

С. В. Знаменский

Глобальная идентификация данных в долговременной перспективе

Аннотация. Ретроспективный подход к хранению данных характеризуется высокими неподдельностью данных, доступностью истории и живучестью. Эти качества особенно привлекательны для будущих масштабируемых систем, долгоживущих в сложно изменяющихся условиях.

Описываются система элементарных запросов к ретроспективному хранилищу, источники ветвления информационных структур. Приводятся рекомендации по выбору системы глобальной идентификации объектов ретроспективного хранилища.

Ключевые слова и фразы: система версионирования, хранилище данных, структура B+tree, метка времени, история изменений, долгоживущие информационные системы.

Введение

Целостная неподдельная история изменений незаменима и широко используется

- для эффективного согласования распределённых данных,
- для ускорения отладки нового кода,
- для быстрого восстановления при поломках,
- для упрощения расследования вторжений и других происшествий,
- для повышения ответственности персонала.

в организации более эффективного доступа к истории изменений перестраивающихся структур данных [1, 2]. Рост сложности и объёмов накапливаемых данных затрудняют замену используемой информационной систем. Всё чаще исследователи обращаются к идее информационной системы, способной к гибкой эволюции [3].

Работа проводилась при финансовой поддержке Министерства образования и науки Российской Федерации.

При этом накопление структурных изменений в истории данных способно экспоненциально увеличивать объём хранимой информации, существенно нелинейно увеличивая затраты на обработку запросов.

Мы сосредоточим внимание на истории изменений иерархических структур данных, обеспечивающих универсальность в известном смысле: Любые существующие структуры данных общепринято описывать в формате XML. Такое описание имеет иерархическую структуру, однозначно и полностью фиксирующую исходную структуру данных. Если исходная структура перестраивается, то изменяется и иерархия XML-описания. Таким образом, любые перестройки произвольных структур данных могут качественно сохраняться в виде перестроек иерархической структуры. Иерархическая структура данных здесь понимается как дерево, с узлами которого ассоциируются данные, сериализованные строками произвольной длины, а порядок следования детских узлов любого родителя фиксирован.

Следуя [4], рассмотрим задачу эффективного построения ретроспективной СУБД для хранения изменений иерархической структуры данных. Объектом будем называть ветку иерархии. Система в целом рассматривается как объект, содержащий все остальные объекты. Перечислим особенности ретроспективного доступа к данным, которые позволяют рассматривать его как перспективную основу для долгоживущей и масштабируемой в сложных условиях системы:

Неподделность: Любые данные сохраняются с фиксацией логического времени [6], настоящего или будущего и неизменны до их окончательного удаления из системы;

Живучесть: Окончательное удаление излишне детальной информации и задержки исполнения порождают «белые пятна истории объекта» — относительно малые легко обходимые промежутки времени, верное состояние информации внутри которых не обеспечено; острый недостаток ресурсов для хранения или обработки информации приведёт к многократному снижению качества, выраженном в увеличении белых пятен, но не к потере функциональности или доступности сервиса;

Доступность истории: Для каждого объекта доступен список всех моментов его изменения, список белых пятен и его состояние на любой доступный момент времени;

На основе качественно реализованного ретроспективного доступа могут реализовываться системы, эмулирующие изменение данных и эксперименты с прошлым и будущим, и другие будущие механизмы доступа к темпоральной информации.

Подход разрабатывается для приложений, приоритетно нуждающихся в сочетании повышенной доступности, жизнеспособности и гибкой модифицируемости информационного обеспечения.

Если например он будет положен в основу операционной системы, то любые системные обновления будут включаться или отменяться мгновенно, не отвлекая без необходимости пользователя и не создавая угроз необратимой потери доступа к данным. При этом система на компьютере будет частью распределённой системы разработки и поддержки этой ОС.

Для реализации такой системы ключевой является проработанность основ организации, обеспечивающих разумную экономию ресурсов. Предлагаемая статья описывает возможный подход к глобальной идентификации данных, независимой от их размещения в распределённой системе.

1. Основа ретроспективного доступа

Структура идентификаторов данных обязана обеспечить эффективность обработки данных в долговременной перспективе. Рассмотрим структуру запросов на ввод и вывод этих данных.

1.1. Элементарные ретроспективные запросы и архитектура локальной части хранилища

Каждый элементарный запрос может сопровождаться пометкой времени, относительно которого определяется актуальность данных. При обращении время может уточняться, об этом ниже.

- (1) чтение актуальных данных объекта, не входящих в подобъекты, и времени их последнего изменения;
- (2) чтение актуальных данных последнего подобъекта (включая ранее удалённые);
- (3) чтение актуальных данных предыдущего объекта (возможно удалённый ранее);
- (4) чтение всех согласованных данных объекта в целом, вместе с подобъектами;
- (5) чтение нуждающихся в согласовании данных объекта в целом, вместе с подобъектами;
- (6) запись объекта с пометкой времени;

Трёх первых видов запросов достаточно для получения полной информации об истории изменений данных например, на языке SPARQL [8]¹. Четвёртый вид предназначен для оптимизации удалённого дублирования данных, а пятый- для их согласования.

¹Для ускорения обработки полезно ввести служебные атрибуты, индексирующие информацию; например, время последнего изменения ветки, растущей из родительского узла.

Это означает, что перемещение объектов осуществляется копированием и удалением².

Поскольку согласованные данные в согласовании уже не нуждаются, то представляется разумным разделить локальное хранилище на две части: оперативно обновляемое хранилище текущей информации, участвующей в согласовании, и периодически заменяемое основное хранилище, оптимизированное на обработку запросов (1)–(3). При этом оперативно обновляемое хранилище должно поддерживать и операцию удаления данных, перенесённых в основное хранилище. Из-за более развитой функциональности продуктивность оперативно обновляемое хранилище сложнее. Но поскольку более 99% информации находится в основной неизменной структуре, доступ к которому легко сделать очень быстрым, то выигрыш от такого разделения должен с избытком компенсировать затраты на его поддержку.

1.2. Автономные объекты и белые пятна

Глобальная память системы предстаёт иерархией объектов. Это означает, что объекты связаны отношением непосредственной вложенности, граф которого G является деревом, направленным к корню - универсальному объекту u , содержащему все остальные. Каждому объекту выделяется сегмент адресов линейно упорядоченного адресного пространства A возможных глобальных идентификаторов информации. Вложенность объектов сохраняется для диапазонов их адресов.

Среди объектов выделяются первичные (в них информация обновляется извне, например, при заполнении веб-формы). Любой объект, не являющийся первичным, изменяется по ассоциированному с ним алгоритму, анализирующему данные других контекстов. Алгоритм тоже может время от времени заменяться.

Мы стараемся избежать организации обмена сообщениями между изолированными частями, чтобы избежать ошибок из-за изменяющейся адресации. Однако принцип модульности заставляет выделить в графе A априорно более тесно связанные части. Для этого из дерева G удаляются рёбра, выходящие из вершин $o_1, \dots, o_k$, и получается $k + 1$ компонент связности $C_1, \dots, C_k$. Начала удалённых рёбер назовём автономными объектами, а подмножества в A ,

²Экономное дублирование может быть в дальнейшем реализовано на описываемой основе через создание объектов особого типа - символических ссылок на объект-время и пометок об удалении. Ссылки и пометки об удалении будут располагаться в особом системном объекте, используемом при обработке неадминистративных запросов. К сожалению, это несколько замедлит выполнение всех запросов.

получаемые исключением из диапазона адресов автономного объекта диапазоны адресов вложенных в него объектов, — автономными контекстами.

Согласованное изменение объекта занимает время. Априорная оценка этого времени фиксируется в памяти как шаг дискретизации автономного объекта, действующий в автономном контексте. Шаг дискретизации становится ценой деления шкалы меток времени. Например, редактируемый пользователем текст может сохраняться каждую секунду, а сводить и сохранять статистику суточной активности достаточно раз в сутки и указания секунд времени её сохранения не требуется.

В некоторых частях объекта могут быть более подробные шкалы времени. Это обеспечивает контроль промежуточных состояний в ходе обработки. Однако запрос такого промежуточного времени не должен пониматься буквально если требуется согласованная информация объекта. В таких случаях считается, что запрос попал в белое пятно истории и запрашиваемое время уточняется.

Если время обработки изменений объекта оказалось больше ожидаемого, то белое пятно истории выходит за рамки промежутка между делениями шкалы. Такие белые пятна для запоминаются в автономных объектах и учитываются при чтении данных из остальных объектов контекста.

Будем считать, что шкалы согласованы в том смысле, что из шкал двух любых объектов одна всегда содержит другую

В системной части памяти для каждого автономного объекта поддерживается список объектов, подписавшихся на изменения, semaфор, показывающий наличие необработанных изменений и индикатор процесса, ответственного за обработку данных автономного контекста. Ответственный процесс:

- (1) Для каждого из ближайших вложенных автономных объектов при наличии необработанных в них изменений и отсутствия признаков жизни ответственного процесса убивает зависший процесс (если он не завершился сам) и запускает новый (либо, возможно, сам делает работу за него).
- (2) Если необработанных изменений больше нет, то помечает конец белого пятна ближайшей последующей меткой времени и завершает работу (либо засыпает до следующей отметки шкалы и возвращается к (1)).
- (3) Определяет отметку времени $t < t'$ последнего из всех изменений необработанных данных, предшествовавших последней прошедшей отметке t' временной шкалы.

- (4) Отметка шкалы $t'' > t$, непосредственно следующая за t , запоминается как начало белого пятна текущей обработки. Если последнее перекрывает часть предыдущего, то предыдущее белое пятно обрезается до начала белого пятна текущей обработки.
- (5) Повторяет цикл сначала.

Проверка наличия необработанных изменений делается по следующей схеме:

- (1) Сдвигаемся во времени назад до предыдущей отметки шкалы. Если эта отметка совпадает с началом завершённой обработки объекта, то заканчиваем: необработанных изменений нет.
- (2) Обходя все объекты, задействованные в алгоритме обновления, возвращаемся к (1) если только обнаружим, что отметка попала на текущую обработку данных.
- (3) Вот мы дошли до той отметки времени, когда все входные данные окончательно определены. Осталось проверить, были ли в период с начала предыдущей обработки изменения в объектах.

Мы предполагаем, что родительский объект определяется по идентификатору объекта. Это, в частности, позволяет запрашивать родителей по цепочке, чтобы получить информацию о статусе автономного контекста.

Таким образом, каждому обновлению автономного объекта предшествует нетривиальная проверка своевременности. Она нацелена на многократное сокращение ресурсоёмких обработок, в особенности при перегрузках и непредвиденном возрастании задержек, линейным увеличением длины белых пятен истории.

1.3. Уточнение времени запроса

При запросе на чтение прошлое время, попадающее в белое пятно, по умолчанию сдвигается на начало этого пятна. Это гарантирует мгновенное получение согласованных данных. Если, например, много пользователей обращаются с одинаковыми запросами, требующими много времени на обработку, то каждый из них будет мгновенно видеть последний из имеющихся согласованных результатов и время, на которое произведена обработка. Точность реляционных СУБД при этом сочетается с оперативностью NoSQL-сервисов. Хуже если запрос не очень частый, а нужна последняя актуальная информация, тогда придётся подождать и повторить запрос. Запрос может иметь признак, указывающий на необходимость чтения ближайших актуальных, пусть и не согласованных данных либо на чтение с актуальностью в конце пятна.

При запросе на запись время также уточняется: прошлое в пределах белого пятна фиксируется как время записи с уточнением в

виде настоящего времени, а прошлое или будущее время за пределами текущего белого пятна делают запись невозможной и запрос теряется.

Основное хранилище организует бинарное дерево для доступа к упорядоченному списку. Список изменений упорядочивается по времени и быстрый доступ к актуальной версии данного осуществляется делением всего упорядоченного списка изменений пополам до тех пор, пока не будет обнаружено изменение, актуальное на запрошенный момент времени. Это единственный известный способ перейти к актуальной на заданный момент времени информации для запроса (1),(4) с логарифмической сложностью.

Обработку (2)-(3) облегчает семантически насыщенная система идентификации узлов. При лексикографическом упорядочении таких идентификаторов родитель предшествует ребёнку и порядок детей одного родителя сохраняется, а идентификаторы узлов любой ветки заполняют интервал упорядоченного списка. Операция чтения дополнительной структуры на диске заменяется вычислением функции, задающей границы этого интервала, и чтением основной структуры.

2. Ветвления и их идентификация

2.1. Семантическая идентификация узлов

Система идентификации узлов, обеспечивающая связность множества идентификаторов узлов любой ветки, не является чем-то доселе неизвестным. Знакомой всем системой таких идентификаторов является идентификация файлов полным путём в файловой системе ISO 9660. Схема проста: имена содержащих файл директорий и имя самого файла конкатенируются с разделителем '/', обладающим важным свойством: этот разделитель не должен при лексикографическом упорядочении оказаться между возможными продолжениями имени любого имени файла или директории. Именно это свойство обеспечивает связность множества идентификаторов, а его нарушение угрожает нарушением этой связности.

Возможны и альтернативные варианты без разделителей. Можно, например, условиться, что имя файла или директории всегда начинается с заглавной латинской буквы, а продолжается только строчными латинскими буквами.

Эти варианты нуждаются в поправке в нашей ситуации, когда начало ветки также является узлом, с которым могут ассоциироваться данные. Без поправки начало отрывается от ветки. Поправить ситуацию можно всегда добавляя в конец пути лексикографически большую строку, чем возможные продолжения имени. Наиболее экономный вариант поправки в первом варианте — использование в качестве разделителя большего байта, чем возможные символы продолжения имени, а во втором варианте — такое изменение

кодировки, при котором заглавные латинские буквы займут позиции в конце кодовой таблицы.

Возможны и иные варианты. Постепенно выявляются ограничения, требующие учёта при выборе варианта. Одно из них — необходимость использовать буквы национальных алфавитов в идентификаторах с соблюдением алфавитного порядка [11].

Вопрос об обоснованном выделении схемы семантической идентификации узлов дерева требует таким образом отдельного исследования.

2.2. Уровни версионирования

Иерархия объектов не является единственной основой ветвления данных. Возможны дополнительные ветвления, связанные с различными версиями данных в системе

2.2.1. Административные версии системы

При анализе инцидентов и подготовке обновлений требуется проведение экспериментов на копии системы или её распределённой части. Многие такие эксперименты могут проводиться на имеющейся физической основе на виртуальной копии системы. При этом принципиальное отсутствие возможности изменить старые версии открывает безопасный путь создания виртуальной копии без фактического копирования данных. Речь идёт об одновременном поддержании нескольких версий состояний системы, отличающихся поздними изменениями.

2.2.2. Управляемое пользователями версионирование

Объект может содержать информацию, подлежащую контролю версий, например, систему управления версиями исходного кода исполняемых модулей. В этом случае разработчики выделяют основную и иные версии, которые могут иметь значительные общие части, но изменяться изолированно. Так же, как система в целом, объект сначала появляется и изменяется, а часто уже после возникает необходимость выделения веток и версий.

Универсальным решением было бы лексикографически разделить множество идентификаторов составляющих объекта и множества версий этого объекта. Если используется строка-разделитель '/', то её достаточно продублировать. Например, `name//01.01.34/a` это версия 01.01.34 части `name/a/`, а `//test/` это тестовая версия системы в целом.

2.2.3. Автоматическое версионирование

При обработке транзакций и других сложных вычислениях возникают версии данных, контролируемые не людьми, а вычислительными процессами. Они могут ветвиться без сознательного участия пользователей[5]. Во избежание ошибок их надо изолировать от идентификаторов, формируемых прикладными программистами. Это обеспечивается формированием идентификатора в виде конкатенации с использованием символа-разделителя, не входящего в диапазон символов используем в идентификаторах объектов.

2.3. Схемы версионирования

В зависимости от потребностей организаторов маркировка версий может отличаться[10]. При этом задача семантической идентификации узлов дерева версий по сути ничем не отличается от аналогичной задачи для дерева иерархии данных.

Единственный подводный камень, который здесь просматривается, — это общепринятая нелексикографическая маркировка версий одного уровня. К примеру, версия 10 лексикографически меньше, чем версия 2. Общепринятый путь решения этой проблемы состоит в приписывании достаточного числа нулей, например, 010 уже больше, чем 002.

Этот путь рано или поздно заведёт в тупик исчерпания доступных номеров версий. В этой ситуации обычно в пожарном порядке меняют схему именования версий вместо того, чтобы использовать простой выход из этого тупика — добавление разряда вместе с буквой. Например, a1001 лексикографически больше, чем 999.

Комбинирование предлагаемых конструкций для идентификаторов даёт общий подход к реализации различных комбинированных вариантов версионирования на основе В+ деревьев, обобщающий [9]. Например, В+ дерево, ключи которого являются конкатенацией идентификатора данного и идентификатора логической версии времени с разделителем, выходящим за диапазон используемых в идентификаторах символов, практически эквивалентно предложенному в [9] ВТ-дереву. Более сложный пример даёт использование в качестве ключей конкатенации с описанными разделителями идентификатора версии системы, идентификатора кода приложения, идентификатора версии кода, идентификатора файла кода и штампа времени.

3. Направления дальнейших исследований

К сожалению, проблему идентификация ретроспективных данных удалось охарактеризовать лишь в общих чертах. Для успешной реализации прототипа системы необходимо проработать

- систему идентификации служебной информации, ссылок и пометок об удалении,
- систему идентификации элементов одномерных и разреженных многомерных и ассоциативных массивов данных,
- кодировку, удобную для реализации системы,
- формальные описания алгоритмов сохранения и чтения информации.

По итогам предварительной проработки должен быть создан и опубликован проект стандарта.

4. Выводы

- (1) Выделены особенности ретроспективного подхода, перспективные для использования в будущих масштабируемых долгоживущих в сложно изменяющихся условиях системах.
- (2) Описана система элементарных запросов к ретроспективному хранилищу.
- (3) Описаны источники ветвления информационных структур.
- (4) Сформулированы рекомендации по выбору системы идентификации объектов ретроспективного хранилища.

Список литературы

- [1] Miles S., Groth P., Munroe S., Moreau L. *PrIME: A Methodology for Developing Provenance-Aware Applications* // ACM Transactions on Software Engineering and Methodology, 2009, no. 3, p. 1-42 ↑
- [2] Kieseberg P., Schrittwieser S., Mulazzani M., Huber M., Weippl E. *Trees Cannot Lie: Using Data Structures for Forensics Purposes* // Intelligence and Security Informatics Conference (EISIC), 2011, p. 282–285 ↑
- [3] Nierstrasz O., Denker M., Girba T., Lienhard A., Röthlisberger D. *Change-Enabled Software Systems* // Challenges for Software-Intensive Systems and New Computing Paradigms. LNCS 5380 : Springer Verlag, 2008, p. 64–79 ↑
- [4] Знаменский С. В. *Ретроспективная основа совместной реорганизации сложных информационных ресурсов* // Электронные библиотеки: перспективные методы и технологии, электронные коллекции // Труды XIII Всероссийской научной конференции RCDL-2011. — Воронеж : Воронежский госуниверситет, 2011, с. 93–101 ↑
- [5] Landes T. *Tree Clocks: An Efficient and Entirely Dynamic Logical Time System* // Proceedings of the IASTED International Conference on Parallel and Distributed Computing and Networks / ed. Burkhart H. — Innsbruck, Austria : ACTA Press, 2007, p. 375–380 ↑2.2.3

- [6] Lamport L. *Time, Clocks, and the Ordering of Events in a Distributed System*, 1978. Vol. **21**, no. 4, p. 558–565 ↑[1](#)
- [7] Torres-Rojas F. J., Ahamad M., Raynal M. *Timed consistency for shared distributed objects* // Proceedings of the Eighteenth Annual ACM Symposium on Principles of Distributed Computing // PODC '99. — New York : ACM, 1999, p. 163–172 ↑
- [8] Préz J., Arenas M., Gutierrez C. *nSPARQL: A navigational language for RDF.e* // Web Semantics: Science, Services and Agents on the World Wide Web, 2009, p. 55–270 ↑[1.1](#)
- [9] Jiang L., Salzberg B., Lomet D., Barrena M. *The BT-Tree: A Branched and Temporal Access Method* // VLDB '00 Proceedings, 2000, p. 451–460 ↑[2.3](#)
- [10] Zhu N. *Data Versioning Systems*, Stony Brook University, 2003 ↑[2.3](#)
- [11] Знаменский С. В. *Процессный подход к эволюционированию информационной системы. Ретроспективное индексирование* // Программные системы: теория и приложения, 2011. Т. **2**, № 4(8), с. 115–125 ↑[2.1](#)

Рекомендовал к публикации

д.ф.-м.н. Н. Н. Непейвода

Об авторе:



Сергей Витальевич Знаменский

Автор критерия разрешимости уравнений свёртки в пространстве функций, голоморфных на множестве, понятий выпуклости в направлении и $\mathbb{C}$ -выпуклости, русификации $\text{\TeX}$ для журналов Отделения математики РАН, графического пакета `mfpic3d` и нового подхода к построению информационных систем.

e-mail:

svz@latex.pereslavl.ru

Образец ссылки на эту публикацию:

С. В. Знаменский. Глобальная идентификация данных в долговременной перспективе // Программные системы: теория и приложения : электрон. научн. журн. 2012. Т. 3, № 2(11), с. 77–88.

URL:

http://psta.psiras.ru/read/psta2012_2_77-88.pdf

S. V. Znamenskiy. *Global data identification in a long term perspective.*

ABSTRACT. The retrospective approach to data storage is characterized by high data authenticity, availability and survivability of history. This particular features are promising for use in future scalable long-lived systems for difficult changing conditions.

The basic requests to the retrospective storage and possible sources of branching for the information structures are described. Recommendations are given for the choice of a global object identification system for a retrospective storage. (in Russian)

Key Words and Phrases: versioning, data storage, B+tree structure, timestamp, changes history, long term information system.