



Sfruttare il valore dei dati su larga scala nel settore dei servizi finanziari



: Sfruttare il valore dei dati su larga scala nel settore dei servizi finanziari

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

I marchi e l'immagine commerciale di Amazon non possono essere utilizzati in relazione a prodotti o servizi che non siano di Amazon, in una qualsiasi modalità che possa causare confusione tra i clienti o in una qualsiasi modalità che denigri o discrediti Amazon. Tutti gli altri marchi non di proprietà di Amazon sono di proprietà dei rispettivi proprietari, che possono o meno essere affiliati, collegati o sponsorizzati da Amazon.

Table of Contents

Introduzione	1
Obiettivi aziendali specifici	2
Parametri operativi	4
Panoramica della soluzione	7
Un framework ML scalabile	7
Un hub centrale per i metadati	9
SageMaker convalida	10
Best practice	11
Una visione di successo	13
Domande frequenti	14
Fasi successive	16
Resources	17
Cronologia dei documenti	18
.....	xix

Sfruttare il valore dei dati su larga scala nel settore dei servizi finanziari

Brian Cavanagh, Maira Ladeira Tanke, Amine Ait el harraj, Junaid Baba, Maren Suilmann, Pauline Ting e Sokratis Kartakis, Amazon Web Services (AWS)

Settembre 2022 ([cronologia dei documenti](#))

Il settore dei servizi finanziari (FS) sta affrontando gravi sconvolgimenti a causa delle società di tecnologia finanziaria (fintech) e delle banche esclusivamente digitali che stanno aprendo la strada alla trasformazione digitale del settore bancario. Questa trasformazione è sempre più caratterizzata dallo sviluppo di prodotti e servizi bancari basati su tecnologie di intelligenza artificiale e apprendimento automatico (AI/ML). Secondo le [prospettive AI-bank del futuro: le banche possono affrontare la sfida dell'IA?](#) from McKinsey & Company, le organizzazioni FS devono implementare AI/ML tecnologie su larga scala se vogliono rimanere pertinenti. Per implementare AI/ML tecnologie su larga scala, le organizzazioni FS devono innanzitutto sbloccare il valore aziendale dei propri dati.

Le organizzazioni FS detengono grandi quantità di dati, ma hanno difficoltà a trarre valore da questi dati su larga scala. Sfruttando il valore dei dati, le organizzazioni FS possono generare offerte, approfondimenti e supporto più approfonditi e personalizzati ai propri clienti e partner. Il valore sbloccato può anche aiutare le organizzazioni FS a smascherare rapidamente le inefficienze nei processi correnti, dai mercati dei capitali alle operazioni di produzione, fornendo al contempo informazioni prioritarie sulle aree che richiedono l'ottimizzazione. Questa strategia descrive come sbloccare il valore dei dati su larga scala sviluppando funzionalità di machine learning nella propria organizzazione. Il pubblico a cui si rivolge questa strategia include CEO, CFO, CIO e dirigenti senior del settore bancario e della gestione patrimoniale.

Questa strategia aiuta a comprendere quanto segue:

- Risultati aziendali derivanti dal lancio di funzionalità di machine learning nella tua organizzazione
- Parametri e punteggi target per il successo operativo
- Un framework di ML scalabile per trasformare le funzionalità ML della tua organizzazione
- AWS migliori pratiche per la scalabilità (basate su centinaia di implementazioni presso i clienti)

Obiettivi aziendali specifici

Il principale risultato aziendale è l'adozione di processi ottimizzati in grado di trasformare i dati dell'organizzazione in valore aziendale. Nello specifico, questa strategia può aiutarti a raggiungere i seguenti obiettivi aziendali specifici:

- Dimensiona le funzionalità di machine learning e operazioni (MLOP) in tutta l'organizzazione e sviluppa una piattaforma di ML per supportare centinaia di data scientist e data engineer nell'esecuzione dei flussi di lavoro ML nella tua organizzazione.
- Implementa un'implementazione di infrastruttura scalabile, sicura, economica e sostenibile utilizzando AWS Service Catalog per creare un'infrastruttura gestita su richiesta con Amazon AI. SageMaker
- Standardizza il processo di sviluppo e implementazione del modello di ML tra più team.
- Riduci il debito tecnico dei modelli esistenti e crea artefatti riutilizzabili per accelerare lo sviluppo di modelli futuri.
- Riduci il tempo di conversione dall'idea al valore per i casi d'uso di dati e analisi fino a 12-16 settimane.
- Riduci il tempo di creazione per gli ambienti dei casi d'uso del ML a 1-2 giorni.

Quando si dimensiona l'uso dell'analisi avanzata in azienda, lo sviluppo e il rilascio di modelli e soluzioni ML in produzione possono richiedere troppo tempo. Collaborando con AWS, riceverai indicazioni e supporto che possono aiutarti a progettare, creare e lanciare rapidamente una piattaforma self-service moderna, sicura, scalabile e sostenibile per lo sviluppo e la produzione ML-based di servizi a supporto della tua azienda e dei tuoi clienti. [AWS I servizi professionali](#) possono collaborare con te per accelerare l'implementazione delle best practice di AWS per l'utilizzo dell'[SageMaker intelligenza artificiale](#) per creare, addestrare e distribuire modelli di machine learning per casi d'uso personalizzati in base alle tue esigenze organizzative.

La collaborazione prevede quanto segue:

- Un DevOps-driven approccio federato e self-service per l'infrastruttura e il codice applicativo, con un percorso chiaro che può ridurre i tempi di implementazione da settimane a minuti
- Un ambiente sicuro, controllato e basato su modelli per accelerare l'innovazione con modelli e approfondimenti di machine learning che utilizzano le best practice del settore e artefatti condivisi a livello bancario

- La capacità di accedere e condividere i dati in modo più semplice e coerente all'interno dell'organizzazione
- Un set di strumenti moderno basato su un'architettura gestita su richiesta in grado di ridurre al minimo i requisiti di elaborazione, ridurre i costi e consentire lo sviluppo e le operazioni di machine learning sostenibili, con la flessibilità necessaria per integrare nuovi AWS prodotti e servizi per soddisfare i continui casi d'uso e i requisiti di conformità
- Supporto all'adozione, al coinvolgimento e alla formazione che consente ai team di data science e data engineering di essere autosufficienti e di dimensionare i prodotti di ML in tutta l'organizzazione

Parametri operativi

AWS I Servizi professionali possono collaborare con i team interni per stabilire un meccanismo di governance della distribuzione per tracciare i progressi dell'organizzazione rispetto alle metriche operative definite nella tabella seguente. I parametri sono basati su valori reali e possono aiutarti a definire il tuo coinvolgimento MLOps. I valori nella tabella seguente sono solo linee guida. I valori target dipenderanno dall'implementazione dell'organizzazione.

Area	Sottoarea	Parametro	Obiettivo di 6 mesi dal lancio della piattaforma
1. Efficienza operativa	Distribuzione più rapida	Time to value	3 mesi in media
	Trasparenza dei costi attraverso una chiara allocazione dei costi alle divisioni aziendali	Costi di gestione AWS	Inferiore a quelli iniziali
2. Boundary-less consegna	Libertà di accesso garantita senza soluzione di continuità in modo controllato	Time to access	1 giorno
	Support per CI/CD	Time to value (idea to live)	3 mesi in media
	Distribuzione route-to-live semplificata	Durata della distribuzione (dopo la fase beta)	3 mesi in media
	On-demand rotazione in laboratorio up/down	Tempo di avvio dell'ambiente	24 ore
3. Controllato	Accesso allineato a ciascun caso d'uso e	Caso d'uso relativo al momento dell'accesso	< 1 giorno

	concesso o rimosso dinamicamente		
	Ambiente di produzione e persistente che consente lo sviluppo basato su spoke	Tempo di implementazione	2 settimane
	Costi di archiviazione ed elaborazione gestiti nel punto di utilizzo e attribuiti al centro di costo	Actuals/forecast trasparente	Varianza della precisione delle previsioni consentita
	Conformità agli standard di sicurezza interni (valutazione e controllo automatici)	Numero di ambienti non conformi	0
	Proprietari dei dati responsabili del contributo e della condivisione dei dati (nell'ambito dello SLA)	Copie dei dati nel data lake	0 copie dei dati nel data lake
	Elaborazione comune con supporto per la separazione normativa dei dati	Accesso conforme	0 risultati normativi
4. Schemi ripetibili e in continua evoluzione	Un ambiente di sviluppo di modelli con standard in continua evoluzione	Numero di aggiornamenti all'anno	6

5. Per tutta la banca	Capacità di impacchettare e implementare ambienti a piacimento utilizzando gli spoke	Tempo di avvio dell'ambiente	2 ore
	Layout chiaro dei servizi e dell'infrastruttura utilizzati	Non-tracked infrastruttura	0
	Numero ridotto di ambienti di modellazione distinti	Numero di ambienti	< 2
	Collaborazione semplificata in team attraverso i confini organizzativi	Tempo di accesso (ad esempio, dati)	1 giorno
	Funzionamento comune e singolare model/framework	Numero di domini controllati dal modello	> 3

Panoramica della soluzione

Un framework ML scalabile

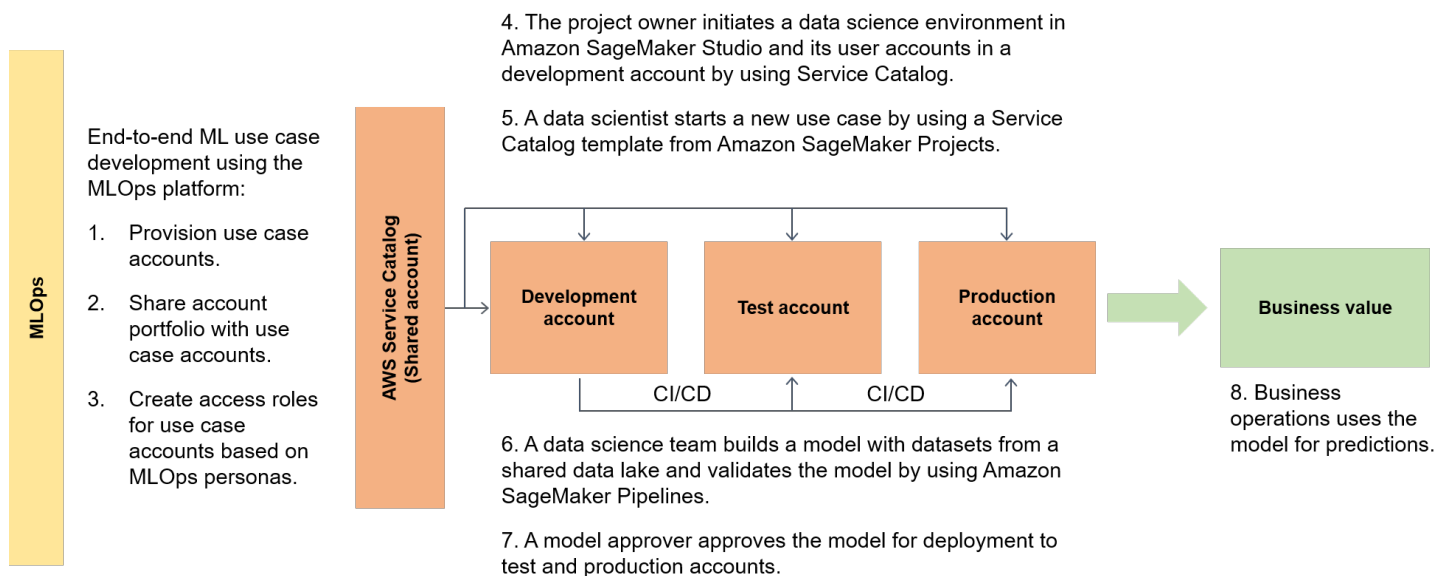
In un'azienda con milioni di clienti distribuiti su più linee di business, i flussi di lavoro ML richiedono l'integrazione di dati posseduti e gestiti da team isolati che utilizzano strumenti diversi per sbloccare il valore aziendale. Le banche si impegnano a proteggere i dati dei propri clienti. Allo stesso modo, anche l'infrastruttura utilizzata per lo sviluppo di modelli ML è soggetta a elevati standard di sicurezza. Questa sicurezza aggiuntiva aggiunge ulteriore complessità e influisce sul time to value dei nuovi modelli ML. In un framework ML scalabile, è possibile utilizzare un set di strumenti modernizzato e standardizzato per ridurre lo sforzo necessario per combinare diversi strumenti e semplificare il processo di route-to-live per i nuovi modelli ML.

Normalmente, la gestione e il supporto delle attività di data science nel settore FS sono controllati da un team di piattaforma centrale che raccoglie i requisiti, fornisce risorse e mantiene l'infrastruttura per i team di dati in tutta l'organizzazione. Per dimensionare rapidamente l'uso del machine learning nei team federati dell'organizzazione, puoi utilizzare un framework ML scalabile per fornire funzionalità self-service agli sviluppatori di nuovi modelli e pipeline. Ciò consente a questi sviluppatori di implementare un'infrastruttura moderna, preapprovata, standardizzata e sicura. In definitiva, queste funzionalità self-service riducono la dipendenza dell'organizzazione dai team di piattaforma centralizzati e velocizzano il time-to-value per lo sviluppo di modelli di machine learning.

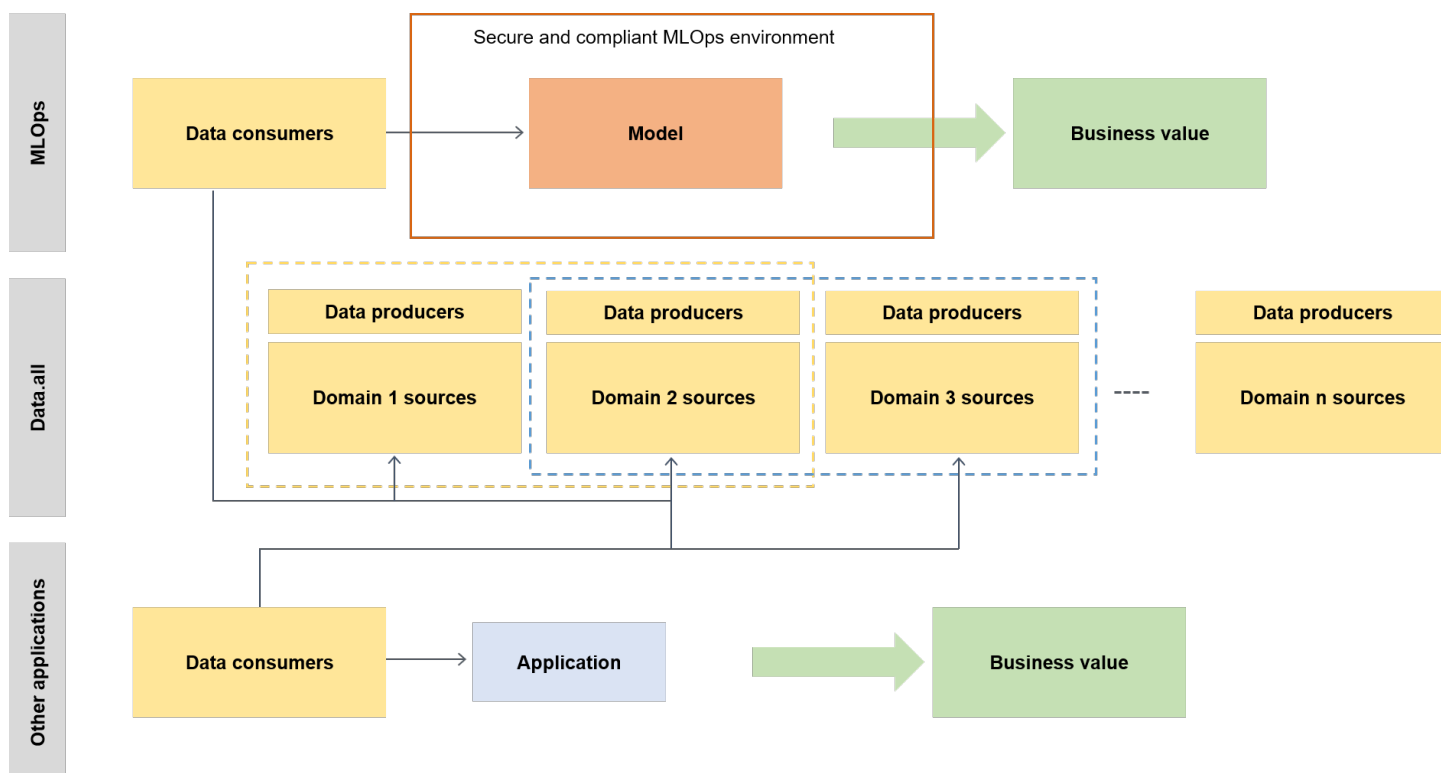
Il framework ML scalabile consente ai consumatori di dati (ad esempio, data scientist o ingegneri di machine learning) di sbloccare il valore aziendale offrendo loro la possibilità di fare quanto segue:

- Esplorare e scoprire i dati preapprovati necessari per la formazione dei modelli
- Accedere ai dati preapprovati in modo rapido e semplice
- Utilizzare dati preapprovati per dimostrare la fattibilità del modello
- Mettere in produzione il modello collaudato per consentirne l'utilizzo da parte di altri

Il diagramma seguente evidenzia il flusso completo del framework e il percorso semplificato di implementazione per i casi d'uso del machine learning.



In un contesto più ampio, i consumatori di dati utilizzano un acceleratore serverless chiamato data.all per reperire dati su più data lake e quindi utilizzare i dati per addestrare i propri modelli, come illustrato nel diagramma seguente.



A un livello inferiore, il framework ML scalabile contiene quanto segue:

- Self-service implementazione dell'infrastruttura: riduci la dipendenza dai team centralizzati.

- Sistema centrale di gestione dei pacchetti Python: rendi disponibili pacchetti Python preapprovati per lo sviluppo di modelli.
- CI/CD pipeline per lo sviluppo e la promozione di modelli: riduci il time to live includendo pipeline di integrazione continua e pipeline continue (CI/CD) come parte dei modelli Infrastructure as Code (IaC).
- Funzionalità di test dei modelli: sfrutta le funzionalità di test unitari, test dei modelli, test di integrazione e test completi che sono automaticamente disponibili per i nuovi modelli.
- Disaccoppiamento e orchestrazione dei modelli: [evita il calcolo non necessario e rendi le distribuzioni più solide disaccoppiando le fasi del modello in base ai requisiti delle risorse di calcolo e orchestrazione delle diverse fasi utilizzando Amazon AI Pipelines. SageMaker](#)
- Standardizzazione del codice: migliora la qualità del codice utilizzando l'integrazione della CI/CD pipeline per la convalida degli standard [Python Enhancement](#) Proposal (PEP 8).
- Quick-start modelli ML generici: [ottieni modelli di Service Catalog che istanziano i tuoi ambienti di modellazione ML \(sviluppo, preproduzione e produzione\) e le pipeline associate con un semplice clic utilizzando SageMaker AI Projects per l'implementazione.](#)
- Monitoraggio della qualità di dati e modelli: assicurati che i tuoi modelli soddisfino i requisiti operativi e rispettino il tuo livello di tolleranza al rischio utilizzando [Amazon SageMaker AI Model Monitor per monitorare automaticamente la variazione dei dati e della qualità del modello.](#)
- Monitoraggio della polarizzazione: consenti ai proprietari dei modelli di prendere decisioni giuste ed eque controllando automaticamente gli squilibri dei dati e se i cambiamenti nel mondo hanno introdotto distorsioni nel modello.

Un hub centrale per i metadati

[Data.all](#) è un acceleratore serverless che puoi integrare con i data lake AWS esistenti per raccogliere metadati in un hub centrale. Un'interfaccia utente semplice e facile da usare in data.all mostra i metadati associati ai set di dati provenienti da più data lake esistenti. Ciò consente sia agli utenti non tecnici che a quelli tecnici di cercare, sfogliare e richiedere l'accesso a dati importanti che possono essere utilizzati nei propri laboratori di machine learning. Data.all utilizza AWS Lake Formation AWS Lambda, Amazon Elastic Container Service (Amazon ECS) AWS Fargate, OpenSearch Amazon Service e. AWS Glue

SageMaker convalida

Per dimostrare le capacità dell' SageMaker intelligenza artificiale in una vasta gamma di architetture di elaborazione dati e ML, il team che implementa le funzionalità seleziona, insieme al team dirigenziale del settore bancario, casi d'uso di varia complessità provenienti da diverse divisioni dei clienti bancari. I dati dei casi d'uso vengono offuscati e resi disponibili in un bucket di dati locale di [Amazon Simple Storage Service \(Amazon S3\)](#) nell'account di sviluppo dei casi d'uso per la fase di verifica delle funzionalità.

Una volta completata la migrazione del modello dall'ambiente di formazione originale a un'architettura di SageMaker intelligenza artificiale, il data lake ospitato nel cloud rende i dati disponibili per essere letti dai modelli di produzione. Le previsioni generate dai modelli di produzione vengono quindi riscritte nel data lake.

Dopo la migrazione dei casi d'uso candidati, il framework ML scalabile utilizza un valore di riferimento iniziale per i parametri di destinazione. Puoi confrontare il valore di riferimento con le precedenti tempistiche on premise o di altri provider di servizi cloud come prova dei miglioramenti temporali consentiti dal framework ML scalabile.

Best practice

Questa sezione fornisce una panoramica delle AWS migliori pratiche per MLOP.

Gestione e separazione degli account

[AWS le migliori pratiche](#) per la gestione degli account consigliano di dividere gli account in quattro account per ogni caso d'uso: sperimentazione, sviluppo, test e produzione. È inoltre consigliabile disporre di un account di governance per fornire risorse MLOps condivise in tutta l'organizzazione e un account data lake per fornire l'accesso centralizzato ai dati. La logica alla base di ciò è quella di separare completamente gli ambienti di sviluppo, test e produzione, evitare i ritardi causati dai limiti del servizio dovuti ai molteplici casi d'uso e ai team di data science che condividono lo stesso set di account e fornire una panoramica completa dei costi per ogni caso d'uso. Infine, è consigliabile separare i dati a livello di account, poiché ogni caso d'uso ha il proprio set di account.

Standard di sicurezza

Per soddisfare i requisiti di sicurezza, è consigliabile disattivare l'accesso pubblico a Internet e crittografare tutti i dati con chiavi personalizzate. Quindi, puoi distribuire un'istanza sicura di Amazon SageMaker AI Studio sull'account di sviluppo in pochi minuti utilizzando Service Catalog. Puoi anche ottenere funzionalità di controllo e monitoraggio dei modelli per ogni caso d'uso utilizzando l'SageMaker intelligenza artificiale tramite modelli distribuiti con SageMaker AI Projects.

Funzionalità dei casi d'uso

Una volta completata la configurazione dell'account, i data scientist della tua organizzazione possono richiedere un nuovo modello di caso d'uso utilizzando SageMaker AI Projects in SageMaker AI Studio. Questo processo implementa l'infrastruttura necessaria per disporre delle funzionalità MLOPS nell'account di sviluppo (con un supporto minimo da parte dei team centrali), come CI/CD pipeline, unit test, test dei modelli e monitoraggio dei modelli.

Ogni caso d'uso viene quindi sviluppato (o rifattorizzato nel caso di una base di codice applicativa esistente) per essere eseguito in un'architettura di intelligenza artificiale utilizzando SageMaker funzionalità di SageMaker intelligenza artificiale come il tracciamento degli esperimenti, la spiegabilità dei modelli, il rilevamento delle distorsioni e il monitoraggio della qualità. data/model [Puoi aggiungere queste funzionalità a ciascuna pipeline di casi d'uso utilizzando le fasi delle pipeline in AI Pipelines.](#) SageMaker

Percorso di maturità di MLOps

Il percorso di maturità di MLOps definisce le funzionalità MLOps necessarie rese disponibili in una configurazione a livello aziendale per garantire un flusso di lavoro basato su modelli end-to-end. Il percorso di maturità si compone di quattro fasi:

1. **Iniziale:** in questa fase, crei l'account di sperimentazione. Inoltre, ti assicuri un nuovo AWS account all'interno della tua organizzazione in cui puoi sperimentare SageMaker Studio e altri nuovi servizi. AWS
2. **Ripetibile:** in questa fase, si standardizzano gli archivi di codice e lo sviluppo di soluzioni di ML. Adotterai inoltre un approccio di implementazione multi-account e standardizzerai i tuoi repository di codice per supportare la governance e gli audit dei modelli man mano che ampli l'offerta impiegando la scalabilità orizzontale. Ti consigliamo di adottare un approccio allo sviluppo di modelli pronti per la produzione con soluzioni standard fornite da un account di governance. I dati vengono archiviati in un account di data lake e i casi d'uso vengono sviluppati in due account. Il primo account riguarda la sperimentazione durante il periodo di esplorazione della data science. In questo resoconto, i data scientist scoprono modelli per risolvere i problemi aziendali e sperimentano molteplici possibilità. L'altro aspetto riguarda lo sviluppo, che avviene dopo che il modello migliore è stato identificato e il team di data science è pronto a lavorare nella pipeline di inferenza.
3. **Affidabile:** in questa fase, introduci il test, l'implementazione e l'implementazione multi-account. È necessario comprendere i requisiti MLOps e introdurre test automatizzati. Implementa le best practice MLOPS per garantire che i modelli siano robusti e sicuri. Durante questa fase, introduci due nuovi account per i casi d'uso: un account di test per testare i modelli sviluppati in un ambiente che emula l'ambiente di produzione e un account di produzione per eseguire l'inferenza dei modelli durante le operazioni aziendali. Infine, utilizza il test, l'implementazione e il monitoraggio automatizzati dei modelli in una configurazione multi-account per garantire che i modelli soddisfino gli standard di alta qualità e prestazioni che hai impostato.
4. **Scalabile:** in questa fase, si modellano e si producono più soluzioni di ML. Diversi team e casi d'uso del ML iniziano ad adottare MLOps durante il processo di creazione del modello end-to-end. Per raggiungere la scalabilità in questa fase, è inoltre necessario aumentare il numero di modelli nella libreria di modelli grazie al contributo di una base più ampia di data scientist, ridurre il time-to-value dall'idea al modello di produzione per più team dell'organizzazione e iterare man mano che si dimensiona.

Per ulteriori informazioni sul modello di maturità MLOPS, consulta la [roadmap di base MLOPS per le imprese con Amazon SageMaker AI](#) sul Machine Learning AWS Blog.

Una visione di successo

Per sfruttare con successo il valore dei dati, l'organizzazione deve sviluppare un framework ML scalabile che consenta di raggiungere i risultati di business, allineandosi al contempo alle migliori pratiche definite in questa strategia. AWS Tale struttura potrebbe aiutare l'organizzazione a raggiungere i seguenti obiettivi:

- Un processo self-service per la creazione di un'infrastruttura standardizzata a livello di produzione per lo sviluppo di analisi avanzate, carichi di lavoro basati sulla scienza dei dati e un percorso di vita chiaro per tali carichi di lavoro DevOps-driven
- Team dei casi d'uso che possono contare su una governance semplificata, che riduce il time-to-value e allo stesso tempo libera risorse tecniche preziose
- Una libreria di data science per accedere e condividere modelli di ML riutilizzabili in grado di ridurre rapidamente i tempi di distribuzione end-to-end utilizzando modelli che possono essere perfezionati e condivisi con altri
- Una libreria di riferimento per i dati su cloud (DataHub) che fornisce metadati su tutti i dati presenti su più data lake, velocizzando non solo la scoperta AWS dei dati ma anche i tempi di accesso ai dati necessari per addestrare i modelli
- Un DevOps team dedicato formato per gestire la piattaforma e il toolkit sottostante, supportare lo sviluppo di nuove funzionalità e garantire la conformità continua ai controlli chiave
- Real-life casi d'uso forniti agli account a livello di produzione
- Un trasferimento di conoscenze tra il team di progetto e i team della linea di business ricevente (team spoke), che consente ai team aziendali di gestire autonomamente lo sviluppo continuo e il processo route-to-live
- Completamento delle attività iniziali di coinvolgimento e adozione con un gruppo mirato di team aziendali
- Centinaia di sessioni di AWS formazione in aula a supporto dell'adozione

Domande frequenti

Quando è necessario creare una piattaforma MLOps?

È giunto il momento di standardizzarsi su una piattaforma MLOps quando noti che i tuoi ingegneri dedicano più tempo alla ricerca e all'ottenimento di approvazioni per l'acquisizione degli strumenti che alla creazione di modelli di ML.

Posso integrare altri strumenti di ML nella piattaforma MLOps?

Sì. È possibile integrare strumenti diversi da AWS strumenti nella piattaforma. Sebbene SageMaker AI Studio sia il fulcro della piattaforma MLOps, puoi comunque integrare altri prodotti con la suite di servizi SageMaker AI Studio.

In che modo la mia organizzazione può semplificare i requisiti di governance per accelerare l'innovazione?

Nell'ambito dei casi d'uso candidati che selezioni per dimostrare la struttura della tua piattaforma MLOps, assicurati che i casi d'uso siano di complessità sufficiente, richiedano varie classificazioni dei dati e richiedano grandi volumi di dati. In questo modo, non solo dimostri la capacità della piattaforma, ma svolgi anche il lavoro più impegnativo dal punto di vista della governance nell'ambito della release iniziale della piattaforma. Se riuscirai a farlo, i team che adotteranno la piattaforma MLOps come parte del rollout avranno un carico di governance più leggero, in quanto utilizzano una piattaforma che ha già soddisfatto i requisiti di governance per casi d'uso complessi.

Di quale team ho bisogno per creare una piattaforma MLOps?

Una solida base MLOps, che definisce chiaramente l'interazione tra più persone e tecnologie, può aumentare il time-to-value, ridurre i costi e consentire ai data scientist di concentrarsi sull'innovazione. Avere il team giusto può fare la differenza tra fallimento e successo dello sviluppo della piattaforma MLOps. A causa della natura di MLOPS, è necessario coinvolgere molti ruoli, come data scientist, ingegneri ML, DevOps professionisti, proprietari di dati, proprietari IT, analisti aziendali e proprietari di prodotti. Assicurati che tutte le parti interessate interagiscano in un team interfunzionale per garantire i migliori risultati per la tua piattaforma MLOps.

Come posso iniziare il mio percorso MLOps?

Puoi iniziare creando un ambiente di sperimentazione sicuro in cui i data scientist ricevano uno snapshot dei dati. I data scientist possono utilizzare l' SageMaker intelligenza artificiale per

sperimentare e, in ultima analisi, dimostrare che il machine learning può risolvere un problema aziendale specifico.

Una trasformazione MLOps dovrebbe essere guidata da un approccio top-down o bottom-up in un'organizzazione?

Sebbene gli approcci bottom-up possano avere successo, il supporto della leadership è essenziale per il successo dello sviluppo della piattaforma MLOps. Con un approccio top-down, puoi garantire una standardizzazione più rapida della soluzione sviluppata, ridurre i costi e ottenere una maggiore scalabilità e riutilizzabilità tra i modelli sviluppati dai diversi team dell'organizzazione.

Fasi successive

Per sbloccare il valore aziendale dei tuoi dati adottando un framework di ML scalabile, considera i seguenti passaggi successivi:

1. [Contatta un rappresentante AWS di vendita](#) o [contatta un rappresentante dei servizi AWS professionali](#).
2. Cerca opportunità di formazione e certificazione per i membri del team appropriati della tua organizzazione:
 - [AWS formazione e certificazioni](#)
 - [MLops Engineering su AWS](#)
 - [AWS Big Data certificati - Specialità](#)
 - [AWS Machine Learning certificato - Specialità](#)
 - [AWS DevOps Ingegnere certificato - Professionista](#)
 - [AWS Formazione privata - Richiesta di informazioni](#)

Resources

- [Cosa c'è di nuovo con il Machine Learning?](#) (AWS documentazione)
- [Automatizza gli MLOP con progetti di SageMaker intelligenza artificiale](#) () YouTube
- [Metti a frutto i tuoi dati sul cloud più scalabile, affidabile e sicuro \(documentazione\)](#) AWS
- [Roadmap di base di MLOps per le aziende con Amazon SageMaker AI](#) (AWS Machine Learning Blog)
- [Parte 1: Come NatWest Group ha creato una piattaforma MLOPS scalabile, sicura e sostenibile](#) (AWS Machine Learning Blog)
- [AI-bank del futuro: le banche possono affrontare la sfida dell'IA?](#) (McKinsey e azienda)

Cronologia dei documenti

La tabella seguente descrive le modifiche significative apportate a questa guida. Per ricevere notifiche sugli aggiornamenti futuri, puoi abbonarti a un [feed RSS](#).

Modifica	Descrizione	Data
Sostituito DataHub con data.all.	—	7 novembre 2023
Pubblicazione iniziale	—	16 settembre 2022

Le traduzioni sono generate tramite traduzione automatica. In caso di conflitto tra il contenuto di una traduzione e la versione originale in Inglese, quest'ultima prevarrà.