

Полтавцева Мария Анатольевна¹,
профессор, д-р техн. наук, доцент;
Платонов Владимир Владимирович²,
доцент, канд. техн. наук, доцент;
Супрун Александр Федорович³,
доцент, канд. техн. наук, доцент;
Семьянов Павел Валентинович⁴,
ст. преподаватель

МЕТОДЫ АНОНИМИЗАЦИИ ПЕРСОНАЛЬНЫХ ДАННЫХ

^{1,2,3,4} Россия, Санкт-Петербург,
Санкт-Петербургский политехнический университет Петра Великого,
Высшая школа кибербезопасности,
¹ poltavtseva@ibks.spbstu.ru, ² plato@ibks.spbstu.ru,
³ suprun@ibks.spbstu.ru, ⁴ psw@ibks.spbstu.ru

Аннотация. Обеспечение безопасности персональных данных, в том числе при их обработке в ГИС, вычислении статистических показателей, научных целях или при обучении интеллектуальных систем является сложной задачей. Для сохранения приватности в таких случаях используются методы анонимизации данных. Однако достаточно большое различие методов в сочетании с отсутствием универсального решения и достаточно общими нормативными требованиями усложняют данную задачу. В работе проводится анализ нормативных требований и систематизация методов анонимизации данных применительно к задаче обеспечения приватности при использовании персональных данных. Уточняется понятие полезности данных и формируются сравнительные критерии выбора метода анонимизации для отдельных задач.

Ключевые слова: информационная безопасность, персональные данные, обезличивание, анонимизация, атаки логического вывода, атаки повторной идентификации, приватность, модель злоумышленника.

*Maria A. Poltavtseva*¹,
Professor, Doctor of Technical Sciences;
*Vladimir V. Platonov*²,
Associate Professor, Candidate of Technical Sciences;
*Aleksandr F. Suprun*³,
Associate Professor, Candidate of Technical Sciences;
*Pavel V. Semyanov*⁴,
Senior Lecturer

METHODS OF PERSONAL DATA ANONYMIZATION

^{1,2,3,4} High School of Cybersecurity,
Peter the Great St. Petersburg Polytechnic University, St. Petersburg, Russia,
¹ poltavtseva@ibks.spbstu.ru, ² plato@ibks.spbstu.ru,
³ suprun@ibks.spbstu.ru, ⁴ psw@ibks.spbstu.ru

Abstract. Ensuring the security of personal data, including when processing them in GIS, calculating statistical indicators, for scientific purposes or for training intelligent systems is a complex task. To preserve privacy in such cases, data anonymization methods are used. However, the rather large number of different methods combined with the lack of a universal solution and rather general regulatory requirements complicate this task. In this paper we analyze the regulatory requirements and systematize data anonymization methods with respect to the problem of privacy preservation in the use of personal data. The notion of data utility is clarified and comparative criteria for selecting anonymization methods for individual tasks are formed.

Keywords: information security, personal data, depersonalization, anonymization, logical inference attacks, re-identification attacks, privacy, attacker model.

Введение

Лавинообразный рост количества данных в мире не останавливается последние десятилетия. При этом, с одной стороны растет количество задач и способов использования информации: предоставление услуг на основе данных; внутренние оценки, маркетинговые стратегии и другие использование на нужды компании; расчет и публикация статистики; предоставление публичных данных по организациям, доходам и т. д.; подготовка интеллектуальных систем на основе машинного обучения, в том числе с обменом данными между организациями и множество других. Данные собирают как государственные и муниципальные органы, так и частные компании, и организации.

В то же время растет число утечек данных, которые вызваны как прямой кражей информации, так и объединением нескольких открытых наборов данных из разных источников (например, социальных сетей) для их «обогащения» и, таким образом, получения злоумышленником конфиденциальных сведений.

Конечно, очевидный способ защиты – шифрование и контроль доступа к данным. Однако, эти способы не всегда применимы. Есть несколько видов использования информации, для которых традиционные методы обеспечения безопасности не являются панацеей. Атаки класса логического вывода (inference attack) продолжают оставаться значимой угрозой. Особенностью таких атак является тот факт, что они проводятся без нарушения политики безопасности. Аналитик или сторона, делающая «вывод», не совершает никаких действий, чтобы нарушить правила доступа, а пользуется доступными ей данными для получения конфиденциальной информации – построения вывода.

Борьба с такими атаками, с одной стороны, является важной частью безопасности и сохранения приватности в современном мире, а с другой – полная защита от логического вывода является пр-полной задачей [1]. Цель данной работы систематизация подходов методов защиты от логического вывода в современных реалиях.

1. Атаки логического вывода и атаки повторной идентификации

Реализация атак логического вывода связана с построением нарушителем ассоциаций между данными [2]. В основе таких ассоциаций – закономерности предметной области, функциональные зависимости между данными и даже интеллектуальные методы (data mining). Техника интеллектуального анализа данных здесь используется для поиска шаблонов и построения правил на основе данных и ассоциаций.

Фрагмент классификации атак логического вывода приведен на рисунке 1. Среди всего множества таких атак выделим атаки идентификации. Это атаки, которые позволяют в наборе данных (в том числе, обезличенном, из которого удалена персональная информация) установить кортежи, или сведения, принадлежащие конкретному субъекту. Вероятность такого вывода показана в работе [3].

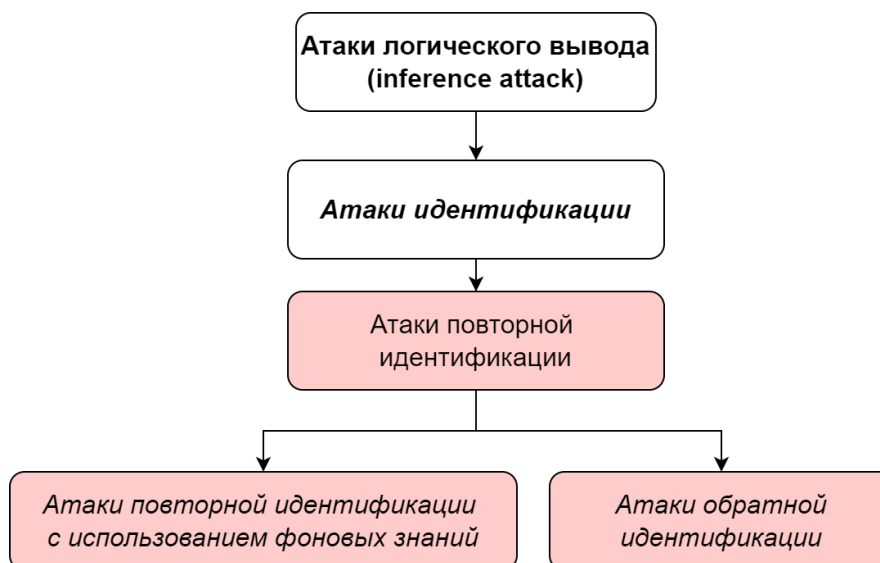


Рис. 1. Место атак повторной идентификации среди атак логического вывода

Повторной идентификацией называется установление по защищенным данным сведений, относящихся к конкретному лицу (объекту). Установить принадлежность конфиденциальных данных конкретному объекту («связать» данные) возможно двумя способами идентификации: обратной идентификацией или идентификацией с использованием фоновых значений.

Обратная идентификация – это установление по защищенным данным информации о конкретном объекте путем восстановления исходных данных на основе только обезличенного набора. Надо сказать, что обратная идентификация в каком-то смысле тоже может относиться к атакам повторной идентификации с использованием фоновых знаний, так как зачастую она реализуется за счет знания злоумышленником алгоритма обезличивания или его особенностей.

Повторная идентификация с использованием фоновых знаний – это установление по защищенным данным сведений, относящихся к конкретному лицу (объекту) путем восстановления исходных данных с использованием внешних знаний и/или данных. К последнему случаю относятся ситуации объединения нескольких наборов данных, использование личных знаний нарушителя или знаний о субъектах данных из других источников. Например, использование знания предметной области, использование характеристик базы данных из которой получена информация (сведений о ключах, связях, зависимостях и других данных).

Необходимость защиты данных от атак идентификации закреплена в законодательстве о персональных данных. Федеральный закон № 152 «О персональных данных», в статье 3 «Основные понятия, используемые в настоящем Федеральном законе», пункт 8 появляется понятие «обезличивание персональных данных» и само определение соответствующего понятия, как действий, в результате которых становится невозможным без использования дополнительной информации определить принадлежность персональных данных конкретному субъекту персональных данных.

Важным в этом определении является указание, что при применении обезличивания наличие любых дополнительных знаний злоумышленника не учитывается. Следовательно, применяемые для обезличивания данных методы должны защищать только от атак обратной идентификации, причем таких, для которых знание алгоритма обезличивания или его особенностей не является необходимым условием реализации угрозы нарушения конфиденциальности.

Прочие нормативные акты устанавливают методику обезличивания (Приказ Роскомнадзора № 996), требование удаления ассоциаций (Стандарт ISO 25237-2008, Стандарт ГОСТ Р 55036-2012) и другие. Однако все эти требования либо являются крайне высокоуровневыми, либо говорят об устранении идентификаторов и не учитывают модель нарушителя, что

вызывает сомнения в их практической эффективности без дополнительных действий. В то же время, сегодня уже появились ведомственные документы, описывающие мелодики практической анонимизации (см. например, Приказ Министерства здравоохранения РФ от 24 ноября 2023 г. № 630н). Однако для отдельных лиц и организаций задача выбора соответствующего метода является достаточно сложной.

2. Анализ методов защиты данных путем анонимизации

На сегодняшний день продолжает расти число работ по анонимизации данных: исследователи рассматривают задачи обезличивания потоковых данных [4], графов включая графы знаний и социальных сетей [5, 6], каждый из которых обладает некоторой спецификой применения. В различных работах можно найти такие методы, как исключение критических атрибутов, добавление шума к данным [7]. Большинство методов анонимизации с технической точки зрения можно разделить на две группы:

- 1) методы, связанные с добавлением шума к данным;
- 2) методы, связанные с обобщением данных.

В первом случае сохраняется точность данных, однако конкретные значения искажаются. Во втором – изменяется само значение, зачастую вместе с удалением части (или даже всей) информацией о конкретном атрибуте или записи.

На рис. 2 приведена обобщенная систематизация современных методов анонимизации данных. В основе классификации – модель нарушителя, определенная на основе его намерений, которая позволяет соотнести методы защиты данных с терминологией в области анонимизации (обезличивания).



Рис. 2. Классификация методов анонимизации данных с учетом модели злоумышленника

В первом случае рассматривается нарушитель со случайным доступом и для него применяется подход недоступности, или, более корректно, не применения таким нарушителем возможных фоновых знаний для построения логического вывода и проведения атаки повторной идентификации над данными. Можно сказать, что методы, отнесенные к этой группе, относятся именно к методам обезличивания данных и для выполнения задачи обезличивания их применения достаточно. Внутренне разделение таких методов основано на обратимости результата, который может быть получен в результате обезличивания и технологических особенностях обработки информации.

Вторую большую группу представляют собой методы так называемой «гарантированной» или «предполагаемой» анонимизации, основанные на «предположениях» об информированности злоумышленника и обеспечивающие конфиденциальность информации в тех или иных условиях. Они разделены не по принципу обратимости (так как, в общем случае, все эти методы являются необратимыми без ведения соответствующих таблиц соответствия исходных данных и защищенных).

Методы на основе модели доверенного обработчика предполагают, что лицо (организация, программный компонент и т. д.) выполнившее анонимизации является доверенным, и результат такой анонимизации уже подлежит использованию. К этой группе относятся методы k-анонимизации и ее расширений, а также «сервисная» модель работы с данными на основе дифференциальной приватности. В первом случае весь набор данных предоставляется аналитику, во втором – речь идет о сохранении данных у обработчика, а аналитику представляется только ответ на его запросы, а не весь массив информации. К методам на основе модели не доверенного обработчика сегодня можно отнести только метод локальной дифференциальной приватности, направленной на обеспечение конфиденциальности данных уже на этапе получения их от источника для обработки.

Сопоставление методов анонимизации (будем в совокупности называть их так) с методами (подходами), определенными в приказе Роскомнадзора № 996 от 05.09.2013 приведено в таблице 1. Это сопоставление может оказаться полезным при составлении соответствующих документов по выполнению требований Федерального закона «О персональных данных» от 27.07.2006 № 152-ФЗ.

В то же время выбор конкретного метода анонимизации является задачей независимым и зависит от принятой оценки полезности данных.

Таблица 1

**Соотнесение методов, предложенных Роскомнадзором,
и технологических методов анонимизации данных**

Методы обезличивания согласно приказу РКН № 996	Методы анонимизации с технологической точки зрения	Примечания
Введение идентификаторов	Подстановка	Сохранение связи с исходными значениями – ключевой показатель отнесения в эту группу. Подстановка без сохранения связей относится к другой группе методов.
Изменение состава или семантики	Замена на константу Размытие Микроагрегация Криптографические преобразования К – анонимизация и ее расширения Дифференциальная приватность	Вместе с использованием идентификаторов этот метод может позволять обеспечить безусловную (согласно приказу № 996) анонимность, сохранив полноту и семантическую целостность данных при доступе к обоим наборам. Для этого необходимо вынести ключевые атрибуты в отдельную (например) таблицу, а в анонимизируемой таблице преобразовать, связав каждую строку с чувствительными данными только случайным идентификатором.
Декомпозиция	Декомпозиция	Разделение данных на части может происходить по различным принципам и относится к решению задач хранения персональных данных в облачных системах с обеспечением обезличивания.
Перемешивание	Обратимое или необратимое перемешивание	Обратимое перемешивание – по определенному детерминированному алгоритму. Необратимое – с использованием случайных значений при выборе соответствия

3. Выбор метода анонимизации

Можно говорить о двух базовых подходах к гарантированной анонимности: вероятностный подход, основанный на той или иной оценке вероятности логического вывода злоумышленником, и более простой подход на основе различимости или «не различимости» анонимизированных записей. В свою очередь, «гарантировать» хотя бы в некоторой степени (или

с некоторой вероятностью) защиту при обезличивании данных можно только используя оценку информированности злоумышленника, сформированную тем или иным образом. Так как, повторим, сформировать точную оценку такой информированности на сегодняшний день практически невозможно, существующие методы и их расширения направлены на минимизацию риска раскрытия с сохранением полезности данных для дальнейшего использования.

С точки зрения обработки информации все представленные сегодня методы используют один из трех приемов или их сочетание: внесение шума, обобщение, подавление. Первые два применяются самостоятельно, третий метод – подавление – используется как вспомогательный при обеспечении k-анонимности данных или дифференциальной конфиденциальности. При k-анонимизации подавление, как мы и рассматривали, применяется к данным, а при дифференциальной конфиденциальности – к запросам, когда запрос (или тип запросов), которые могут привести к раскрытию определенной схемы обезличивания, могут быть отклонены. Сравнение этих подходов приведено в таблице 2.

Таблица 2

Сравнительная таблица методов обеспечения анонимности

Способ обработки данных	Тип атрибутов		Достоинства	Недостатки	Основная технология
	В*	К*			
Внесение шума	+	+	Разные типы шума, разные модели обработки данных	Подвержена отдельным атакам, искажает данные	Предположительная анонимность, дифференциальная конфиденциальность
Обобщение	+	+	Исходные данные обобщаются, но не искажаются	Подвержена отдельным атакам	К-анонимность
Подавление (данных или запроса)	+	+	Существенно улучшает баланс конфиденциальности и полезности (как вспомогательный метод)	Подвержена отдельным атакам, часть данных исключается	К-анонимность (данных), дифференциальная конфиденциальность (запроса)

* Вещественные и категориальные.

Нельзя не отметить еще несколько характерных особенностей. К-анонимность применяется к данным и является характеристикой данных. В то же время дифференциальная конфиденциальность – это свойство алгоритма. Поэтому при обезличивании данных, в общем случае, может быть применен дифференциально конфиденциальный алгоритм, в результате

использования которого формируется k-анонимизированный набор. На практике эффективный алгоритм такого рода пока не предложен.

Выбор того или иного метода всегда является балансом между полезностью полученных в результате обезличивания данных и безопасностью (конфиденциальностью) частных характеристик. Отметим, что метрика полезности является в некоторой степени субъективной. Она определяется спецификой конкретной задачи в виде набора требований к результату обезличивания (например, сохранение медианного или среднего значения данных после обработки) и дополнительно оценивается как некоторый показатель, сопоставляющий анонимизированные данные и исходные [8, 9].

С точки зрения применимости для различных задач, дифференциально конфиденциальные алгоритмы рандомизации широко используются в различных сферах, однако не закрывают полностью потребностей задач по обезличиванию. Искажение данных на основе шума не всегда является приемлемым, и направление k-анонимности развивается в отношении различных методов и расширений, направленных на устранение уязвимостей этого метода и повышению его устойчивости к конкретным типам атак.

Заключение

В работе были изучены нормативные требования к анонимизации персональных данных и технологические методы обезличивания, представленные в современных исследованиях и практических инструментах. Построена классификация методов анонимизации на основе модели злоумышленника и систематизированы технические решения в каждом классе. Проведено сравнение различных подходов к анонимизации с учетом типа исходных атрибутов.

Список литературы

1. Hale J., Shenoï S. Analyzing FD inference in relational databases // Data & Knowledge Engineering – 1996. – Vol. 18. – Pp. 167–183. – DOI:[https://doi.org/10.1016/0169-023X\(95\)00033-0](https://doi.org/10.1016/0169-023X(95)00033-0).
2. Hylkema M. A survey of database inference attack prevention methods [Electronic resource] // Educational Technology Research. – 2009. – URL: <https://tarjomefa.com/wp-content/uploads/2017/05/6608-English-TarjomeFa.pdf> (access date: 01.05.2024).
3. Rocher L., Hendrickx J. M., de Montjoye Y.-A. Estimating the success of re-identifications in incomplete datasets using generative models // Nat Commun – 2019. – Vol. 10. – Paper 3069. – DOI:<https://doi.org/10.1038/s41467-019-10933-3>.
4. Sopaoglu U., Abul O. A utility based approach for data stream anonymization // J Intell Inf Syst. – 2020. – No. 54. – Pp. 605–631. – DOI:<https://doi.org/10.1007/s10844-019-00577-6>.
5. Kiabod M., Naderi Dehkordi M., Barekatain B., Raahemifar K. FSopt_k: Finding the optimal anonymization level for a social network graph // Applied Sciences. – 2023. – Vol. 13. – No. 6. – Paper 3770. – DOI: 10.3390/app13063770.

6. Hoang A. T., Carminati B., Ferrari E. Protecting privacy in knowledge graphs with personalized anonymization // IEEE Transactions on Dependable and Secure Computing. – 2023. – Vol. 21 – No. 4 – Pp. 2181–2193.
7. Murthy S., Abu Bakar A., Abdul Rahim F., Ramli R. A comparative study of data anonymization techniques // 2019 IEEE 5th Intl Conference on Big Data Security on Cloud (BigDataSecurity), IEEE Intl Conference on High Performance and Smart Computing (HPSC) and IEEE Intl Conference on Intelligent Data and Security (IDS). – IEEE, 2019. – Pp. 306–309. – DOI:10.1109/BigDataSecurity-HPSC-IDS.2019.00063.
8. Domingo-Ferrer J., Sánchez D., Soria-Comas J. Database anonymization: privacy models, data utility, and microaggregation-based inter-model connections. – Springer Nature, 2022. – 119 p.
9. Lee H., Kim S., Kim J. W., Chung Y. Utility-preserving anonymization for health data publishing // BMC Medical Informatics and Decision Making. – 2017. – Vol. 17. – No. 104 – 12 p. –DOI: 10.1186/s12911-017-0499-0.