

Aula 8: Avaliação Crítica, Métricas e Fechamento do Módulo

Aléssio Miranda Júnior

Agosto - 2026/2

Capítulo em Revisão

Este conteúdo está sendo revisado. Algumas informações podem estar preliminares, conter erros de formatação ou sofrer alterações na versão final.

O que você verá neste capítulo

Neste fechamento do primeiro arco, consolidamos os dois últimos pilares da Engenharia de IA: a **Avaliação Crítica** de modelos e o **Deploy** (implantação em produção). Demonstraremos por que a Acurácia é uma métrica enganosa para dados desbalanceados, formalizaremos a **Matriz de Confusão** e as métricas derivadas (Precision, Recall e F1-Score), e introduziremos os conceitos fundamentais de serialização de modelos e implantação via API REST, preparando o terreno para o Trabalho Prático 1 (TP1).

1 A Ilusão da Acurácia

Neste arco, você já sabe construir, treinar, regularizar e realizar inferências com Redes Neurais. O último pilar fundamental é a **Avaliação Crítica**: a capacidade de diagnosticar se o seu modelo é genuinamente bom ou se está enganando a equipe com números superficiais.

Historicamente, olhamos para a *Acurácia* (taxa de acertos gerais) como a métrica definitiva. Mas imagine que você está construindo uma IA para detectar câncer, e em 100 pacientes, apenas 1 tem a doença. Se a sua IA for completamente inútil e chutar “Saudável” para absolutamente todo mundo, ela acertará 99 pacientes e errará 1. Ela terá uma **Acurácia de 99%**, mas será um **fracasso retumbante** na vida real — o único paciente doente não será tratado.

Quando as classes do problema não são equilibradas (ex: 95% saudáveis, 5% doentes; ou 99% transações legítimas, 1% fraudes), a acurácia é uma **mentira estatística**. Um modelo que sempre prediz a classe majoritária terá acurácia altíssima sem nunca ter aprendido nada. Esse é um dos erros mais perigosos e recorrentes em projetos reais de Machine Learning.

Para avaliar modelos em ambientes corporativos e acadêmicos com rigor, “quebramos” os acertos e erros em uma tabela chamada **Matriz de Confusão** (*Confusion Matrix*). Ela revela exatamente *como* o modelo está errando:

		Realidade (Gabarito)	
		Positivo	Negativo
Previsão do Modelo	Positivo	Verdadeiro Positivo (TP)	Falso Negativo (FN)
	Negativo	Falso Positivo (FP)	Verdadeiro Negativo (TN)

- **Verdadeiros Positivos (TP):** Tinha câncer e a IA acertou. **Acerto.**
- **Verdadeiros Negativos (TN):** Estava saudável e a IA acertou. **Acerto.**
- **Falsos Positivos (FP):** Estava saudável, mas a IA disse que tinha câncer. **Alarme Falso.**
- **Falsos Negativos (FN):** Tinha câncer, mas a IA disse que estava saudável. **Falha Fatal.**

Considere uma IA de detecção de fraude que processou 1.000 transações bancárias:

	Real: 1	Real: 0
Prev: 1	TP = 15 Achou a fraude	FN = 5 Escaparam!
Prev: 0	FP = 30 Alarmes falsos	TN = 950 Legítimas OK

Figura 2: Matriz de Confusão preenchida com dados de detecção de fraude bancária. A Acurácia seria $\frac{15+950}{1000} = 96,5\%$, mas das 20 fraudes reais, 5 escaparam (FN=5).

A partir desses valores, calcularemos as três métricas que realmente importam.

3 Precision, Recall e F1-Score

A partir dos quatro quadrantes da Matriz de Confusão, extraímos as três métricas definitivas que substituem a acurácia em qualquer projeto sério:

3.1 Precision (Precisão)

$$\text{Precision} = \frac{TP}{TP + FP} \quad (1)$$

De todas as vezes que a IA **gritou “É Positivo!”**, quantas vezes ela estava certa? A Precision penaliza os **alarmes falsos** (FP). Uma Precision alta significa: “quando a IA diz que é câncer, ela tem razão”.

3.2 Recall (Revocação / Sensibilidade)

$$\text{Recall} = \frac{TP}{TP + FN} \quad (2)$$

De todos os positivos que **realmente existiam** na base, quantos a IA conseguiu encontrar? O Recall penaliza as **omissões cegas** (FN). Um Recall alto significa: “a IA não deixa escapar doentes”.

3.3 F1-Score (Média Harmônica)

$$F1 = 2 \cdot \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (3)$$

É a média harmônica entre Precision e Recall. Se o F1 for alto, significa que a IA é **balanceada**: não é nem “medrosa” (que chuta tudo como positivo para não perder nenhum) nem “cautelosa demais” (que só diz positivo quando tem 100% de certeza, deixando muitos escaparem).

A tabela a seguir ilustra por que a média harmônica é superior à média aritmética simples:

Cenário	Precision	Recall	F1-Score	Diagnóstico
Balanceado	0.85	0.85	0.85	Saudável
Medroso (só acerta quando aposta)	1.00	0.20	0.33	Omisso: deixa escapar
Agressivo (aposta em tudo)	0.10	1.00	0.18	Caótico: alarmes falsos

Tabela 1: Comparação de cenários: o F1-Score (média harmônica) penaliza severamente desequilíbrios que a média aritmética mascararia.

3.4 O Trade-off Precision vs. Recall

Existe um dilema inerente entre Precision e Recall que pode ser visualizado como uma gangorra:

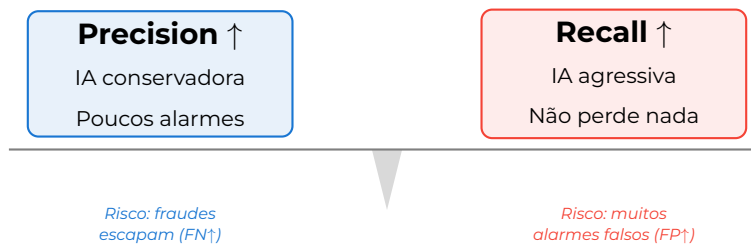


Figura 3: A gangorra Precision–Recall: melhorar uma métrica frequentemente deteriora a outra. O engenheiro escolhe o equilíbrio baseado no custo do erro no domínio.

A escolha de qual métrica priorizar depende inteiramente do **custo do erro** no domínio de aplicação:

Domínio	Prioridade	Justificativa
Diagnóstico Médico	Recall	É melhor gerar alarmes falsos (FP) do que deixar um paciente doente sem tratamento (FN).
Filtro de Spam	Precision	Um e-mail legítimo classificado como spam (FP) é inaceitável para o usuário.
Detecção de Fraude	F1	Equilibrar: não perder fraude nem bloquear clientes injustamente.
Carro Autônomo	Recall	Não identificar um pedestre (FN) pode causar um acidente fatal.

Tabela 2: Guia prático: quando priorizar Precision, Recall ou F1-Score, dependendo do custo do erro.

► Prática: O Classification Report do Scikit-Learn

O Scikit-Learn consolida todas as métricas em um relatório padronizado:

```
1 from sklearn.metrics import classification_report
2
3 # y_test: rotulos reais | y_pred: previsoes do modelo
4 y_pred = model.predict(X_test)
5 y_pred_classes = (y_pred > 0.5).astype(int)
6
7 print(classification_report(y_test, y_pred_classes))
```

A saída será uma tabela formatada:

	precision	recall	f1-score	support
0	0.92	0.88	0.90	150
1	0.85	0.90	0.87	100
accuracy			0.89	250
macro avg	0.88	0.89	0.88	250
weighted avg	0.89	0.89	0.89	250

4 Serialização e Persistência de Modelos

Treinar uma rede neural pode levar horas ou dias. Não é viável retreinar o modelo toda vez que quisermos fazer uma previsão. O Keras permite **salvar** o modelo treinado (pesos + arquitetura) em um arquivo e **carregá-lo** posteriormente para inferência:

```
1 # Salvar o modelo completo
2 model.save('meu_modelo.keras')
3
4 # Em outro script, carregar e usar diretamente
5 from tensorflow.keras.models import load_model
6 modelo_carregado = load_model('meu_modelo.keras')
7
8 # Inferencia imediata
9 previsao = modelo_carregado.predict(nova_amostra)
```

O fluxo completo de serialização e uso em produção pode ser visualizado como um pipeline:



Figura 4: Pipeline de serialização: o treinamento (custoso) acontece uma única vez. O modelo salvo pode ser carregado e usado para inferência instantânea em qualquer ambiente de produção.

◇ Informação

O formato `.keras` (antigo `.h5`) armazena tanto a arquitetura quanto os pesos treinados. Isso permite que o modelo seja distribuído como um “executável” de IA — qualquer máquina com Keras instalado pode carregar e usar o modelo sem precisar dos dados de treinamento originais.

5 Deploy: Da Máquina Local para a Produção

O **Deploy** (implantação) é a etapa que separa o protótipo acadêmico do produto real. Um modelo que só funciona no Jupyter Notebook do criador não tem valor de mercado. A indústria exige que a IA seja acessível via rede, consumível por qualquer aplicação cliente.

A abordagem mais comum é encapsular o modelo em uma **API REST**:

```
1 from flask import Flask, request, jsonify
2 from tensorflow.keras.models import load_model
3 import numpy as np
4
5 app = Flask(__name__)
6 model = load_model('meu_modelo.keras')
7
8 @app.route('/predict', methods=['POST'])
9 def predict():
10     dados = request.json['features']
11     X = np.array(dados).reshape(1, -1)
12     previsao = model.predict(X).tolist()
13     return jsonify({'prediction': previsao})
14
15 if __name__ == '__main__':
16     app.run(host='0.0.0.0', port=5000)
```

Com isso, qualquer sistema — seja um aplicativo mobile, um site web ou outro microserviço — pode enviar dados via HTTP e receber a previsão da IA em formato JSON.

6 O Trabalho Prático 1 (TP1)

A consolidação de todos os conhecimentos deste primeiro arco ocorrerá agora. No seu primeiro Trabalho Prático (TP1), você tem a liberdade de buscar um problema real na plataforma **Kaggle** (ex: *Detecção de fraudes financeiras*, *Previsão de evasão escolar*, *Diagnóstico de doenças cardíacas*).

Você deverá:

- 1 **Tratar os dados** com Pandas (limpeza, análise exploratória, tratamento de valores ausentes).

- 2 Normalizar as features** com `StandardScaler` ou `MinMaxScaler`.
- 3 Dividir em Treino/Validação/Teste** com `train_test_split`.
- 4 Criar a arquitetura** com Keras (camadas Dense, Dropout, ativações).
- 5 Impedir a memorização** com Dropout e `EarlyStopping`.
- 6 Avaliar com rigor:** extrair o `classification_report` completo (Precision, Recall e F1).
- 7 Serializar e servir:** salvar o modelo e disponibilizar uma API REST básica.

★ Checkpoint do Projeto: Entrega do TP1

O TP1 será entregue via repositório GitLab com CI/CD configurado. A pipeline deverá:

- Executar os testes unitários (PyTest).
- Treinar o modelo e gerar o `classification_report`.
- Salvar o modelo serializado como artefato da pipeline.

A nota será composta por: qualidade do código (30%), rigor da avaliação (40%) e apresentação dos resultados (30%).

7 Conclusão

Você concluiu o “Hello World” da Inteligência Artificial. Com as bases numéricas do Perceptron, o entendimento matemático da otimização via Descida do Gradiente, a poderosa abstração do Keras, as defesas contra o Overfitting, a visão computacional via CNNs e a avaliação crítica com métricas profissionais, você está equipado para transicionar para o próximo grande degrau.

No Módulo 2, pararemos de processar planilhas e números e começaremos a processar a **Linguagem Humana**. Exploraremos como transformar palavras, frases e conceitos abstratos em vetores matemáticos (*Embeddings*) que computadores conseguem processar, abrindo as portas para o universo dos *Large Language Models* (LLMs) e da IA Generativa.

✓ Takeaways

- A **Acurácia** é uma métrica enganosa em dados desbalanceados — um modelo que sempre prediz a classe majoritária pode ter 99% sem ter aprendido nada.
- A **Matriz de Confusão** decompõe os erros em quatro quadrantes: TP, TN, FP e FN.
- **Precision** ($\frac{TP}{TP+FP}$) mede a qualidade dos alarmes; **Recall** ($\frac{TP}{TP+FN}$) mede a cobertura dos positivos reais.
- O **F1-Score** (média harmônica) só é alto quando ambas as métricas são altas — é a métrica definitiva para dados desbalanceados.
- **Serializar** modelos com `model.save()` e `load_model()` é obrigatório para deploy em produção.
- O **classification_report** do Scikit-Learn consolida Precision, Recall e F1 num relatório padronizado.

📎 Questões: Autoavaliação

- 1 **[Direta]** Uma IA de diagnóstico médico apresenta: TP=80, FP=20, FN=10, TN=890. Calcule Precision, Recall e F1-Score. A IA é segura para uso clínico?
- 2 **[Direta]** Explique por que a média **harmônica** (F1) é preferida sobre a média aritmética simples para combinar Precision e Recall. Dê um exemplo numérico.
- 3 **[Direta]** Um modelo de filtro de spam tem Precision=0.99 e Recall=0.60. Interprete esses números em termos práticos para o usuário de e-mail.
- 4 **[Reflexiva]** Em um sistema de detecção de fraudes, é pior um **Falso Positivo** (bloquear uma transação legítima) ou um **Falso Negativo** (deixar uma fraude escapar)? A resposta muda se o sistema for de diagnóstico de câncer?
- 5 **[Reflexiva]** Você está entregando um modelo para um cliente. Ele olha a Acurácia de 97% e fica satisfeito. Que perguntas você faria antes de concordar que o modelo é bom?

Lab. 08: Métricas de Avaliação e Serialização

Capítulo em Revisão

Este conteúdo está sendo revisado. Algumas informações podem estar preliminares, conter erros de formatação ou sofrer alterações na versão final.

★ Objetivo do Laboratório

Este é o **último laboratório do Módulo 1** — “Hello World das Redes Neurais”. Aqui você enfrentará a situação mais traiçoeira do Machine Learning: um dataset **propositalmente desbalanceado** (95% vs 5%) onde a Acurácia **mente**. Você aprenderá a avaliar de verdade com **Precision, Recall e F1-Score**, plotará a **Matriz de Confusão** e **serializará** o modelo para uso em produção. Esta atividade consolida todas as técnicas do módulo e prepara o terreno para o **TP1**.

8 Parte 1: O Problema H — Estilo Maratona

Problema H: O Sentinela de Fraudes do NeoBank

Contexto: O *NeoBank*, um banco digital com 2 milhões de clientes, processa 100.000 transações por dia. Dessas, apenas $\approx 5\%$ são fraudulentas — mas cada fraude não detectada custa em média R\$ 15.000,00 ao banco. A equipe de ciência de dados treinou um modelo que atinge **99% de acurácia** e comemorou. Até que o CTO perguntou: “das 5.000 fraudes de ontem, quantas vocês realmente pegaram?” — silêncio total.

Sua missão é construir um classificador para este cenário desbalanceado e demonstrar que a **Acurácia é uma métrica enganosa**. O relatório final deve usar métricas que realmente importam.

Modelo de Entrada (X): Vetor com 15 features de transação:

- $x_1 \dots x_8$: Features informativas (valor, horário, localização, etc.)
- $x_9 \dots x_{12}$: Features redundantes (correlações internas)
- $x_{13} \dots x_{15}$: Features ruidosas (dados espúrios)
- **Desbalanceamento:** 95% legítimas (classe 0) vs 5% fraudes (classe 1)

Modelo de Saída: Três artefatos obrigatórios:

- 1 Relatório `classification_report` com Precision, Recall e F1 por classe
- 2 Gráfico `matriz_confusao.png` com mapa de calor
- 3 Modelo serializado `modelo_fraude.keras`

Critérios de Aprovação:

- 1 F1-Score da classe Fraude $> 30\%$ (acima de um classificador trivial)
- 2 Modelo serializado carrega e produz previsões idênticas
- 3 Previsões binárias válidas (0 ou 1)
- 4 `classification_report` contém ambas as classes

9 Parte 2: A Armadilha Visual da Acurácia

Antes de codar, entenda visualmente **por que a Acurácia engana**. Imagine 1.000 transações processadas pelo modelo:

99% Acurácia — 990 transações classificadas corretamente

10

O modelo aprendeu a “chutar” a classe majoritária (legítima) para **todas** as transações.

Esses 10 “erros”
eram **TODAS**
as fraudes reais!

Um modelo que **sempre diz “legítima”** atinge 95% de acurácia neste dataset — mas não detecta **nenhuma fraude**. A Acurácia é inútil aqui.

10 Parte 3: Dataset Desbalanceado (8 min)

Crie o arquivo `src/avaliacao_lab.py`:

```
1 import numpy as np
2 import os
3 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '2'
4
5 import matplotlib
6 matplotlib.use('Agg')
7 import matplotlib.pyplot as plt
8 from sklearn.datasets import make_classification
9 from sklearn.model_selection import train_test_split
10 from sklearn.preprocessing import StandardScaler
11 from sklearn.metrics import (classification_report,
12                             confusion_matrix, ConfusionMatrixDisplay)
13 from tensorflow.keras.models import Sequential, load_model
14 from tensorflow.keras.layers import Dense, Dropout
15 from tensorflow.keras.callbacks import EarlyStopping
16
17 # Dataset DESBALANCEADO: 95% classe 0, 5% classe 1
18 X, y = make_classification(
19     n_samples=2000, n_features=15,
20     n_informative=8, n_redundant=4,
21     weights=[0.95, 0.05], # 95% vs 5% !!
22     random_state=42, flip_y=0.02
23 )
24
25 print(f"Classe 0 (legítima): {np.sum(y==0)}")
26 print(f"Classe 1 (fraude): {np.sum(y==1)}")
27 print(f"Proporção: {np.sum(y==1)/len(y)*100:.1f}% fraudes")
28
29 # Separar com stratify para manter proporção
30 X_train, X_test, y_train, y_test = train_test_split(
31     X, y, test_size=0.2, random_state=42, stratify=y)
32
33 scaler = StandardScaler()
34 X_train_norm = scaler.fit_transform(X_train)
```

```
35 X_test_norm = scaler.transform(X_test)
```

• Teoria: Por que stratify=y?

O parâmetro stratify=y garante que a proporção 95%/5% seja mantida **tanto no treino quanto no teste**. Sem ele, em dados tão desbalanceados, é possível que o teste fique com zero exemplos da classe minoritária — tornando a avaliação impossível.

11 Parte 4: Modelo com Regularização (8 min)

```
1 def criar_modelo():
2     """Modelo para classificação desbalanceada."""
3     model = Sequential()
4     model.add(Dense(64, activation='relu', input_shape=(15,)))
5     model.add(Dropout(0.3))
6     model.add(Dense(32, activation='relu'))
7     model.add(Dropout(0.2))
8     model.add(Dense(1, activation='sigmoid'))
9
10    model.compile(
11        optimizer='adam',
12        loss='binary_crossentropy',
13        metrics=['accuracy']
14    )
15    return model
16
17 modelo = criar_modelo()
18
19 vigilante = EarlyStopping(
20     monitor='val_loss', patience=10,
21     restore_best_weights=True)
22
23 historico = modelo.fit(
24     X_train_norm, y_train,
25     epochs=200, batch_size=32,
26     validation_split=0.2,
27     callbacks=[vigilante],
28     verbose=0)
29
30 print(f"Parou na época {len(historico.epoch)}")
```

12 Parte 5: Desmascarando a Acurácia (8 min)

Agora, demonstre a armadilha e revele as métricas reais:

```

1 # Avaliação com Acurácia (ENGANOSA)
2 _, acc = modelo.evaluate(X_test_norm, y_test, verbose=0)
3 print(f"\nAcurácia geral: {acc:.2%}")
4 print("Parece ótimo, certo? Mas espere...")
5
6 # Avaliação REAL com classification_report
7 y_pred = modelo.predict(X_test_norm, verbose=0)
8 y_pred_classes = (y_pred > 0.5).astype(int).flatten()
9
10 print("\n" + "="*50)
11 print("RELATÓRIO DE MÉTRICAS (A VERDADE)")
12 print("="*50)
13 report = classification_report(
14     y_test, y_pred_classes,
15     target_names=['Legítima (0)', 'Fraude (1)'])
16 print(report)

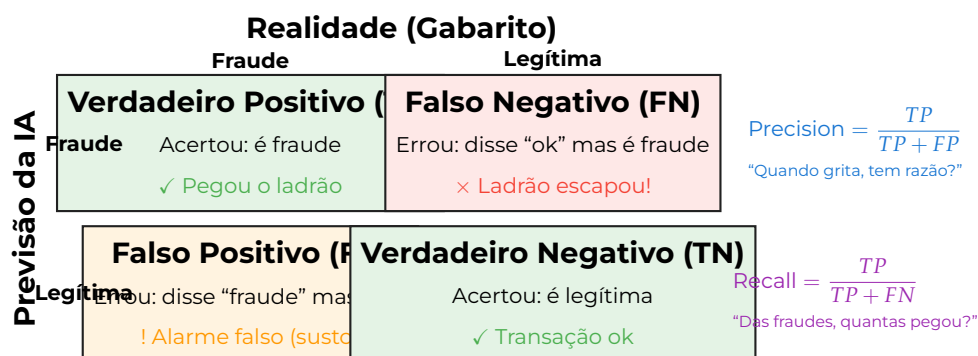
```

△ Importante: O que Observar no Relatório

A acurácia geral será >95%, mas observe o **Recall da classe 1 (Fraude)**. Se for baixo (<50%), significa que mais da metade das fraudes **escaparam** — exatamente o cenário catastrófico que o CTO temia. O F1-Score é o veredicto final que equilibra Precision e Recall.

A. Anatomia da Matriz de Confusão

O diagrama abaixo explica visualmente os 4 quadrantes que você vai plotar:



13 Parte 6: Matriz de Confusão Visual (8 min)

Plote a Matriz de Confusão como um mapa de calor:

```

1 def plotar_matriz_confusao(y_test, y_pred_classes):
2     """Gera a Matriz de Confusão visual."""
3     cm = confusion_matrix(y_test, y_pred_classes)
4     disp = ConfusionMatrixDisplay(
5         confusion_matrix=cm,

```

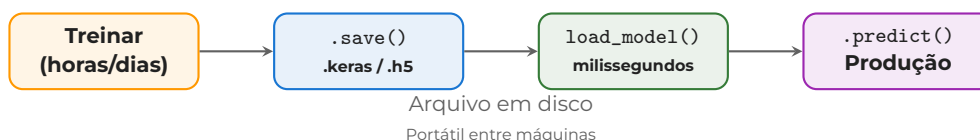
```

6     display_labels=['Legítima', 'Fraude'])
7
8     fig, ax = plt.subplots(figsize=(6, 5))
9     disp.plot(cmap='Blues', ax=ax, colorbar=False)
10    ax.set_title('Matriz de Confusão')
11    plt.tight_layout()
12    plt.savefig('matriz_confusao.png', dpi=150)
13    plt.close()
14    print("Gráfico salvo: matriz_confusao.png")
15    return cm

```

14 Parte 7: Serialização do Modelo (8 min)

O modelo treinado custou tempo e GPU — não queremos retreiná-lo toda vez. A serialização salva o modelo completo (arquitetura + pesos + otimizador) em disco:



```

1 # Salvar o modelo completo
2 modelo.save('modelo_fraude.keras')
3 print("Modelo salvo: modelo_fraude.keras")
4
5 # Recarregar em memória (simulando outro script)
6 modelo_carregado = load_model('modelo_fraude.keras')
7
8 # Verificar que as previsões são idênticas
9 preds_original = modelo.predict(X_test_norm[:5], verbose=0)
10 preds_carregado = modelo_carregado.predict(
11     X_test_norm[:5], verbose=0)
12
13 assert np.allclose(preds_original, preds_carregado), \
14     "Os modelos deveriam produzir previsões idênticas!"
15 print("Verificação OK: modelo salvo e carregado.")

```

• Teoria: .keras vs .h5 — Qual Formato Usar?

Formato	Vantagens	Uso
.keras (novo)	Padrão oficial, suporta tudo	Recomendado (TF 2.16+)
.h5 (legado)	Compatível com código antigo	Projetos legados
SavedModel/	Diretório TF Serving	Deploy em produção

15 Parte 8: Bloco Principal e Execução (5 min)

Monte o bloco principal com o relatório completo:

```

1  if __name__ == "__main__":
2      cm = plotar_matriz_confusao(y_test, y_pred_classes)
3
4      # Extrair os 4 quadrantes
5      tn, fp, fn, tp = cm.ravel()
6      precision = tp / (tp + fp) if (tp + fp) > 0 else 0
7      recall = tp / (tp + fn) if (tp + fn) > 0 else 0
8      f1 = (2*precision*recall/(precision+recall)
9           if (precision+recall) > 0 else 0)
10
11     print(f"\n{'='*50}")
12     print(f"  MÉTRICAS REAIS (Classe Fraude)")
13     print(f"  TP={tp}  FP={fp}  FN={fn}  TN={tn}")
14     print(f"  Precision: {precision:.2%}")
15     print(f"  Recall:    {recall:.2%}")
16     print(f"  F1-Score:   {f1:.2%}")
17     print(f"{'='*50}")

```

16 Parte 9: Garantia de Qualidade — Testes CI/CD (5 min)

Os testes já estão prontos no arquivo tests/test_avaliacao.py:

```

1  import numpy as np
2  import os
3  from sklearn.metrics import f1_score, classification_report
4  from tensorflow.keras.models import load_model
5  from src.avaliacao_lab import (
6      modelo, X_test_norm, y_test, y_pred_classes)
7
8  def test_f1_minimo():
9      """F1-Score da fraude deve ser > 0.30."""
10     f1 = f1_score(y_test, y_pred_classes, pos_label=1)
11     assert f1 > 0.30, \
12         f"F1 da fraude = {f1:.2f}, muito baixo"
13
14  def test_modelo_serializado():
15     """Modelo deve ter sido salvo corretamente."""
16     assert os.path.exists('modelo_fraude.keras'), \
17         "Modelo não foi serializado!"
18     m = load_model('modelo_fraude.keras')
19     preds = m.predict(X_test_norm[:5], verbose=0)
20     assert preds.shape == (5, 1)
21
22  def test_previsoes_validas():

```

```

23     """Previsões devem ser 0 ou 1."""
24     assert set(np.unique(y_pred_classes)).issubset({0, 1})
25
26 def test_report_existe():
27     """0 classification_report deve conter ambas as classes."""
28     report = classification_report(
29         y_test, y_pred_classes, output_dict=True)
30     assert '0' in report and '1' in report
31     assert 'precision' in report['1']

```

- 1 Rode os testes no terminal: `pytest tests/test_avaliacao.py -v`.
- 2 Se a tela exibir **4 passed** em verde, sua implementação está correta.

17 Parte 10: Commit e Push (5 min)

Com todos os testes passando localmente, efetue o envio para o GitLab:

```

1 git add src/avaliacao_lab.py modelo_fraude.keras matriz_confusao.png
2 git commit -m "Lab 08: Métricas reais + Serialização (Fechamento Módulo 1)"
3 git push origin main

```

O GitLab CI interceptará seu push e executará o pytest automaticamente. O **Selo Verde** confirma sua pontuação.

△ Importante: Se o Pipeline Falhar

Os erros mais comuns são: (1) `modelo_fraude.keras` não existir no diretório de trabalho do CI (verifique o path); (2) `threshold de F1 > 0.30` não atingido — tente aumentar as épocas ou ajustar o `patience`; (3) importar de `avaliacao_lab` sem o prefixo `src..`

Conclusão: A Jornada Completa do Módulo 1

Em 8 aulas e 8 labs, você percorreu a jornada completa de um engenheiro de ML:



✓ Takeaways

- Você treinou um classificador em dados **desbalanceados** (95% vs 5%) e demonstrou que a **Acurácia engana** — um modelo trivial atinge 95% sem detectar nenhuma fraude.
- Extraíu **Precision** (“quando grita, tem razão?”), **Recall** (“das fraudes, quantas pegou?”) e **F1-Score** (veredicto final) com o `classification_report`.
- Plotou a **Matriz de Confusão** e analisou os 4 quadrantes: TP, TN, FP e FN — entendendo que o **custo do erro** depende do domínio.
- **Serializou** o modelo com `.save()` e comprovou que o modelo carregado produz previsões idênticas — pronto para deploy.
- Completou a **maratona de 8 problemas** (A–H), passando do Perceptron solitário até métricas profissionais de produção.

◇ Informação

Fechamento do Módulo 1 — “Hello World das Redes Neurais”.

Em 8 aulas, você percorreu a jornada completa: do Perceptron solitário (Lab A) até a avaliação crítica com métricas profissionais e serialização (Lab H), passando por MLPs, Gradiente Descendente, Keras, Overfitting, Regularização, CNNs e Transfer Learning.

No Módulo 2 — “A Revolução da Linguagem e dos Vetores”, abandonaremos números e pixels e entraremos no domínio da **linguagem humana**: como transformar palavras em vetores (*Embeddings*), como funcionam os *Transformers* que revolucionaram a IA, e como chegamos aos *Large Language Models* (LLMs) como GPT e Gemini. A pergunta que guiará o módulo: “como o computador entende que ‘rei’ está para ‘rainha’ assim como ‘homem’ está para ‘mulher’?”