

Думбов Данил Сергеевич, бакалавр
Уфимский университет науки и технологий
Dumbov Danil Sergeevich
Ufa University of Science and Technology

АРХИТЕКТУРНЫЕ ПОДХОДЫ К АСИНХРОННОМУ ВЗАИМОДЕЙСТВИЮ КОМПОНЕНТОВ РАСПРЕДЕЛЁННЫХ ИНФОРМАЦИОННЫХ СИСТЕМ ARCHITECTURAL APPROACHES TO ASYNCHRONOUS INTERACTION BETWEEN COMPONENTS OF DISTRIBUTED INFORMATION SYSTEMS

Аннотация. Рассмотрены архитектурные подходы к асинхронному взаимодействию компонентов распределённых информационных систем. Проанализированы очереди сообщений, publish/subscribe, журнальные модели, подтверждения обработки и гарантии доставки. Показана роль идемпотентности, повторных попыток, dead-letter очередей, партиционирования и обратного давления. Сопоставлены свойства брокерных систем на примере AMQP/RabbitMQ, Apache Kafka и NATS.

Abstract. Architectural approaches to asynchronous interaction between components of distributed information systems are reviewed. Message queues, publish/subscribe, log-based models, acknowledgements, and delivery guarantees are analyzed. The roles of idempotency, retries, dead-letter queues, partitioning, and backpressure are discussed. Broker properties are compared using AMQP/RabbitMQ, Apache Kafka, and NATS as examples.

Ключевые слова: Распределённые системы, асинхронное взаимодействие, очередь сообщений, брокер сообщений, Kafka, RabbitMQ, NATS.

Keywords: Distributed systems, asynchronous interaction, message queue, message broker, Kafka, RabbitMQ, NATS.

Введение

При прямом синхронном взаимодействии сервис А отправляет запрос сервису В и ожидает ответ. Такая схема проста, но связывает доступность и задержку компонентов: отказ В немедленно влияет на А, а всплеск нагрузки передаётся по цепочке вызовов. В распределённых системах альтернативой является асинхронная передача сообщений через промежуточный слой. Производитель фиксирует событие или задание, после чего продолжает работу, а потребитель обрабатывает сообщение независимо по времени.

Асинхронность не устраняет распределённые ошибки, а меняет место их обработки. Возникают вопросы хранения сообщения, подтверждений, повторной доставки, порядка и контроля очереди. Различные брокеры реализуют эти свойства по-разному: AMQP стандартизует модель надёжного обмена сообщениями [2], Kafka использует распределённый журнал [1], а NATS сочетает лёгкий publish/subscribe с дополнительным постоянным слоем JetStream [5, 6]. Цель работы – систематизировать основные архитектурные решения и критерии выбора.

Очереди сообщений и publish/subscribe

В модели очереди производитель помещает сообщение в канал, а один из потребителей выполняет связанную с ним работу. Это удобно для фоновых задач, отправки уведомлений, обработки файлов и других операций, которые не обязаны завершаться в рамках исходного HTTP-запроса. Очередь сглаживает кратковременный всплеск: входная скорость может временно превышать скорость обработки, если брокер способен сохранить накопившиеся сообщения.

Publish/subscribe решает другую задачу. Производитель публикует событие, не зная конкретных получателей, а заинтересованные подписчики получают свои копии. Такая схема снижает связанность и позволяет независимо добавлять аналитику, аудит или новые бизнес-процессы. Core NATS, например, использует subject-based publish/subscribe и предоставляет многие-ко-многим взаимодействие без обязательного постоянного хранения [5].



Журнальная модель Kafka сохраняет упорядоченную последовательность записей внутри partition. Потребитель хранит позицию чтения и может повторно обработать данные. Партиционирование одновременно распределяет нагрузку и задаёт область гарантированного порядка: сообщения упорядочены внутри partition, но не образуют единого глобального порядка между всеми partitions [1, 7]. Такая модель удобна для потоков событий и повторного воспроизведения истории.

Гарантии доставки и идемпотентность

Обычно различают at-most-once, at-least-once и exactly-once семантики. При at-most-once сообщение может быть потеряно, но не повторяется. At-least-once стремится исключить потерю ценой возможных дублей. Exactly-once требует более сложной координации состояния брокера и приложения и всегда должно рассматриваться вместе с конкретными границами гарантии [7].

В RabbitMQ подтверждения потребителя и publisher confirms используются для передачи ответственности за сообщение между участниками. Если подтверждение не пришло, сообщение может быть повторно передано; документация прямо указывает, что при at-least-once потребитель должен быть готов к redelivery и проектироваться идемпотентно [3]. Идемпотентная обработка означает, что повтор одного логического сообщения не создаёт повторный побочный эффект, например второй платёж или вторую запись заказа.

Практически идемпотентность реализуется с помощью уникального идентификатора события, ограничения уникальности в базе данных, таблицы обработанных сообщений или операции upsert. В NATS JetStream подтверждаемая доставка также допускает повторение при отдельных отказах, а exactly-once механизм дополняется идентификаторами публикаций и подтверждением подтверждения [6]. Следовательно, надёжность определяется совместно брокером, протоколом и прикладной логикой.

Повторные попытки и обработка ошибок

Автоматический retry полезен для временных ошибок: кратковременной недоступности базы, сетевого сбоя или ограничения внешнего API. Однако немедленное бесконечное повторение способно создать положительную обратную связь и усилить перегрузку. Поэтому интервалы повторов обычно увеличиваются, а их количество ограничивается. Для ошибок, которые не исчезнут со временем, сообщение необходимо отделить от основного потока.

Dead-letter механизм перемещает проблемные сообщения в отдельный канал после отказа, истечения срока или превышения лимита попыток. RabbitMQ поддерживает dead-letter exchanges, куда могут направляться отклонённые или просроченные сообщения [4]. Отдельный поток ошибок позволяет основной очереди продолжать работу, а оператору – анализировать причину и при необходимости повторно публиковать исправленные сообщения.

Повторная доставка должна учитывать порядок. Если несколько событий изменяют один объект, обработка более нового события до более старого способна нарушить состояние. Решением служит последовательное потребление по ключу, версия объекта или проверка номера события. В Kafka ключ может использоваться для отправки связанных сообщений в одну partition, сохраняя их локальный порядок [7].

Масштабирование и обратное давление

Пусть λ – средняя скорость поступления сообщений, а μ – суммарная скорость обработки потребителей. При $\lambda > \mu$ размер очереди растёт примерно пропорционально разности $\lambda - \mu$. Брокер делает такой дисбаланс наблюдаемым и позволяет временно накопить данные, но не отменяет ограничение производительности. Если перегрузка длительная, требуется увеличить число потребителей, уменьшить входной поток или оптимизировать обработку.



Горизонтальное масштабирование зависит от модели брокера. В рабочих очередях сообщения распределяются между экземплярами потребителя. Kafka масштабирует consumer group через partitions, поэтому количество активно обрабатывающих потребителей ограничено числом partitions. NATS поддерживает queue groups для распределения сообщений между подписчиками, а JetStream добавляет постоянные consumers и управление подтверждениями [5, 6].

Необходимо также ограничивать число сообщений, одновременно переданных одному потребителю. Иначе быстрый брокер может переполнить память медленного приложения. Prefetch, pull-потребление, лимит in-flight сообщений и flow control реализуют формы обратного давления. Такие ограничения обычно уменьшают пиковую пропускную способность, но делают задержку и потребление памяти предсказуемее.

Выбор архитектуры взаимодействия

Очередь заданий подходит, когда требуется выполнить работу одним из группы взаимозаменяемых workers. Publish/subscribe полезен для доменных событий, которые должны независимо обрабатываться несколькими подсистемами. Распределённый журнал оправдан, когда важны длительное хранение, высокий поток, повторное чтение и последовательность событий по ключу. Один продукт может поддерживать несколько паттернов, но архитектуру следует выбирать по семантике задачи, а не по популярности брокера.

RabbitMQ и AMQP удобны для маршрутизации, очередей, подтверждений и классических схем доставки. Kafka естественно подходит для журналов событий и потоковой обработки. Core NATS минимизирует накладные расходы для лёгкого pub/sub, а JetStream добавляет хранение и более сильные гарантии. В любом случае критичными остаются схема сообщения, идентификатор, версия контракта, тайм-ауты, getty и наблюдаемость.

Мониторинг должен включать глубину очереди, возраст старейшего сообщения, скорость публикации и потребления, число повторных доставок и ошибок. Эти показатели показывают не только состояние брокера, но и дисбаланс между частями системы. Асинхронный канал полезен тогда, когда его состояние наблюдаемо и существует понятная политика обработки накопившихся и ошибочных сообщений.

Заключение

Асинхронное взаимодействие снижает временную связанность компонентов и позволяет сглаживать кратковременные пики нагрузки, но требует явного проектирования семантики доставки. Очереди, publish/subscribe и журнальные системы решают разные задачи, поэтому их свойства необходимо сопоставлять с требованиями к хранению, порядку, масштабированию и повторной обработке.

Надёжная архитектура строится не только на возможностях брокера. Она включает идемпотентных потребителей, подтверждения, ограниченные повторы, dead-letter канал, контроль in-flight нагрузки и мониторинг отставания. При таком подходе асинхронный обмен становится средством изоляции отказов и масштабирования, а не дополнительным непрозрачным слоем между сервисами

Список литературы:

1. Kreps J., Narkhede N., Rao J. Kafka: A Distributed Messaging System for Log Processing // Proceedings of NetDB. 2011.
2. OASIS. Advanced Message Queuing Protocol (AMQP) Version 1.0. OASIS Standard. 2012.
3. RabbitMQ Documentation. Reliability Guide // RabbitMQ Docs.
4. RabbitMQ Documentation. Dead Letter Exchanges // RabbitMQ Docs.
5. NATS Documentation. Core NATS // NATS Docs.
6. NATS Documentation. JetStream // NATS Docs.
7. Apache Software Foundation. Apache Kafka Design: Message Delivery Semantics and Partitioning // Apache Kafka Documentation

