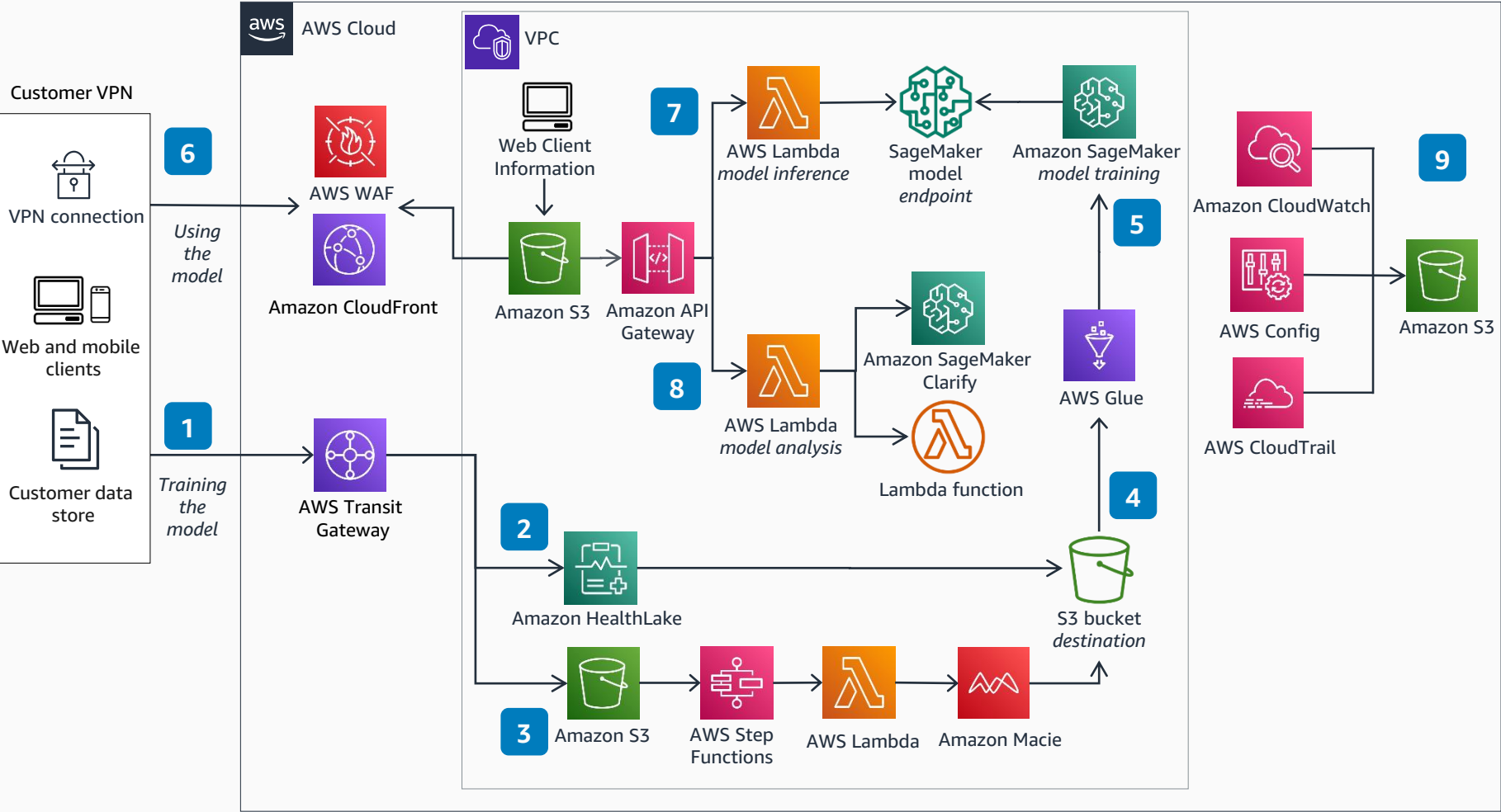


Guidance for Patient Outcome Prediction on AWS

This architecture illustrates the POP application. The application uses patient data to train ML models that predict patient outcomes.



- 1 Access POP and input patient health data such as medical records, insurance claims, lab reports, and doctor's notes through **AWS Transit Gateway**.
- 2 **Amazon HealthLake**, a Health Insurance Portability and Accountability Act (HIPAA)-eligible service, ingests customer health data. **HealthLake** transforms customer health data to make it ready for querying and ML processing.
- 3 A series of **AWS Step Functions** sends customer datasets through **Amazon Macie**, a service that discovers and protects sensitive data.
- 4 **Amazon Simple Storage Service (Amazon S3)** stores customer data from **HealthLake** and **Macie**. **AWS Glue** then processes the data by transforming and preparing it for ingestion as training data. The data is then ready for a custom-trained **Amazon SageMaker** model.
- 5 Custom **SageMaker** models are trained to predict patient outcomes, such as disease progression for potentially undiagnosed patients and hospital readmission probability. **SageMaker** model endpoints are then created for model inference.
- 6 Access the web client frontend through **Amazon CloudFront** and **AWS WAF** to perform model inference.
- 7 When you want to perform model inference from a previously trained model, an **AWS Lambda** trigger lets you pick a **SageMaker** model endpoint to perform predictions.
- 8 When you want to explore model explainability, a **Lambda** trigger lets you look at potential bias in your training data and trained models through **Amazon SageMaker Clarify**.
- 9 To support operational and cost monitoring, **Amazon CloudWatch**, **AWS Config**, and **AWS CloudTrail** record all **Amazon Virtual Private Cloud (VPC)** flow logs, API metrics, and other AWS resource usage.

