



Schémas d'architecture

Composition de l'API parallèle dans AWS



Composition de l'API parallèle dans AWS: Schémas d'architecture

Copyright © 2026 Amazon Web Services, Inc. and/or its affiliates. All rights reserved.

Les marques et la présentation commerciale d'Amazon ne peuvent être utilisées en relation avec un produit ou un service qui n'est pas d'Amazon, d'une manière susceptible de créer une confusion parmi les clients, ou d'une manière qui dénigre ou discrédite Amazon. Toutes les autres marques commerciales qui ne sont pas la propriété d'Amazon appartiennent à leurs propriétaires respectifs, qui peuvent ou non être affiliés ou connectés à Amazon, ou sponsorisés par Amazon.

Table of Contents

Composition de l'API parallèle i

Composition de l'API parallèle dans AWS 1

Suggestions de lecture 2

Historique du diagramme 2

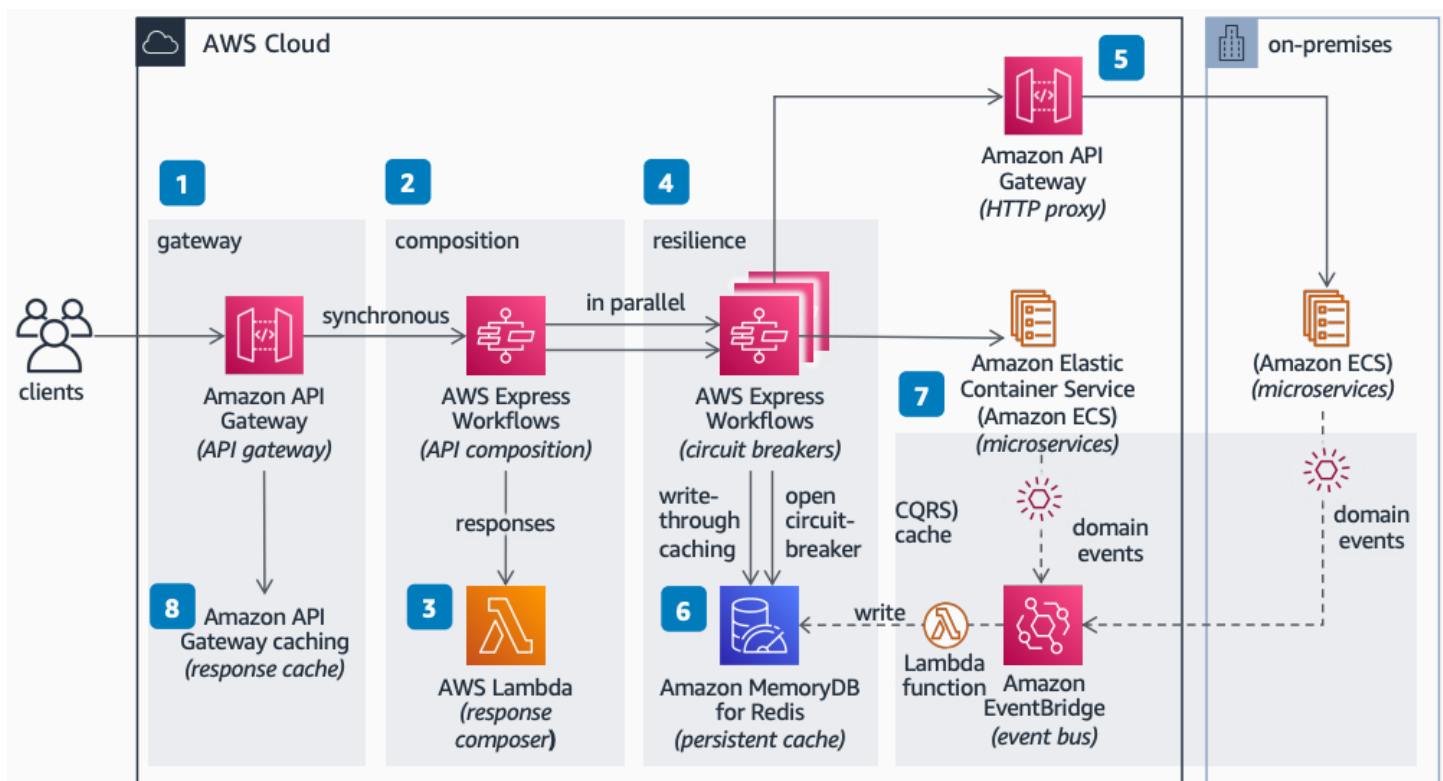
..... iv

Composition de l'API parallèle dans AWS

Date de publication : 17 mai 2022 ([Historique du diagramme](#))

Cette architecture montre comment appeler plusieurs points de terminaison d'API en aval, en parallèle ou en séquence, pour composer une seule réponse agrégée. Vous créez un flux de travail de disjoncteur générique et paramétré à l'aide d'[AWS Step Functions](#) Express Workflows et vous utilisez la séparation des responsabilités des requêtes de commande (CQRS) pour maintenir un cache persistant et éventuellement cohérent.

Composition de l'API parallèle dans AWS



Les étapes suivantes décrivent l'architecture :

1. Créez une passerelle API à l'aide d'[Amazon API Gateway](#) pour permettre aux clients d'envoyer des requêtes Web synchrones à vos microservices.
2. Utilisez Step Functions Express Workflows pour créer un flux de travail comportant des étapes parallèles qui invoquent plusieurs microservices en même temps.
3. Utilisez une [AWS Lambda](#) fonction pour composer les multiples réponses des microservices en aval en une seule réponse agrégée pour les clients.

4. Gérez les demandes parallèles à l'aide de disjoncteurs conçus avec des flux de travail express Step Functions imbriqués et paramétrés pour améliorer la résilience.
5. Utilisez API Gateway pour envoyer des requêtes HTTP par proxy lorsque vous invoquez des microservices locaux.
6. Créez une base de données Amazon MemoryDB pour Redis afin de mettre en cache les réponses des microservices à l'aide du modèle de mise en cache en écriture directe. Lorsque'un microservice en aval devient indisponible ou que sa latence atteint un seuil, le circuit se ferme et lit toutes les réponses directement depuis le cache.
7. Complétez les données mises en cache avec des événements de domaine provenant de microservices en aval à l'aide du modèle CQRS. Créez un bus d'événements avec [Amazon EventBridge](#), puis gérez les événements à l'aide d'une fonction Lambda pour conserver les agrégats de domaines dans MemoryDB.
8. Activez le cache des réponses d'API Gateway et ajustez les valeurs de durée de vie (TTL) et les filtres en fonction de chaque type de demande afin d'améliorer les performances.

Suggestions de lecture

Pour plus d'informations, consultez les ressources suivantes :

- [AWS Icônes de l'architecture](#)
- [AWS Centre d'architecture](#)
- [AWS Well-Architected](#)

Historique du diagramme

Pour être informé des mises à jour de ce schéma d'architecture de référence, abonnez-vous au fil RSS.

Modification	Description	Date
Publication initiale	Schéma d'architecture de référence publié pour la première fois.	17 mai 2022

 Note

Pour vous abonner aux mises à jour RSS, un plug-in RSS doit être activé pour le navigateur que vous utilisez.

Les traductions sont fournies par des outils de traduction automatique. En cas de conflit entre le contenu d'une traduction et celui de la version originale en anglais, la version anglaise prévaudra.