860>, including the source code for our TapSPy framework, the scripts used for our large-scale measurement study, and the resulting dataset of detected risks and feature usage.