

# Google.Professional-Machine-Learning-Engineer.v2022-12-19.q66

|                                                                                                                                                                                                                       |                                               |
|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------|
| 試験コード:                                                                                                                                                                                                                | Professional-Machine-Learning-Engineer        |
| 試験名称:                                                                                                                                                                                                                 | Google Professional Machine Learning Engineer |
| 認定資格:                                                                                                                                                                                                                 | Google                                        |
| 無料問題数:                                                                                                                                                                                                                | 66                                            |
| バージョン:                                                                                                                                                                                                                | v2022-12-19                                   |
| アクセス数:                                                                                                                                                                                                                | 1030                                          |
| ページビュー数:                                                                                                                                                                                                              | 660                                           |
| <a href="https://www.jpnpdf.com/Google.Professional-Machine-Learning-Engineer.v2022-12-19.q66-mondaishu.html">https://www.jpnpdf.com/Google.Professional-Machine-Learning-Engineer.v2022-12-19.q66-mondaishu.html</a> |                                               |

## 最新問題: 1

クラウドに接続されたデバイスを提供することで健康的な睡眠パターンを促進する会社は、現在AWSで睡眠追跡アプリケーションをホストしています。アプリケーションは、デバイスユーザーからデバイス使用情報を収集します。同社のデータサイエンスチームは、ユーザーが会社のデバイスの利用を停止するかどうか、またいつ停止するかを予測するための機械学習モデルを構築しています。このモデルからの予測は、ユーザーに連絡するための最良のアプローチを決定するダウンストリームアプリケーションによって使用されます。

データサイエンスチームは、機械学習モデルの複数のバージョンを構築して、各バージョンを会社のビジネス目標に照らして評価しています。長期的な有効性を測定するために、チームは、モデルによって提供される推論の部分を制御する機能を使用して、モデルの複数のバージョンを長期間並行して実行したいと考えています。

最小限の労力でこれらの要件を満たすソリューションはどれですか？

- A. AmazonSageMakerで複数のモデルをビルドしてホストします。複数の本番バリエーションを使用してAmazonSageMakerエンドポイント構成を作成します。エンドポイント構成を更新することにより、複数のモデルによって提供される推論の部分をプログラムで制御します。
- B. AmazonSageMakerで複数のモデルをビルドしてホストします。複数のモデルにアクセスする単一のエンドポイントを作成します。Amazon SageMakerバッチ変換を使用して、単一のエンドポイントを介したさまざまなモデルの呼び出しを制御します。
- C. AmazonSageMakerで複数のモデルをビルドしてホストします。モデルごとに1つずつ、複数のAmazonSageMakerエンドポイントを作成します。アプリケーション層で推論するために、さまざまなモデルの呼び出しをプログラムで制御します。
- D. Amazon SageMaker Neoで複数のモデルを構築してホストし、さまざまな種類の医療機器を考慮に入れます。医療機器のタイプに基づいて、推論のために呼び出されるモデルをプログラムで制御します。

**Answer: B (メッセージを残す)**

**最新問題: 2**

オンラインファッション会社で働く機械学習スペシャリストは、会社のAmazonS3ベースのデータレイク用のデータ取り込みソリューションを構築したいと考えています。

スペシャリストは、以下で構成される将来の機能を可能にする一連の取り込みメカニズムを作成したいと考えています。

\*リアルタイム分析

\*履歴データのインタラクティブな分析

\*クリックストリーム分析

\*製品の推奨事項

スペシャリストはどのサービスを使用する必要がありますか？

**A.** データカタログとしてのAWS Glue; 履歴データインサイトのための

AmazonKinesisDataStreamsとAmazonKinesisData Analytics; クリックストリーム分析のためにAmazonESに配信するAmazonKinesisData Firehose; パーソナライズされた製品の推奨事項を生成するAmazonEMR

**B.** データカタログとしてのAmazon Athena; 履歴データインサイトのための

AmazonKinesisDataStreamsとAmazonKinesisData Analytics; クリックストリーム分析用のAmazonDynamoDBストリーム。パーソナライズされた製品の推奨事項を生成するAWSGlue

**C.** データカタログとしてのAmazon Athena ;ほぼリアルタイムのデータインサイトのためのAmazon KinesisDataStreamsとAmazonKinesisData Analytics; クリックストリーム分析用のAmazonKinesisData Firehose; パーソナライズされた製品の推奨事項を生成するAWSGlue

**D.** データカタログとしてのAWS Glue; リアルタイムのデータインサイトのための

AmazonKinesisDataStreamsとAmazonKinesisData Analytics; クリックストリーム分析のためにAmazonESに配信するAmazonKinesisData Firehose; パーソナライズされた製品の推奨事項を生成するAmazonEMR

**Answer: D (メッセージを残す)**

**最新問題: 3**

データサイエンスチームは、さまざまな機能、モデルアーキテクチャ、ハイパーパラメータをすばやく実験する必要があります。さまざまな実験の精度指標を追跡し、APIを使用して時間の経過とともに指標をクエリする必要があります。手作業を最小限に抑えながら、実験を追跡および報告するために何を使用する必要がありますか？

**A.** AIPlatformNotebooksを使用して実験を実行します。共有のGoogleスプレッドシートファイルに結果を収集し、GoogleスプレッドシートAPIを使用して結果をクエリします

**B.** Kubeflow Pipelinesを使用して実験を実行しますメトリックファイルをエクスポートし、KubeflowPipelinesAPIを使用して結果をクエリします。

**C.** AI Platform Trainingを使用して実験を実行します精度メトリックをBigQueryに書き込み、BigQueryAPIを使用して結果をクエリします。

D. AI Platform Trainingを使用して実験を実行します。精度メトリックをCloudMonitoringに書き込み、MonitoringAPIを使用して結果をクエリします。

**Answer: C** ([メッセージを残す](#))

#### 最新問題: 4

あなたは大規模なホテルチェーンで働いており、ターゲットを絞ったマーケティング戦略の予測を収集する際にマーケティングチームを支援するように依頼されました。マーケティングをそれに応じて調整できるように、今後30日間のユーザー生涯価値(LTV)について予測する必要があります。顧客データセットはBigQueryにあり、AutoMLテーブルを使用してトレーニングするための表形式のデータを準備しています。このデータには、複数の列にまたがる時報があります。

AutoMLがデータに最適なモデルに適合することをどのように確認する必要がありますか？

- A. 手動変換を実行せずにトレーニング用のデータを送信し、[時間]列として適切な列を指定します。AutoMLが提供された時間信号に基づいてデータを分割できるようにし、検証およびテストセット用に最新のデータを予約します。
- B. 手動変換を実行せずにトレーニング用のデータを送信するAutoMLが適切な変換を処理できるようにするトレーニング、検証、およびテストセット間で分割される自動データを選択する
- C. 時間信号を含むすべての列を手動で配列に結合しますAutoMLがこの配列を適切に解釈できるようにしますトレーニング、検証、およびテストセット間で分割される自動データを選択します
- D. 手動変換を実行せずにトレーニング用にデータを送信する時間信号のある列を使用してデータを手動で分割する検証セットのデータがトレーニングセットのデータから30日後のものであり、データがテストセットでは、検証セットの30日後からです

**Answer: (**[解答を表示する](#)**)**

#### 最新問題: 5

機械学習チームは、Amazon SageMakerを使用して、調査データセットを使用してApacheMXNet手書き数字分類モデルをトレーニングします。チームは、モデルが過剰適合しているときに通知を受け取りたいと考えています。

監査人は、Amazon SageMakerログアクティビティレポートを表示して、不正なAPI呼び出しがないことを確認したいと考えています。

最小限のコードと最小限の手順で要件に対応するには、機械学習チームは何をすべきですか？

- A. AWSLambda関数を実装してAmazonSageMakerAPI呼び出しをAmazonS3に記録します。カスタムメトリクスをAmazonCloudWatchにプッシュするコードを追加します。CloudWatchでAmazonSNSを使用してアラームを作成し、モデルが過剰適合したときに通知を受信します。
- B. AWS CloudTrailを使用して、AmazonSageMakerAPI呼び出しをAmazonS3に記録します。カスタムメトリクスをAmazonCloudWatchにプッシュするコードを追加します。CloudWatchでAmazonSNSを使用してアラームを作成し、モデルが過剰適合したときに通知を受信します。
- C. AWSLambda関数を実装してAmazonSageMakerAPI呼び出しをAWSCloudTrailに記録します。カスタムメトリクスをAmazonCloudWatchにプッシュするコードを追加します。CloudWatchでAmazonSNSを使用してアラームを作成し、モデルが過剰適合したときに通知を受信します。

D. AWS CloudTrailを使用して、AmazonS3へのAmazonSageMakerAPI呼び出しをログに記録します。モデルが過剰適合したときに通知を受信するようにAmazonSNSを設定します

**Answer: C** ([メッセージを残す](#))

#### 最新問題: 6

チームは、画像に運転免許証、パスポート、またはクレジットカードが含まれているかどうかを予測するモデルを構築する必要があります。データエンジニアリングチームはすでにパイプラインを構築し、運転免許証付きの10,000枚の画像、パスポート付きの1,000枚の画像、クレジットカード付きの1,000枚の画像で構成されるデータセットを生成しました。ここで、次のラベルマップを使用してモデルをトレーニングする必要があります ['driverslicense', 'passport', 'credit\_card']。どの損失関数を使用する必要がありますか？

- A. カテゴリヒンジ
- B. バイナリクロスエントロピー
- C. カテゴリクロスエントロピー
- D. スパースカテゴリクロスエントロピー

**Answer: D** ([メッセージを残す](#))

`sesparse_categorical_crossentropy`。上記の3クラス分類問題の例 [1]、[2]、[3]

#### 最新問題: 7

大企業は、さまざまな運用メトリックから収集されたデータを使用してレポートとダッシュボードを生成するBIアプリケーションを開発しました。同社は、経営幹部が自然言語を使用してレポートからデータを取得できるように、高度なエクスペリエンスを提供したいと考えています。同社は、幹部が書面と口頭のインターフェースを使用して質問できるようにしたいと考えています。この会話型インターフェースを構築するために使用できるサービスの組み合わせはどれですか？ (3つ選択してください。)

- A. アマゾン理解
- B. Amazon Connect
- C. Amazon Lex
- D. Amazon Polly
- E. Alexa for Business
- F. Amazonトランスクリプト

**Answer: A,B,F** ([メッセージを残す](#))

#### 最新問題: 8

あなたはあなたの会社のeコマースウェブサイトで買い物客のためのMLレコメンデーションモデルを設計しています。推奨事項AIを使用して、システムを構築、テスト、および展開します。ベストプラクティスに従いながら、収益を増やすための推奨事項をどのように作成する必要がありますか？

- A. 「よく一緒に購入する」推奨タイプを使用して、注文ごとにショッピングカートのサイズを大きくします。

- B. ユーザーイベントをインポートしてから製品カタログをインポートし、最高品質のイベントストリームがあることを確認します
- C. あなたが好きかもしれない他の製品」推奨タイプを使用してクリック率を上げます
- D. 商品データの収集と記録には時間がかかるため、商品カタログのプレースホルダー値を使用して、モデルの実行可能性をテストします。

**Answer: B** ([メッセージを残す](#))

**最新問題: 9**

あなたは最近、新しいプロジェクトを間もなくリリースする機械学習チームに参加しました。プロジェクトのリーダーとして、MLコンポーネントの生産準備が整っているかどうかを判断するように求められます。チームはすでに機能とデータ、モデル開発、およびインフラストラクチャをテストしています。チームに推奨する追加の準備チェックはどれですか？

- A. モデルのパフォーマンスが監視されていることを確認します
- B. すべてのハイパーパラメータが調整されていることを確認します
- C. トレーニングが再現可能であることを確認します
- D. 機能の期待がスキーマに取り込まれていることを確認します

**Answer: (**[解答を表示する](#)**)**

**最新問題: 10**

ニューラルネットワークのバッチトレーニング中に、損失に振動があることに気づきます。モデルが確実に収束するように、モデルをどのように調整する必要がありますか？

- A. 学習率ハイパーパラメータを増やします
- B. トレーニングバッチのサイズを小さくします
- C. 学習率ハイパーパラメータを減らします
- D. トレーニングバッチのサイズを増やします

**Answer: (**[解答を表示する](#)**)**

**最新問題: 11**

あなたのチームは、何百万もの顧客が使用するグローバル銀行向けのアプリケーションを構築しています。3日後のcustomers1アカウントの残高を予測する予測モデルを作成しました。チームは、アカウントの残高が25ドルを下回る可能性がある場合にユーザーに通知する新機能で、結果を使用します。どのように予測を提供する必要がありますか？

- A. 1.Firebaseで通知システムを構築する  
2.各ユーザーをFirebaseCloudMessagingサーバーのユーザーIDに登録します。これにより、すべてのアカウント残高予測の平均が25ドルのしきい値を下回ったときに通知が送信されます。
- B. 1.Firebaseで通知システムを構築する  
2.各ユーザーをFirebaseCloudMessagingサーバーのユーザーIDに登録します。これにより、モデルがユーザーのアカウント残高が25ドルのしきい値を下回ると予測したときに通知が送信されます。
- C. 1.ユーザーごとにPub/Subトピックを作成します

2.ユーザーのアカウント残高が25ドルのしきい値を下回るとモデルが予測したときに通知を送信する、クラウド関数をデプロイします。

D. 1.ユーザーごとにPub/Subトピックを作成します

2.ユーザーのアカウント残高が25ドルのしきい値を下回るとモデルが予測したときに通知を送信する、AppEngine標準環境にアプリケーションをデプロイします

**Answer:** [\(解答を表示する\)](#)

#### 最新問題: 12

トレーニング時間を短縮して、セマンティック画像セグメンテーションの深層学習モデルをトレーニングしています。ディープラーニングVMイメージを使用しているときに、次のエラーが発生します。リソース'projects / deeplearning-platform / zones / europe-west4-c / AcceleratorTypes/nvidia-tesla-k80'が見つかりませんでした。あなたは何をするべきか？

A. 選択したリージョンにプリアンプティブGPUクォータがあることを確認します。

B. 選択したリージョンにGPUクォータがあることを確認します。

C. 選択したGPUにワークロードに十分なGPUメモリがあることを確認します。

D. 必要なGPUが選択したリージョンで使用可能であることを確認します。

**Answer:** B ([メッセージを残す](#))

#### 最新問題: 13

最近、組織のフレームワークに固有の重要な依存関係を使用するカスタムニューラルネットワークを設計および構築しました。GoogleCloudのマネージドトレーニングサービスを使用してモデルをトレーニングする必要があります。ただし、MLフレームワークと関連する依存関係は、AIPlatformTrainingではサポートされていません。また、モデルとデータの両方が大きすぎて、単一のマシンのメモリに収まりません。選択したMLフレームワークは、スケジューラー、ワーカー、およびサーバーの分散構造を使用します。あなたは何をするべきか？

A. AIPlatformTrainingで分散トレーニングジョブを実行するためのカスタムコンテナを構築します

B. AIPlatformTrainingでサポートされている依存関係を持つMLフレームワークにコードを再構成します

C. AIPlatformTrainingで利用可能な組み込みモデルを使用する

D. AIPlatformTrainingでジョブを実行するためのカスタムコンテナを構築する

**Answer:** A ([メッセージを残す](#))

#### 最新問題: 14

このグラフは、ニューラルネットワークのエポックに対するトレーニングと検証の損失を示しています。

トレーニング中のネットワークは次のとおりです。

\* 2つの密な層、1つの出力ニューロン

\*各層に100個のニューロン

\*100エポック

\*重みのランダムな初期化



検証セットの精度の観点からモデルのパフォーマンスを向上させるために使用できる手法はどれですか？

- A. 適切なシードを使用した重みのランダムな初期化
- B. エポック数を増やす
- C. 100個のニューロンで別のレイヤーを追加する
- D. 早期打ち切り

Answer: ([解答を表示する](#))

最新問題: 15

次のジョブ送信スクリプトを使用してテキストを要約するために、AIPlatformでLSTMベースのモデルをトレーニングしています。

```
gcloud ai-platform jobs submit training $JOB_NAME \
  --package-path $TRAINER_PACKAGE_PATH \
  --module-name $MAIN_TRAINER_MODULE \
  --job-dir $JOB_DIR \
  --region $REGION \
  --scale-tier basic \
  -- \
  --epochs 20 \
  --batch_size=32 \
  --learning_rate=0.001 \
```

モデルの精度を大幅に損なうことなく、トレーニング時間を最小限に抑える必要があります。あなたは何をすべきか？

- A. 「scale-tier」パラメータを変更します
- B. バッチサイズのパラメータを変更します
- C. 「学習率」パラメータを変更します
- D. 「エポック」パラメータを変更します

**Answer: B** ([メッセージを残す](#))

#### 最新問題: 16

ある会社がAmazon SageMaker環境をセットアップしています。企業のデータセキュリティポリシーでは、インターネットを介した通信は許可されていません。

Amazon SageMakerノートブックインスタンスへの直接インターネットアクセスを有効にせずに、Amazon SageMakerサービスを有効にするにはどうすればよいですか？

- A. 企業VPC内にNATゲートウェイを作成します。
- B. Amazon SageMakerトラフィックをオンプレミスネットワーク経由でルーティングします。
- C. 企業VPC内にAmazon SageMaker VPCインターフェースエンドポイントを作成します。
- D. Amazon SageMakerをホストしているAmazon VPCとのVPCピアリングを作成します。

**Answer: (**[解答を表示する](#)**)**

説明/参照 : <https://docs.aws.amazon.com/sagemaker/latest/dg/sagemaker-dg.pdf> 46)

有効な **Professional-Machine-Learning-Engineer** 問題集は GoShiken.com が提供された合格しやすい Professional-Machine-Learning-Engineer 試験問題集！ GoShiken.com が最新の **Professional-Machine-Learning-Engineer** 試験問題集を提供しています。GoShiken.com Professional-Machine-Learning-Engineer 試験問題は最新で、解答が正確でございます。最新の GoShiken.com Professional-Machine-Learning-Engineer 問題集をゲットする人はこちら:

**最新問題: 17**

計算コストの高い前処理操作を必要とするデータセットでモデルをトレーニングしました。予測時に同じ前処理を実行する必要があります。高スループットのオンライン予測のために、モデルをAIPlatformにデプロイしました。どのアーキテクチャを使用する必要がありますか？

**A.** 着信予測リクエストデータをCloudSpannerにストリーミングします

\*前処理ロジックを抽象化するビューを作成します。

\*新しいレコードについて毎秒ビューをクエリします

\*変換されたデータを使用してAIPlatformに予測リクエストを送信します

\*予測をアウトバウンドPub/Subキューに書き込みます。

**B.** 前処理されたデータでトレーニングしたモデルの精度を検証します

\*生データを使用し、リアルタイムで利用できる新しいモデルを作成します

\*オンライン予測のために新しいモデルをAIPlatformにデプロイします

**C.** 着信予測リクエストをPub/Subトピックに送信します

\*メッセージがPub/Subトピックに公開されたときにトリガーされるクラウド関数を設定します。

\*クラウド関数に前処理ロジックを実装する

\*変換されたデータを使用してAIPlatformに予測リクエストを送信します

\*予測をアウトバウンドPub/Subキューに書き込みます

**D.** 着信予測リクエストをPub/Subトピックに送信します

\*データフロージョブを使用して受信データを変換します

\*変換されたデータを使用してAIPlatformに予測リクエストを送信します

\*予測をアウトバウンドPub/Subキューに書き込みます

**Answer:** ([解答を表示する](#))

**最新問題: 18**

BigQuery MLで線形回帰モデルを構築して、顧客が会社の製品を購入する可能性を予測します。モデルは、主要な予測コンポーネントとして都市名変数を使用します。モデルをトレーニングして提供するには、データを列に整理する必要があります。予測可能な変数を維持しながら、最小限のコーディングを使用してデータを準備する必要があります。あなたは何をするべきか？

**A.** Dataprepを使用して、ワンホットエンコーディング方式を使用して状態列を変換し、各都市をバイナリ値の列にします。

**B.** Cloud Data Fusionを使用して各都市を1、2、3、4、または5rのラベルが付いた地域に割り当て、その番号を使用してモデル内の都市を表します。

**C.** TensorFlowを使用して、語彙リストを含むカテゴリ変数を作成します。語彙ファイルを作成し、モデルの一部としてBigQueryMLにアップロードします。

**D.** 都市情報を含む列を含まないBigQueryを使用して新しいビューを作成します

**Answer:** B ([メッセージを残す](#))

#### 最新問題: 19

カスタマーサポートの電子メールを分類するためのモデルを開発しています。オンプレミスシステムで小さなデータセットを使用してTensorFlowEstimatorsでモデルを作成しましたが、高いパフォーマンスを確保するには、大きなデータセットを使用してモデルをトレーニングする必要があります。モデルをGoogleCloudに移植し、コードのリファクタリングとインフラストラクチャのオーバーヘッドを最小限に抑えて、オンプレミスからクラウドへの移行を容易にします。あなたは何をするべきか？

- A. トレーニング用にDataprocでクラスターを作成します
- B. Kubeflowパイプラインを使用してGoogleKubernetesEngineクラスターでトレーニングします。
- C. 分散トレーニングにAIPlatformを使用する
- D. 自動スケーリングを使用してマネージドインスタンスグループを作成します

**Answer:** ([解答を表示する](#))

#### 最新問題: 20

企業の機械学習スペシャリストは、TensorFlowを使用して時系列予測モデルのトレーニング速度を向上させる必要があります。トレーニングは現在、単一のGPUマシンで実装されており、完了するまでに約23時間かかります。トレーニングは毎日実行する必要があります。

モデルの精度は許容範囲内ですが、同社はトレーニングデータのサイズが継続的に増加し、モデルを毎日ではなく1時間ごとに更新する必要があると予想しています。同社はまた、コーディングの労力とインフラストラクチャの変更を最小限に抑えたいと考えています。

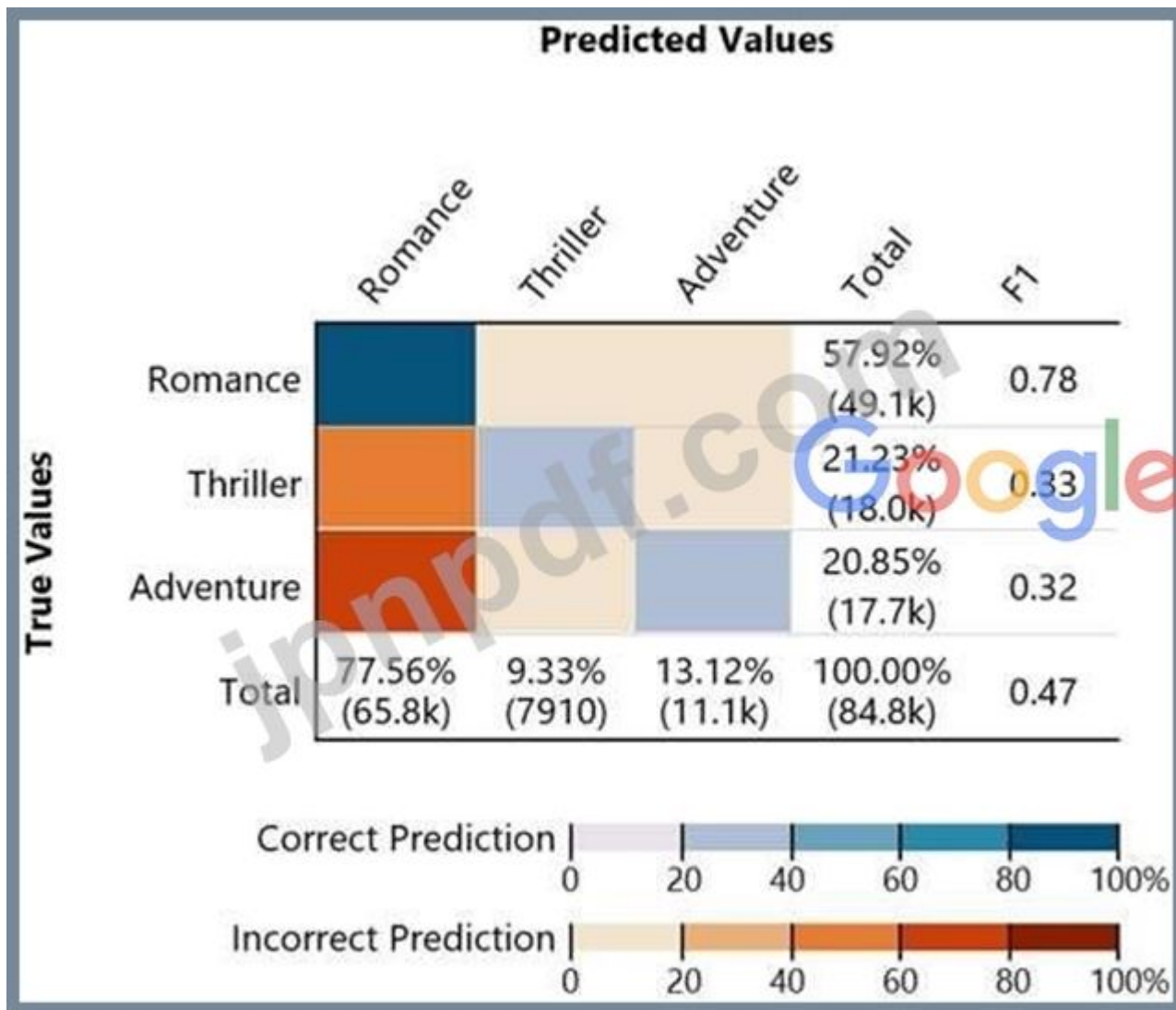
機械学習スペシャリストは、将来の需要に合わせて拡張できるように、トレーニングソリューションに対して何をすべきですか？

- A. トレーニングをAmazon EMRに移動し、ビジネス目標を達成するために必要な数のマシンにワークロードを分散します。
- B. TensorFlowコードを変更しないでください。トレーニングをスピードアップするために、マシンをより強力なGPUを備えたマシンに変更します。
- C. TensorFlowコードを変更して、AmazonSageMakerでサポートされるHorovod分散フレームワークを実装します。ビジネス目標を達成するために必要な数のマシンにトレーニングを並列化します。
- D. 組み込みのAWSSageMakerDeepARモデルの使用に切り替えます。ビジネス目標を達成するために必要な数のマシンにトレーニングを並列化します。

**Answer:** C ([メッセージを残す](#))

#### 最新問題: 21

映画分類モデルの次の混同行列を考えると、ロマンスの真のクラス頻度とアドベンチャーの予測クラス頻度はどれくらいですか？



- A. ロマンズの真のクラス頻度は77.56%で、アドベンチャーの予測クラス頻度は20.85%
- B. ロマンズの真のクラス頻度は57.92%で、アドベンチャーの予測クラス頻度は13.12%
- C. ロマンズの真のクラス頻度は77.56% \* 0.78であり、アドベンチャーの予測クラス頻度は20.85% \* 0.32
- D. ロマンズの真のクラス頻度は0.78で、アドベンチャーの予測クラス頻度は 0.47-0.32)

**Answer:** ([解答を表示する](#))

#### 最新問題: 22

AI Platformを使用してMLモデルを開発し、それを本番環境に移行したいと考えています。1秒あたり数千のクエリを処理し、遅延の問題が発生しています。着信リクエストは、Google Kubernetes Engine (GKE)で実行されている複数のKubeflowCPU専用ポッドに分散するロードバランサーによって処理されます。目標は、基盤となるインフラストラクチャを変更せずに、サービスのレイテンシを改善することです。あなたは何をすべきか？

- A. max\_enqueued\_batchesTensorFlowサービングパラメーターを大幅に増やします
- B. TensorFlow Servingのtensorflow-model-server-universalバージョンに切り替えます

C. max\_batch\_sizeTensorFlowサービングパラメーターを大幅に増やします

D. CPU固有の最適化をサポートするためにソースを使用してTensorFlowサービングを再コンパイルするサービングノードに適切なベースライン最小CPUプラットフォームを選択するようにGKEに指示します

**Answer: D** ([メッセージを残す](#))

#### 最新問題: 23

あなたは公共交通機関で働いており、複数の輸送ルートの遅延時間を推定するためのモデルを構築する必要があります。予測は、アプリ内のユーザーにリアルタイムで直接提供されます。季節や人口の増加がデータの関連性に影響を与えるため、毎月モデルを再トレーニングします。Googleが推奨するベストプラクティスに従う必要があります。予測モデルのエンドツーエンドアーキテクチャをどのように構成する必要がありますか？

A. トレーニングからモデルのデプロイまでのマルチステップワークフローをスケジュールするようにKubeflowパイプラインを構成します。

B. Cloud Composerを使用して、トレーニングからモデルのデプロイまでのワークフローを実行するDataflowジョブをプログラムでスケジュールします

C. CloudSchedulerによってトリガーされるAIプラットフォームでトレーニングとデプロイのジョブを起動するCloudFunctionsスクリプトを記述します

D. BigQuery MLでトレーニングおよびデプロイされたモデルを使用し、BigQueryのスケジュールされたクエリ機能を使用して再トレーニングをトリガーします

**Answer: (**[解答を表示する](#)**)**

#### 最新問題: 24

機械学習スペシャリストは、AmazonSageMakerにカスタムアルゴリズムを導入したいと考えています。スペシャリストは、AmazonSageMakerでサポートされているDockerコンテナにアルゴリズムを実装します。

Amazon SageMakerがトレーニングを正しく起動できるように、スペシャリストはDockerコンテナをどのようにパッケージ化する必要がありますか？

A. トレーニングプログラムをディレクトリ/opt/ml/trainにコピーします

B. コンテナ内のbash\_profileファイルを変更し、bashコマンドを追加してトレーニングプログラムを開始します

C. DockerfileでCMD configを使用して、トレーニングプログラムをイメージのCMDとして追加します

D. トレーニングプログラムをENTRYPOINTという名前のトレインとして構成します

**Answer: C** ([メッセージを残す](#))

#### 最新問題: 25

あなたは広告会社で働いていて、あなたの会社の最新の広告キャンペーンの効果を理解したいと思っています。500MBのキャンペーンデータをBigQueryにストリーミングしました。テーブルを

クエリしてから、AIPlatformノートブックのpandasデータフレームを使用してそのクエリの結果を操作します。あなたは何をすべきか？

- A. テーブルをBigQueryからローカルCSVファイルとしてダウンロードし、AIPlatformノートブックインスタンスにアップロードします。パンダを使用します。read\_csvを使用して、ファイルをpandasデータフレームとして取り込みます
- B. AI Platform NotebooksのBigQueryセルマジックを使用してデータをクエリし、結果をパンダデータフレームとして取り込みます
- C. AI Platformノートブックのbashセルから、bqextractコマンドを使用してテーブルをCSVファイルとしてCloudStorageにエクスポートし、gsutilcpを使用してデータをノートブックにコピーします。pandasを使用します。read\_csvを使用して、ファイルをpandasデータフレームとして取り込みます
- D. テーブルをCSVファイルとしてBigQueryからGoogleドライブにエクスポートし、GoogleドライブAPIを使用してファイルをノートブックインスタンスに取り込みます

**Answer: A** ([メッセージを残す](#))

#### 最新問題: 26

あなたは、AIプラットフォームを使用する50人以上のデータサイエンティストからなる成長中のチームで働いています。クリーンでスケーラブルな方法でジョブ、モデル、およびバージョンを整理するための戦略を設計しています。どの戦略を選ぶべきですか？

- A. ラベルを使用して、リソースを説明的なカテゴリに整理します。作成された各リソースにラベルを適用して、ユーザーがリソースを表示または監視するときにラベルで結果をフィルタリングできるようにします。
- B. AIプラットフォームノートブックに制限付きのIAM権限を設定して、特定のインスタンスにアクセスできるのは1人のユーザーまたはグループのみになるようにします。
- C. 各データサイエンティストの作業を異なるプロジェクトに分けて、各データサイエンティストによって作成されたジョブ、モデル、およびバージョンにそのユーザーだけがアクセスできるようにします。
- D. AIプラットフォームのリソース使用状況に関する情報を取得するために適切にフィルタリングされるクラウドログログ用のBigQueryシンクを設定します。BigQueryで、ユーザーを使用しているリソースにマッピングするSQLビューを作成します

**Answer: B** ([メッセージを残す](#))

#### 最新問題: 27

あなたは最近、数千のデータセットを持つ企業規模の企業に加わりました。BigQueryの各テーブルには正確な説明があることを知っており、AIプラットフォームで構築しているモデルに使用する適切なBigQueryテーブルを検索しています。必要なデータをどのように見つける必要がありますか？

- A. データカタログを使用して、テーブルの説明でキーワードを使用してBigQueryデータセットを検索します。

- B. テーブルの説明をテーブルIDにマッピングするルックアップテーブルをBigQueryで維持します。ルックアップテーブルをクエリして、必要なデータの正しいテーブルIDを見つけます。
- C. BigQueryでクエリを実行し、BigQueryにネイティブなINFORMATION\_SCHEMAメタデータテーブルを使用して、プロジェクト内の既存のすべてのテーブル名を取得します。結果を使用して、必要なテーブルを見つけます。
- D. AIプラットフォーム上の各モデルとバージョンのリソースに、トレーニングに使用されたBigQueryテーブルの名前をタグ付けします。

**Answer: D** ([メッセージを残す](#))

#### 最新問題: 28

機械学習スペシャリストは、ストリーミングデータを取り込んで、調査と分析のためにApacheParquetファイルに保存できる必要があります。

次のサービスのうち、このデータを正しい形式で取り込み、保存するのはどれですか？

- A. AWS DMS
- B. AmazonKinesisデータストリーム
- C. Amazon Kinesis Data Firehose
- D. Amazon Kinesis Data Analytics

**Answer:** ([解答を表示する](#))

#### 最新問題: 29

TPUを使用してAIPlatformでResnetモデルをトレーニングし、自動車エンジンの欠陥のタイプを視覚的に分類しています。Cloud TPUプロファイラープラグインを使用してトレーニングプロファイルをキャプチャし、それが高度に入力にバインドされていることを確認します。ボトルネックを減らし、モデルのトレーニングプロセスをスピードアップしたいと考えています。tf.dataデータセットにどのような変更を加える必要がありますか？

2つの答えを選択してください

- A. 変換のバッチサイズ引数を減らします
- B. 繰り返しパラメータの値を減らします
- C. シャッフルオプションのバッファサイズを増やします。
- D. データの読み取りにインターリーブオプションを使用します
- E. プリフェッチオプションをトレーニングバッチサイズと同じに設定します

**Answer: D,E** ([メッセージを残す](#))

#### 最新問題: 30

機械学習スペシャリストは、AmazonSageMakerを使用して時系列予測を実行するモデルを構築しています。スペシャリストはモデルのトレーニングを終了し、エンドポイントで負荷テストを実行して、モデルバリエーションの自動スケーリングを構成できるようにすることを計画しています。

スペシャリストが負荷テスト中に遅延、メモリ使用率、およびCPU使用率を確認できるようにするアプローチはどれですか。

- A. AmazonAthenaとAmazonQuickSightを利用してAmazonS3に書き込まれたSageMakerログを確認し、作成中のログを視覚化します。
- B. Amazon CloudWatchダッシュボードを生成して、Amazon SageMakerによって出力されるレイテンシー、メモリ使用率、およびCPU使用率のメトリックスの単一のビューを作成します。
- C. カスタムAmazon CloudWatchログを作成し、Amazon ESとKibanaを利用して、AmazonSageMakerによって生成されたログデータをクエリして視覚化します。
- D. AmazonSageMakerによって生成されたAmazonCloudWatchログをAmazonESに送信し、Kibanaを使用してログデータをクエリおよび視覚化します。

**Answer: B** ([メッセージを残す](#))

説明/リファレンス <https://docs.aws.amazon.com/sagemaker/latest/dg/monitoring-cloudwatch.html>

#### 最新問題: 31

機械学習スペシャリストは、AmazonAthenaを使用してAmazonS3上のデータセットをクエリするプロセスを構築する必要があります。データセットには、プレーンテキストのCSVファイルとして保存された800,000を超えるレコードが含まれています。各レコードには200列で、サイズは約1.5MBです。ほとんどのクエリは5~10列のみに及びます。機械学習スペシャリストは、クエリの実行時間を最小限に抑えるためにデータセットをどのように変換する必要がありますか？

- A. レコードをApacheParquet形式に変換します。
- B. レコードをJSON形式に変換します。
- C. レコードをGZIPCSV形式に変換します。
- D. レコードをXML形式に変換します。

**Answer: (**[解答を表示する](#)**)**

圧縮を使用すると、Amazon Athenaによってスキャンされるデータの量が減り、S3バケットストレージも減ります。これは、AWSの請求書にとってWin-Winです。サポートされている形式：GZIP、LZO、SNAPPY (Parquet)、およびZLIB。

参照 <https://www.cloudforecast.io/blog/using-parquet-on-athena-to-save-money-on-aws/>

有効な **Professional-Machine-Learning-Engineer** 問題集は GoShiken.com が提供された合格しやすい Professional-Machine-Learning-Engineer 試験問題集！ GoShiken.com が最新の **Professional-Machine-Learning-Engineer** 試験問題集を提供しています。GoShiken.com Professional-Machine-Learning-Engineer 試験問題は最新で、解答が正確でございます。最新の GoShiken.com Professional-Machine-Learning-Engineer 問題集をゲットする人はこちら：  
<https://www.goshiken.com/Google/Professional-Machine-Learning-Engineer-mondaishu.html>  
(**30030%OFF**問題集溶と正解付きで 30%w 特別割引コード: **Freepdfdumps**)

#### 最新問題: 32

GoogleCloudでディープニューラルネットワークモデルをトレーニングしました。モデルのトレーニングデータの損失は少ないですが、検証データのパフォーマンスが低下しています。モデルが過剰適合に対して弾力性があるようにする必要があります。モデルを再トレーニングするときに、どの戦略を使用する必要がありますか？

- A. AI Platformでハイパーパラメータ調整ジョブを実行して学習率を最適化し、ニューロンの数を2倍に増やします。
- B. AI Platformでハイパーパラメータ調整ジョブを実行して、L2正則化およびドロップアウトパラメータを最適化します
- C. ドロップアウトパラメータ0.2を適用し、学習率を10分の1に減らします。
- D. 0.4のL2正則化パラメータを適用し、学習率を10分の1に減らします。

**Answer: C** ([メッセージを残す](#))

#### 最新問題: 33

トラック会社は、世界中のトラック群からライブ画像データを収集しています。データは急速に増加しており、毎日約100GBの新しいデータが生成されます。同社は、特定のIAMユーザーのみがデータにアクセスできるようにしながら、機械学習のユースケースを調査したいと考えています。

どのストレージオプションが最も処理の柔軟性を提供し、IAMによるアクセス制御を可能にしますか？

- A. Amazon DynamoDBなどのデータベースを使用してイメージを保存し、目的のIAMユーザーのみにアクセスを制限するようにIAMポリシーを設定します。
- B. Amazon S3でバックアップされたデータレイクを使用して生の画像を保存し、バケットポリシーを使用して権限を設定します。
- C. Hadoop分散ファイルシステム (HDFS) を使用してAmazon EMRをセットアップし、ファイルを保存し、IAMポリシーを使用してEMRインスタンスへのアクセスを制限します。
- D. IAMポリシーを使用してAmazon EFSを設定し、IAMユーザーが所有するAmazon EC2インスタンスでデータを利用できるようにします。

**Answer: C** ([メッセージを残す](#))

説明

#### 最新問題: 34

データサイエンティストは、複数のクラスを持つデータセットで多層知覚 (MLP) をトレーニングしています。対象のターゲットクラスは、データセット内の他のクラスと比較して一意ですが、許容可能なリコールメトリックを達成していません。データサイエンティストは、MLPの非表示レイヤーの数とサイズを変更しようとしたますが、結果は大幅に改善されていません。リコールを改善するためのソリューションは、可能な限り迅速に実装する必要があります。

これらの要件を満たすためにどの手法を使用する必要がありますか？

- A. MLPの代わりにXGBoostモデルをトレーニングする
- B. MLPの代わりに異常検出モデルをトレーニングする

- C. Amazon Mechanical Turkを使用してより多くのデータを収集し、再トレーニングします
- D. MLPの損失関数にクラスの重みを追加してから、再トレーニングします

**Answer: A** ([メッセージを残す](#))

#### 最新問題: 35

データエンジニアは、顧客のクレジットカード情報を含むデータセットを使用してモデルを構築する必要があります。データエンジニアは、データが暗号化されたままであり、クレジットカード情報が安全であることをどのように確認できますか？

- A. カスタム暗号化アルゴリズムを使用してデータを暗号化し、VPCのAmazonSageMakerインスタンスにデータを保存します。SageMaker DeepARアルゴリズムを使用して、クレジットカード番号をランダム化します。
- B. IAMポリシーを使用して、AmazonS3バケットとAmazonKinesisのデータを暗号化し、クレジットカード番号を自動的に破棄して偽のクレジットカード番号を挿入します。
- C. VPCのSageMakerインスタンスにデータがコピーされたら、AmazonSageMaker起動設定を使用してデータを暗号化します。SageMaker主成分分析 (PCA) アルゴリズムを使用して、クレジットカード番号の長さを減らします。
- D. AWSKMSを使用してAmazonS3とAmazonSageMakerのデータを暗号化し、AWSGlueを使用して顧客データからクレジットカード番号を編集します。

**Answer: C** ([メッセージを残す](#))

説明/参照 <https://docs.aws.amazon.com/sagemaker/latest/dg/pca.html>

#### 最新問題: 36

あなたは規制対象の保険会社のMLエンジニアです。潜在的な顧客からの保険申請を受け入れるか拒否する保険承認モデルを開発するように求められます。モデルを構築する前に、どのような要素を考慮する必要がありますか？

- A. 差分プライバシー連合学習と説明性
- B. トレーサビリティ、再現性、説明性
- C. 連合学習、再現性、説明性
- D. 編集、再現性、説明性

**Answer: B** ([メッセージを残す](#))

#### 最新問題: 37

AIPlatformに複数のバージョンの画像分類モデルを展開しました。モデルバージョンのパフォーマンスを時間の経過とともに監視する必要があります。この比較をどのように実行する必要がありますか？

- A. What-Ifツールを使用して、各モデルの受信者動作特性 (ROC) 曲線を比較します。
- B. 連続評価機能を使用して、モデル全体の平均平均精度を比較します
- C. 検証データの各モデルの損失パフォーマンスを比較します
- D. 保持されたデータセットの各モデルの損失パフォーマンスを比較します。

**Answer:** ([解答を表示する](#))

**最新問題: 38**

あなたのチームは、DNN回帰モデルをトレーニングしてテストし、良好な結果を得ました。展開から6か月後、入力データの分布が変化したため、モデルのパフォーマンスが低下しました。本番環境での入力の違いにどのように対処する必要がありますか？

- A. モデルを再トレーニングし、ハイパーパラメーター調整サービスを使用してL2正則化パラメーターを選択します
- B. モデルで特徴選択を実行し、より少ない特徴でモデルを再トレーニングします
- C. モデルで特徴選択を実行し、より少ない特徴で月ごとにモデルを再トレーニングします
- D. スキューを監視するアラートを作成し、モデルを再トレーニングします。

**Answer: A** ([メッセージを残す](#))

**最新問題: 39**

クレジットカード会社は、新しいクレジットカード申請者がクレジットカードの支払いをデフォルトにするかどうかを予測するのに役立つクレジットスコアリングモデルを構築したいと考えています。同社は、何千もの生の属性を持つ多数のソースからデータを収集しています。分類モデルをトレーニングするための初期の実験では、多くの属性が高度に相関しており、多数の機能によってトレーニング速度が大幅に低下し、過剰適合の問題があることが明らかになりました。このプロジェクトのデータサイエンティストは、元のデータセットから多くの情報を失うことなく、モデルのトレーニング時間を短縮したいと考えています。

データサイエンティストが目的を達成するために使用する必要がある特徴工学手法はどれですか？

- A. すべての機能で自己相関を実行し、相関性の高い機能を削除します
- B. すべての数値を0から1の間に正規化します
- C. オートエンコーダーまたは主成分分析 (PCA) を使用して、元の機能を新しい機能に置き換えます
- D. k-meansを使用して生データをクラスター化し、各クラスターのサンプルデータを使用して新しいデータセットを構築します

**Answer: B** ([メッセージを残す](#))

**最新問題: 40**

大規模なモバイルネットワーク運営会社は、サービスの購読を解除する可能性のある顧客を予測するための機械学習モデルを構築しています。解約のコストはインセンティブのコストよりはるかに高いため、同社はこれらの顧客にインセンティブを提供する予定です。

モデルは、100人の顧客のテストデータセットで評価した後、次の混同行列を生成します。

|                  |                 |    |
|------------------|-----------------|----|
| n= 100           | PREDICTED CHURN |    |
|                  | Yes             | No |
| ACTUAL Churn Yes | 10              | 4  |
| Actual No        | 10              | 76 |

モデルの評価結果に基づいて、なぜこれが生産のための実行可能なモデルなのでしょうか？

- A. モデルの精度は86%で、モデルの精度よりも高くなっています。
- B. モデルの精度は86%であり、モデルの精度よりも低くなっています。
- C. モデルは86%正確であり、誤検知の結果として会社が負担するコストは誤検知よりも少なくなります。
- D. モデルは86%正確であり、誤検知の結果として会社が負担するコストは誤検知よりも少なくなります。

**Answer:** [\(解答を表示する\)](#)

#### 最新問題: 41

Compute EngineでGPUを搭載した仮想マシンを使用して、特定のイメージに存在する政府IDのタイプを予測するコンピュータービジョンモデルをトレーニングする必要があります。次のパラメータを使用します。

\*オプティマイザー :SGD

\*画像の形状=224x224

\*バッチサイズ=64

\*エポック=10

\*詳細=2

トレーニング中に、次のエラーが発生します。ResourceExhaustedError :テンソルの割り当て時にメモリ不足 (oom)。あなたは何をすべきか？

- A. オプティマイザを変更します
- B. バッチサイズを小さくします
- C. 画像の形状を縮小します
- D. 学習率を変更する

**Answer: A** [\(メッセージを残す\)](#)

#### 最新問題: 42

機械学習スペシャリストは、小さなデータサンプルを使用して企業の概念実証を完了しました。これで、スペシャリストはAmazonSageMakerを使用してAWSでエンドツーエンドのソリューションを実装する準備が整いました。過去のトレーニングデータはAmazonRDSに保存されます。スペシャリストは、そのデータを使用してモデルをトレーニングするためにどのアプローチを使用する必要がありますか？

- A. データをAmazon DynamoDBに移動し、ノートブック内でDynamoDBへの接続を設定して、データをプルします。

- B. ノートブック内のSQLデータベースへの直接接続を書き込み、データをプルインします
- C. AWS DMSを使用してデータをAmazon ElastiCacheに移動し、ノートブック内に接続を設定して、高速アクセスのためにデータをプルします。
- D. AWS データパイプラインを使用してMicrosoft SQL ServerからAmazon S3にデータをプッシュし、ノートブック内のS3の場所を提供します。

**Answer: D** ([メッセージを残す](#))

#### 最新問題: 43

機械学習スペシャリストは、さまざまな経済的要因に基づいて将来の雇用率を予測するモデルを構築しています。データを調べている間、スペシャリストは入力特徴の大きさが大きく異なることに気づきます。スペシャリストは、より大きな大きさの変数がモデルを支配することを望んでいません。

モデルトレーニング用のデータを準備するには、スペシャリストは何をすべきですか？

- A. 分位ビニングを適用してデータをカテゴリビンにグループ化し、大きさを分布に置き換えることでデータの関係を維持します。
- B. デカルト積変換を適用して、大きさに依存しないフィールドの新しい組み合わせを作成します。
- C. 正規化を適用して、各フィールドの平均が0になり、分散が1になるようにして、有意な大きさを削除します。
- D. 直交スパースバイグラム (OSB) 変換を適用して、固定サイズのスライディングウィンドウを適用し、同様の大きさの新しい特徴を生成します。

**Answer: C** ([メッセージを残す](#))

説明/参照 <https://docs.aws.amazon.com/machine-learning/latest/dg/data-transformations-reference.html>

#### 最新問題: 44

あなたのチームは、畳み込みニューラルネットワーク (CNN) ベースのアーキテクチャをゼロから構築しています。オンプレミスのCPUのみのインフラストラクチャで実行された予備実験は有望でしたが、収束が遅くなりました。市場投入までの時間を短縮するために、モデルトレーニングをスピードアップするように求められました。より強力なハードウェアを活用するために、Google Cloudで仮想マシン (VM) を試してみたいと考えています。コードには手動のデバイス配置が含まれておらず、Estimatorモデルレベルの抽象化にラップされていません。モデルをトレーニングする環境はどれですか？

- A. すべてのライブラリがプリインストールされたより強力なCPU `e2-highcpu-16` マシンを備えたディープラーニングVM。
- B. すべての依存関係が手動でインストールされたCompute Engineおよび8 GPU上のAVM。
- C. すべてのライブラリがプリインストールされた `n1-standard-2` マシンと1 GPUを備えたディープラーニングVM。
- D. Compute Engine上のAVMとすべての依存関係が手動でインストールされた1つのTPU。

**Answer:** ([解答を表示する](#))

**最新問題: 45**

あなたは大企業のコンタクトセンターのMLエンジニアです。録音された電話での会話から顧客の感情を予測する感情分析ツールを構築する必要があります。コンタクトセンターに電話をかけた顧客の性別、年齢、文化の違いがモデル開発パイプラインと結果のどの段階にも影響を与えないようにしながら、モデルを構築するための最良のアプローチを特定する必要があります。あなたは何をするべきか？

- A. 音声をテキストに変換し、単語に基づいてモデルを構築します
- B. 音声録音から直接感情を抽出します
- C. 音声をテキストに変換し、文章に基づいて感情を抽出します
- D. 音声をテキストに変換し、構文分析を使用して感情を抽出します

**Answer: C** ([メッセージを残す](#))

**最新問題: 46**

テクノロジーの新興企業は、複雑なディープニューラルネットワークとGPUコンピューティングを使用して、各顧客の習慣と相互作用に基づいて、既存の顧客に会社の製品を推奨しています。このソリューションは現在、ローカルで実行されている会社のGitリポジトリからプルされたTensorFlowモデルにデータをロードする前に、AmazonS3バケットから各データセットをプルします。このジョブは、進行状況を同じS3バケットに継続的に出力しながら、数時間実行されます。ジョブは、障害が発生した場合にいつでも一時停止、再開、および続行でき、中央キューから実行されます。

上級管理職は、ソリューションのリソース管理の複雑さと、プロセスを定期的に繰り返すことに伴うコストについて懸念しています。彼らは、ワークロードを自動化して、月曜日から金曜日の営業終了までに週に1回実行するように求めています。

最小のコストでソリューションを拡張するには、どのアーキテクチャを使用する必要がありますか？

- A. AWS Deep Learning Containersを使用してソリューションを実装し、スポットインスタンスで実行されているAWS Fargateを使用してワークロードを実行し、組み込みのタスクスケジューラを使用してタスクをスケジュールします
- B. 低コストのGPU互換のAmazon EC2インスタンスを使用してソリューションを実装し、AWS インスタンススケジューラを使用してタスクをスケジュールします
- C. AWS Deep Learning Containersを使用してソリューションを実装し、GPU互換のスポットインスタンスでAWSBatchを使用してコンテナをジョブとして実行します
- D. スポットインスタンスで実行されているAmazon ECSを使用してソリューションを実装し、ECSサービススケジューラを使用してタスクをスケジュールします

**Answer: A** ([メッセージを残す](#))

有効な **Professional-Machine-Learning-Engineer** 問題集は GoShiken.com が提供された合格しやすい Professional-Machine-Learning-Engineer 試験問題集！ GoShiken.com が最新の **Professional-Machine-Learning-Engineer** 試験問題集を提供しています。GoShiken.com Professional-Machine-Learning-Engineer 試験問題は最新で、解答が正確でございます。最新の GoShiken.com Professional-Machine-Learning-Engineer 問題集をゲットする人はこちら：  
<https://www.goshiken.com/Google/Professional-Machine-Learning-Engineer-mondaishu.html>  
(**30030%OFF**問題集溶と正解付きで **30%w**特別割引コード: **Freepdfdumps**)

**最新問題: 47**

現在BigQueryに保存されているいくつかの構造化データセットに対して分類ワークフローを構築する必要があります。分類を数回実行するため、コードを記述せずに次の手順を実行する必要があります。探索的データ分析、特徴選択、モデル構築、トレーニング、ハイパーパラメーターの調整と提供。あなたは何をするべきか？

- A. 分類タスクを実行するようにAutoMLテーブルを構成します
- B. BigQuery MLタスクを実行して、分類のロジスティック回帰を実行します
- C. AI Platformを使用して、ハイパーパラメータ調整用に構成された分類モデルジョブを実行します
- D. AI Platform Notebooksを使用して、パンダライブラリで分類モデルを実行します

**Answer: D** ([メッセージを残す](#))

**最新問題: 48**

個人情報 (PII) を含む可能性のあるファイルをGoogleCloudにストリーミングするリアルタイム予測エンジンを構築しています。Cloud Data Loss Prevention (DLP) APIを使用してファイルをスキャンするとします。許可されていない個人がPIIにアクセスできないようにするにはどうすればよいですか？

- A. すべてのファイルをGoogle CloudTにストリーミングしてから、データをBigQueryに書き込みます。DLPAPIを使用して、テーブルの一括スキャンを定期的に行います。
- B. DLP APIを使用してそのバケットの一括スキャンを定期的に行い、データを機密バケットまたは非機密バケットに移動します
- C. 機密データと非機密データの2つのバケットを作成するすべてのデータを非機密バケットに書き込むDLP APIを使用してそのバケットの一括スキャンを定期的に行い、機密データを機密バケットに移動します
- D. すべてのファイルをGoogle Cloudにストリーミングし、データのバッチをBigQueryに書き込みます。データがBigQueryに書き込まれている間に、DLPAPIを使用してデータの一括スキャンを実行します。
- E. データの3つのバケットを作成します：隔離 機密、および非機密すべてのデータを隔離バケットに書き込みます。

**Answer: (**[解答を表示する](#)**)**

#### 最新問題: 49

機械学習スペシャリストが適切なものを決定したい

エンドポイント自動スケーリング構成のSageMakerVariantInvocationsPerInstancesetting。

スペシャリストは、単一のインスタンスで負荷テストを実行し、サービスの低下がない場合の1秒あたりのピークリクエスト (RPS)は約20RPSであると判断しました。これは最初の展開であるため、スペシャリストは呼び出しの安全率を0.5に設定する予定です。

記載されているパラメーターに基づいて、インスタンスごとの呼び出し設定が1分ごとに測定される場合、スペシャリストはSageMakerVariantInvocationsPerInstance設定として何を設定する必要がありますか？

- A. 600
- B. 10
- C. 2,400
- D. 30

**Answer: A** ([メッセージを残す](#))

#### 最新問題: 50

ある会社は、Amazon SageMakerのデフォルトの組み込み画像分類アルゴリズムのトレーニング中に、低い精度を観察しています。データサイエンスチームは、ResNetアーキテクチャの代わりにInceptionニューラルネットワークアーキテクチャを使用したいと考えています。

次のうちどれがこれを達成しますか？ 2つ選択してください。)

- A. SageMakerチームでサポートケースを作成して、デフォルトの画像分類アルゴリズムをInceptionに変更します。
- B. 開始ネットワークコードをダウンロードしてapt-getでAmazon EC2インスタンスにインストールし、このインスタンスをAmazonSageMakerのJupyterノートブックとして使用します。
- C. 組み込みの画像分類アルゴリズムをカスタマイズしてInceptionを使用し、これをモデルトレーニングに使用します。
- D. AmazonSageMakerのカスタムコードとTensorFlowEstimatorを使用して、モデルにInceptionネットワークをロードし、これをモデルのトレーニングに使用します。
- E. InceptionネットワークがロードされたTensorFlow EstimatorにDockerコンテナをバンドルし、これをモデルトレーニングに使用します。

**Answer: (**[解答を表示する](#)**)**

#### 最新問題: 51

機械学習スペシャリストは、アプリケーション用のカスタムビデオレコメンデーションモデルを開発しています。このモデルのトレーニングに使用されるデータセットは非常に大きく、数百万のデータポイントがあり、AmazonS3バケットでホストされています。

スペシャリストは、移動に数時間かかり、ノートブックインスタンスに接続されている5GBのAmazonEBSボリュームを超えるため、このすべてのデータをAmazonSageMakerノートブックインスタンスにロードしないようにしたいと考えています。

スペシャリストがすべてのデータを使用してモデルをトレーニングできるようにするアプローチはどれですか？

**A.** データの小さなサブセットをSageMakerノートブックにロードし、ローカルでトレーニングします。トレーニングコードが実行されており、モデルパラメータが妥当であると思われることを確認します。AWS Deep LearningAMIを使用してAmazonEC2インスタンスを起動し、S3バケットをアタッチして完全なデータセットをトレーニングします。

**B.** AWS Deep LearningAMIを使用してAmazonEC2インスタンスを起動し、S3バケットをインスタンスにアタッチします。少量のデータをトレーニングして、トレーニングコードとハイパーパラメータを確認します。Amazon SageMakerに戻り、完全なデータセットを使用してトレーニングします

**C.** AWS Glueを使用して、データの小さなサブセットを使用してモデルをトレーニングし、データがAmazonSageMakerと互換性があることを確認します。パイプ入力モードを使用して、S3バケットからの完全なデータセットを使用してSageMakerトレーニングジョブを開始します。

**D.** データの小さなサブセットをSageMakerノートブックにロードし、ローカルでトレーニングします。トレーニングコードが実行されており、モデルパラメータが妥当であると思われることを確認します。パイプ入力モードを使用して、S3バケットからの完全なデータセットを使用してSageMakerトレーニングジョブを開始します。

**Answer: D** ([メッセージを残す](#))

#### 最新問題: 52

100を超える入力特徴を持ち、すべて-1から1までの値を持つ線形モデルを構築しています。多くの特徴が有益ではないと思われます。情報を提供する機能を元の形式のままにして、情報を提供しない機能をモデルから削除する必要があります。どのテクニックを使うべきですか？

**A.** 反復ドロップアウト手法を使用して、削除したときにモデルを劣化させない機能を特定します。

**B.** 主成分分析を使用して、最も情報量の少ない機能を排除します。

**C.** モデルを作成したら、シャープレイ値を使用して、どの機能が最も有益かを判断します。

**D.** L1正則化を使用して、情報量の少ない特徴の係数を0に減らします。

**Answer: (**[解答を表示する](#)**)**

#### 最新問題: 53

Compute EngineでGPUを搭載した仮想マシンを使用して、特定のイメージに存在する政府IDのタイプを予測するコンピュータービジョンモデルをトレーニングする必要があります。次のパラメータを使用します。

\*オプティマイザー :SGD

\*画像の形状= 224x224

\*バッチサイズ= 64

\*エポック= 10

\*詳細= 2

トレーニング中に、次のエラーが発生します。ResourceExhaustedError :テンソルの割り当て時にメモリ不足 (oom)。あなたは何をするべきか？

- A. 学習率を変更する
- B. 画像の形状を縮小します
- C. バッチサイズを小さくします
- D. オプティマイザを変更します

**Answer:** ([解答を表示する](#))

#### 最新問題: 54

あなたはモバイルアプリケーションを持っている銀行のMLエンジニアです。経営陣は、指紋に基づいて顧客のIDを検証するアプリのMLベースの生体認証を構築するように依頼しました。指紋は機密性の高い個人情報と見なされ、銀行のデータベースにダウンロードして保存することはできません。このMLモデルをトレーニングしてデプロイするために、どの学習戦略を推奨する必要がありますか？

- A. データ損失防止API
- B. 連合学習
- C. データを暗号化するMD5
- D. 差分プライバシー

**Answer:** ([解答を表示する](#))

#### 最新問題: 55

データサイエンティストは、企業のeコマースプラットフォームの不正なユーザーアカウントを特定する必要があります。同社は、新しく作成されたアカウントが以前に知られている不正なユーザーに関連付けられているかどうかを判断する機能を望んでいます。

データサイエンティストは、AWS Glueを使用して、取り込み中に会社のアプリケーションログをクレンジングしています。

データサイエンティストが不正なアカウントを特定できるようにする戦略はどれですか？

- A. 組み込みのFindDuplicatesAmazonAthenaクエリを実行します。
- B. AWSGlueでFindMatches機械学習トランスフォームを作成します。
- C. AWS Glueクローラーを作成して、ソースデータ内の重複するアカウントを推測します。
- D. AWSGlueデータカタログで重複するアカウントを検索します。

**Answer: B** ([メッセージを残す](#))

説明/参照 <https://docs.aws.amazon.com/glue/latest/dg/machine-learning.html>

#### 最新問題: 56

金融サービス会社は、AmazonS3上に堅牢なサーバーレスデータレイクを構築しています。データレイクは柔軟性があり、次の要件を満たす必要があります。

\*AmazonAthenaおよびAmazonRedshiftSpectrumを介したAmazonS3の新旧データのクエリをサポートします。

\*イベント駆動型ETLパイプラインをサポート

\*メタデータを理解するための迅速かつ簡単な方法を提供します

これらの要件を満たすアプローチはどれですか？

- A. AWS Glueクローラーを使用してS3データをクロールし、AWS Lambda関数を使用してAWS バッチジョブをトリガーし、外部のApacheHiveメタストアを使用してメタデータを検索および検出します。
- B. AWS Glueクローラーを使用してS3データをクロールし、AWS Lambda関数を使用して AWSGlue ETLジョブをトリガーし、AWSGlueデータカタログを使用してメタデータを検索および検出します。
- C. AWS Glueクローラーを使用してS3データをクロールし、AmazonCloudWatchアラームを使用してAWSGlue ETLジョブをトリガーし、外部のApacheHiveメタストアを使用してメタデータを検索および検出します。
- D. AWS Glueクローラーを使用してS3データをクロールし、Amazon CloudWatchアラームを使用してAWSバッチジョブをトリガーし、AWSGlueデータカタログを使用してメタデータを検索および検出します。

**Answer:** ([解答を表示する](#))

#### 最新問題: 57

ゲーム会社は、人々が無料でプレイを開始できるオンラインゲームを開始しましたが、特定の機能を使用することを選択した場合は料金を支払う必要があります。同社は、新規ユーザーが1年以内に有料ユーザーになるかどうかを予測する自動システムを構築する必要があります。同社は、100万人のユーザーからラベル付きのデータセットを収集しました。

トレーニングデータセットは、1,000のポジティブサンプル（1年以内に支払いを済ませたユーザーからのもの）と

999,000のネガティブサンプル（有料機能を使用しなかったユーザーから）。各データサンプルは、ユーザーの年齢、デバイス、場所、プレイパターンなどの200の機能で構成されています。

このデータセットをトレーニングに使用して、データサイエンスチームはランダムフォレストモデルをトレーニングしました。

トレーニングセットで99%の精度。ただし、テストデータセットの予測結果は満足のいくものではありませんでした。この問題を軽減するために、データサイエンスチームは次のどのアプローチを採用する必要がありますか？ 2つ選択してください。）

- A. 誤検知が誤検知よりもコスト値に大きな影響を与えるように、コスト関数を変更します。
- B. トレーニングデータセットのテストデータセットにサンプルのコピーを含めます。
- C. 誤検知が誤検知よりもコスト値に大きな影響を与えるように、コスト関数を変更します。
- D. モデルがより多くの機能を学習できるように、ランダムフォレストにさらに深い木を追加します。
- E. ポジティブサンプルを複製し、複製されたデータに少量のノイズを追加することにより、よりポジティブなサンプルを生成します。

**Answer:** ([解答を表示する](#))

#### 最新問題: 58

金融サービス会社は、AmazonS3上に堅牢なサーバーレスデータレイクを構築しています。データレイクは柔軟性があり、次の要件を満たす必要があります。

\* AmazonAthenaおよびAmazonRedshiftSpectrumを介したAmazonS3の新旧データのクエリをサポートします。

\* イベント駆動型ETLパイプラインをサポート

\* メタデータを理解するための迅速かつ簡単な方法を提供します

これらの要件を満たすアプローチはどれですか？

**A.** AWS Glueクローラーを使用してS3データをクロールし、AWS Lambda関数を使用してAWS バッチジョブをトリガーし、外部のApacheHiveメタストアを使用してメタデータを検索および検出します。

**B.** AWS Glueクローラーを使用してS3データをクロールし、AmazonCloudWatchアラームを使用してAWSGlue ETLジョブをトリガーし、外部のApacheHiveメタストアを使用してメタデータを検索および検出します。

**C.** AWS Glueクローラーを使用してS3データをクロールし、Amazon CloudWatchアラームを使用してAWSバッチジョブをトリガーし、AWSGlueデータカタログを使用してメタデータを検索および検出します。

**D.** AWS Glueクローラーを使用してS3データをクロールし、AWSLambda関数を使用してAWSGlue ETLジョブをトリガーし、AWSGlueデータカタログを使用してメタデータを検索および検出します。

**Answer: D** ([メッセージを残す](#))

#### 最新問題: 59

あなたはビジュアル検索エンジンを作成しているオンライン小売会社で働いています。画像に会社の商品が含まれているかどうかを分類するために、GoogleCloudにエンドツーエンドのMLパイプラインを設定しました。近い将来に新製品がリリースされることを期待して、パイプラインに再トレーニング機能を構成し、新しいデータをMLモデルにフィードできるようにしました。また、AI Platformの継続的な評価サービスを使用して、モデルがテストデータセットで高精度であることを確認する必要があります。あなたは何をするべきか？

**A.** 評価指標が事前に決定されたしきい値を下回ったときに、新しい製品の画像でテストデータセットを更新します。

**B.** 新しい製品が再トレーニングに組み込まれた場合でも、元のテストデータセットを変更しないでください

**C.** 再トレーニングが導入されたときに、テストデータセットを新しい製品の画像に置き換えます。

**D.** 再トレーニングが導入されたときに、新しい製品の画像を使用してテストデータセットを拡張します

**Answer: C** ([メッセージを残す](#))

#### 最新問題: 60

オンラインファッション会社で働く機械学習スペシャリストは、会社のAmazonS3ベースのデータレイク用のデータ取り込みソリューションを構築したいと考えています。

スペシャリストは、以下で構成される将来の機能を可能にする一連の取り込みメカニズムを作成したいと考えています。

- \*リアルタイム分析

- \*履歴データのインタラクティブな分析

- \*クリックストリーム分析

- \*製品の推奨事項

スペシャリストはどのサービスを使用する必要がありますか？

**A.** データカタログとしてのAWS Glue; リアルタイムのデータインサイトのための

AmazonKinesisDataStreamsとAmazonKinesisData Analytics; クリックストリーム分析のためにAmazonESに配信するAmazonKinesisData Firehose; パーソナライズされた製品の推奨事項を生成するAmazonEMR

**B.** データカタログとしてのAmazon Athena ;ほぼリアルタイムのデータインサイトのための

Amazon KinesisDataStreamsとAmazonKinesisData Analytics; クリックストリーム分析用のAmazonKinesisData Firehose; パーソナライズされた製品の推奨事項を生成するAWSGlue

**C.** データカタログとしてのAWS Glue; 履歴データインサイトのための

AmazonKinesisDataStreamsとAmazonKinesisData Analytics; クリックストリーム分析のためにAmazonESに配信するAmazonKinesisData Firehose; パーソナライズされた製品の推奨事項を生成するAmazonEMR

**D.** データカタログとしてのAmazon Athena; 履歴データインサイトのための

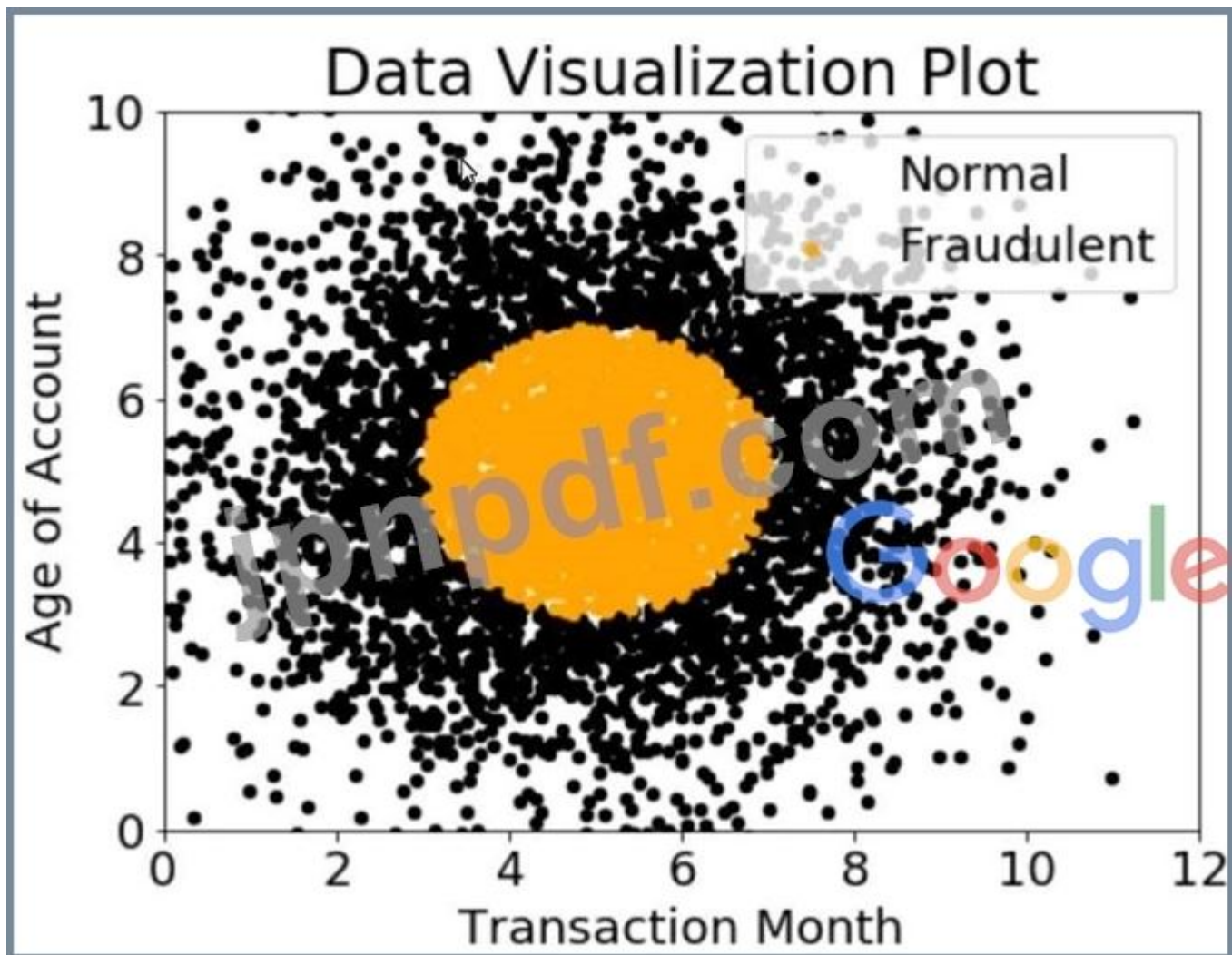
AmazonKinesisDataStreamsとAmazonKinesisData Analytics; クリックストリーム分析用のAmazonDynamoDBストリーム。パーソナライズされた製品の推奨事項を生成するAWSGlue

**Answer: A** ([メッセージを残す](#))

説明

最新問題: 61

企業は、ユーザーの行動を不正または通常のいずれかに分類したいと考えています。内部調査に基づいて、機械学習スペシャリストは、アカウントの年齢とトランザクション月の2つの機能に基づいてバイナリ分類器を構築したいと考えています。これらの機能のクラス分布は、提供された図に示されています。



この情報に基づいて、どのモデルが最も精度が高いでしょうか？

- A. 非線形カーネルを備えたサポートベクターマシン (SVM)
- B. スケーリングされた指数線形単位 (ELU) を使用した長短期記憶 (LSTM) モデル
- C. ロジスティック回帰
- D. tanh活性化関数を備えた単一パーセプトロン

Answer: ([解答を表示する](#))

有効な **Professional-Machine-Learning-Engineer** 問題集は GoShiken.com が提供された合格しやすい Professional-Machine-Learning-Engineer 試験問題集！ GoShiken.com が最新の **Professional-Machine-Learning-Engineer** 試験問題集を提供しています。GoShiken.com Professional-Machine-Learning-Engineer 試験問題は最新で、解答が正確でございます。最新の GoShiken.com Professional-Machine-Learning-Engineer 問題集をゲットする人はこちら：  
<https://www.goshiken.com/Google/Professional-Machine-Learning-Engineer-mondaishu.html>  
(30030%OFF問題集溶と正解付きで 30%w 特別割引コード: **Freepdfdumps**)

あなたはソーシャルメディア会社で働いています。投稿された画像に車が含まれているかどうかを検出する必要があります。各トレーニング例は、正確に1つのクラスのメンバーです。オブジェクト検出ニューラルネットワークをトレーニングし、評価のためにモデルバージョンをAIPlatformPredictionにデプロイしました。デプロイする前に、評価ジョブを作成し、それをAIPlatformPredictionモデルバージョンに添付しました。精度がビジネス要件で許可されているよりも低いことに気づきました。精度を上げるために、モデルの最終層のソフトマックスしきい値をどのように調整する必要がありますか？

- A. リコールを増やす
- B. フォールスネガティブの数を減らします
- C. 誤検知の数を増やします
- D. リコールを減らします。

**Answer: B** ([メッセージを残す](#))

#### 最新問題: 63

大企業内のデータサイエンスチームは、AmazonSageMakerノートブックを使用してAmazonS3バケットに保存されているデータにアクセスします。ITセキュリティチームは、インターネット対応のノートブックインスタンスがセキュリティの脆弱性を生み出し、インスタンスで実行されている悪意のあるコードがデータのプライバシーを侵害する可能性があることを懸念しています。同社は、すべてのインスタンスがインターネットアクセスのない安全なVPC内にとどまるように義務付けており、データ通信トラフィックはAWSネットワーク内にとどまる必要があります。データサイエンスチームは、これらの要件を満たすためにノートブックインスタンスの配置をどのように構成する必要がありますか？

- A. AmazonSageMakerノートブックをVPCのプライベートサブネットに関連付けます。VPCにS3VPCエンドポイントとAmazonSageMakerVPCエンドポイントが接続されていることを確認します。
- B. AmazonSageMakerノートブックをVPCのプライベートサブネットに関連付けます。AmazonSageMakerエンドポイントとS3バケットを同じVPC内に配置します。
- C. AmazonSageMakerノートブックをVPCのプライベートサブネットに関連付けます。IAMポリシーを使用して、AmazonS3およびAmazonSageMakerへのアクセスを許可します。
- D. AmazonSageMakerノートブックをVPCのプライベートサブネットに関連付けます。VPCにNATゲートウェイと関連するセキュリティグループがあり、AmazonS3とAmazonSageMakerへのアウトバウンド接続のみを許可していることを確認します。

**Answer: D** ([メッセージを残す](#))

#### 最新問題: 64

モデルのトレーニングと予測の前に、Dataflowを使用して生データを前処理する需要予測パイプラインが本番環境にあります。前処理中に、BigQueryに保存されているデータにZスコアの正規化を採用し、BigQueryに書き戻します。新しいトレーニングデータは毎週追加されます。計算時間と手動による介入を最小限に抑えることで、プロセスをより効率的にしたいと考えています。あなたは何をするべきか？

- A. TensorFlowのFeatureColumnAPIでnormalizer\_fn引数を使用します
- B. GoogleKubernetesEngineを使用してデータを正規化する
- C. BigQueryのDataprocコネクタを使用してApacheSparkでデータを正規化する
- D. BigQueryで使用するために正規化アルゴリズムをSQLに変換します

**Answer:** [\(解答を表示する\)](#)

**最新問題: 65**

機械学習チームは、AmazonSageMakerで独自のトレーニングアルゴリズムを実行します。トレーニングアルゴリズムには外部資産が必要です。チームは、独自のアルゴリズムコードとアルゴリズム固有のパラメーターの両方をAmazonSageMakerに送信する必要があります。

チームがAmazonSageMakerでカスタムアルゴリズムを構築するために使用するサービスの組み合わせは何ですか？

(2つ選択してください。)

- A. Amazon S3
- B. Amazon ECS
- C. AWSシークレットマネージャー
- D. AWS CodeStar
- E. Amazon ECR

**Answer:** A,E ([メッセージを残す](#))

**最新問題: 66**

ある会社は、入手可能な過去の販売データに基づいて住宅の販売価格を予測したいと考えています。会社のデータセットのターゲット変数は販売価格です。機能には、ロットサイズ、居住エリアの測定値、非居住エリアの測定値、寝室の数、浴室の数、築年数、郵便番号などのパラメーターが含まれます。同社は、多変数線形回帰を使用して住宅販売価格を予測したいと考えています。

分析に関係のない機能を削除し、モデルの複雑さを軽減するために、機械学習スペシャリストはどの手順を実行する必要がありますか？

- A. 特徴のヒストグラムをプロットし、それらの標準偏差を計算します。分散の大きい機能を削除します。
- B. ターゲット変数に対してすべての機能の相関チェックを実行します。ターゲット変数の相関スコアが低い機能を削除します。
- C. データセットとそれ自体の相関関係を示すヒートマップを作成します。相互相関スコアが低い機能を削除します。
- D. 特徴のヒストグラムをプロットし、それらの標準偏差を計算します。分散の少ない機能を削除します。

**Answer:** B ([メッセージを残す](#))

**Valid Professional-Machine-Learning-Engineer Dumps** shared by GoShiken.com for Helping Passing Professional-Machine-Learning-Engineer Exam! GoShiken.com now offer the **newest Professional-Machine-Learning-Engineer exam dumps**, the GoShiken.com Professional-Machine-Learning-Engineer exam **questions have been updated** and **answers have been corrected** get the **newest** GoShiken.com Professional-Machine-Learning-Engineer dumps with Test Engine here: <https://www.goshiken.com/Google/Professional-Machine-Learning-Engineer-mondaishu.html> (**300** Q&As Dumps, **30%OFF** Special Discount: **Freepdfdumps**)