

Lab. 04: Ecossistema Keras

Aléssio Miranda Júnior

Agosto - 2026/2

Lab. 04: Ecossistema Keras

★ Objetivo do Laboratório

Nos Labs 01–03, você construiu neurônios, redes MLP e loops de treinamento com NumPy puro — traduzindo cada equação em código artesanal. Hoje você abandona a oficina e entra na **fábrica**: o ecossistema Keras/TensorFlow. A missão é construir o **mesmo pipeline** de Machine Learning, mas agora com ferramentas industriais — pandas para dados, StandardScaler para normalização e Sequential para a rede. Ao final, a pergunta que o Lab responde é: *“se o Keras faz tudo sozinho, por que precisei aprender na mão?”*

1 Parte 1: O Problema D — Estilo Maratona

Problema D: A Previsão de Churn Bancário

Contexto: O banco *DigitalBank* está perdendo clientes a um ritmo alarmante. A diretoria contratou sua equipe de Engenharia de IA para construir um modelo preditivo que identifique, **antes** do cancelamento, quais clientes têm maior risco de abandonar o banco (*churn*). Os dados históricos já foram coletados — sua missão é construir o **pipeline completo** de classificação binária.

Modelo de Entrada (X): Vetor com 2 features numéricas por cliente:

- x_1 : Idade (inteiro entre 18 e 70)
- x_2 : Salário Mensal em R\$ (inteiro entre 2.000 e 350.000)

Regra de Negócio (Rótulo): Um cliente cancelou ($y = 1$) se, e somente se, ambas as condições forem satisfeitas simultaneamente:

$$y = \begin{cases} 1 & \text{se } x_1 > 40 \text{ E } x_2 > 50.000 \\ 0 & \text{caso contrário} \end{cases}$$

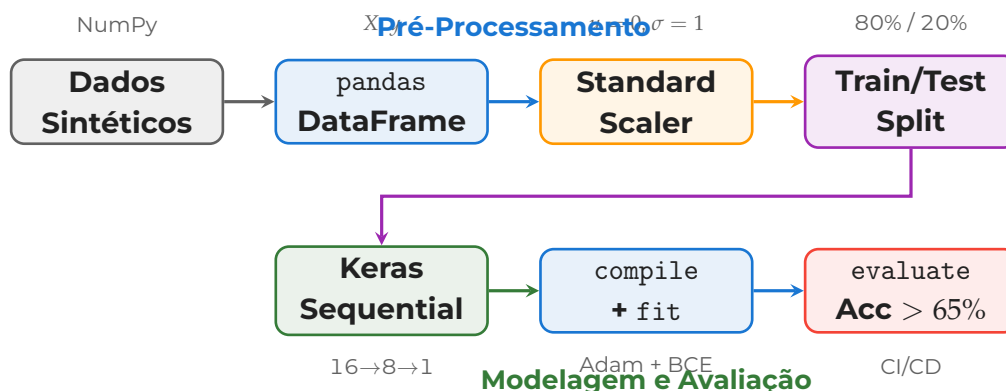
Modelo de Saída ($\hat{y}$): Probabilidade de churn via Sigmoid $\in (0,1)$:

- $\hat{y} > 0.5 \rightarrow$ Cliente em risco de churn
- $\hat{y} \leq 0.5 \rightarrow$ Cliente seguro

Critério de Aprovação: Acurácia no conjunto de teste $> 65\%$.

A. O Pipeline Industrial

Diferente dos Labs anteriores onde cada componente era programado à mão, agora usaremos ferramentas especializadas para cada etapa. Observe o pipeline completo que você vai construir:



• Teoria: Do Artesanal ao Industrial

Compare mentalmente com os Labs anteriores: no Lab 01, você escreveu a função degrau; no Lab 02, a classe `CamadaDensa` com forward pass; no Lab 03, o loop de épocas com update rule. Agora, a linha `model.fit(X, y, epochs=80)` executa **tudo isso internamente** — Forward Pass, Loss, Backpropagation e Update Rule — com aceleração em GPU. O que mudou não é a matemática, é a **escala**.

2 Parte 2: Gerando e Estruturando os Dados (8 min)

Antes de codificarmos, certifique-se de que seu arquivo `requirements.txt` está atualizado. Como o ecossistema Python evolui rápido, misturar versões pode causar quebras no ambiente do GitLab (especialmente em compatibilidade com o `numpy>=2.0.0`).

△ Importante: Atualize o `requirements.txt` (Recomendações)

Abra o arquivo `requirements.txt` do seu projeto e adicione (ou atualize) as bibliotecas industriais usadas neste lab. Cuidado para não deixar versões antigas duplicadas:

- `numpy>=2.0.0` (Obrigatório para álgebra linear)
- `pandas>=3.0.0` (Obrigatório para manipulação de CSVs)
- `scikit-learn>=1.5.0` (Obrigatório para o `StandardScaler` e `Split`)
- `tensorflow-cpu>=2.17.0` (Obrigatório — usamos a versão CPU para evitar warnings de driver CUDA em máquinas convencionais)

Após atualizar o arquivo, não se esqueça de rodar: `pip install -r requirements.txt`.

Como o GitLab executará seu script em um container isolado sem internet, criaremos dados sintéticos em memória que simulam o problema bancário de churn. Crie o arquivo `src/churn_keras.py` e adicione a fundação:

```
1 import numpy as np
2 import pandas as pd
3 import os
4
5 os.environ['TF_CPP_MIN_LOG_LEVEL'] = '2' # Silencia logs em C++ do
    TF
6 os.environ['TF_ENABLE_ONEDNN_OPTS'] = '0' # Silencia warnings do
    oneDNN
7
8 from sklearn.model_selection import train_test_split
9 from sklearn.preprocessing import StandardScaler
10 from tensorflow.keras.models import Sequential
11 from tensorflow.keras.layers import Dense, Input
12
13 # Fixando a semente para reprodutibilidade
14 np.random.seed(42)
15
16 # Gerando 500 clientes fictícios
17 idades = np.random.randint(18, 70, size=500)
18 salarios = np.random.randint(2000, 350000, size=500)
19
20 # Regra de churn: idade > 40 E salário > 50k => cancelou (1)
21 alvo = np.where((idades > 40) & (salarios > 50000), 1, 0)
22
23 # Organizando em DataFrame pandas
24 df = pd.DataFrame({
25     'idade': idades,
26     'salario': salarios,
27     'churn': alvo
28 })
29
30 # Separando Features (X) e Target (y)
31 X = df[['idade', 'salario']].values
32 y = df['churn'].values
```

△ Importante: Por que Dados Sintéticos?

Em ambiente CI/CD, o script roda em um container sem acesso à internet. Gerar dados em memória garante reprodutibilidade e independência. Em projetos reais, você substituiria este bloco por `pd.read_csv('dados_reais.csv')`. A semente `np.random.seed(42)` garante que todos os alunos geram exatamente os mesmos dados.

B. Verificação: Distribuição dos Dados

Antes de continuar, verifique a distribuição de classes no seu dataset:

Estatística	Valor	Significado
Total de clientes	500	Gerados com <code>np.random</code>
Clientes sem churn ($y = 0$)	≈ 350	Idade ≤ 40 OU Salário $\leq 50k$
Clientes com churn ($y = 1$)	≈ 150	Idade > 40 E Salário $> 50k$
Features	2	Idade, Salário

3 Parte 3: Dividindo e Normalizando (10 min)

C. Train/Test Split

Separe os dados em treino (80%) e teste (20%). Adicione ao seu arquivo:

```
1 X_train, X_test, y_train, y_test = train_test_split(
2     X, y, test_size=0.2, random_state=42)
```

D. StandardScaler — A Normalização Obrigatória

Agora aplique a lei áurea do pré-processamento. Lembre-se do vocabulário do *scikit-learn* que vimos nos slides: `fit` é o verbo **aprender** (μ e σ), e `transform` é o verbo **aplicar**. Observe a diferença crucial:

```
1 scaler = StandardScaler()
2
3 # 1. fit_transform no TREINO: Aprende a escala e ja aplica
4 X_train_escalado = scaler.fit_transform(X_train)
5
6 # 2. transform no TESTE: Aplica a MESMA escala (sem aprender de novo)
7 X_test_escalado = scaler.transform(X_test)
```

△ Importante: Erro Grave — Data Leakage

Se você usar `fit_transform()` no teste, o scaler vai "esquecer" a escala do treino e recalcular a média usando os dados novos. Isso "vaza" informação do futuro e destrói a matemática da rede neural. Use **sempre** apenas `transform()` no conjunto de teste. Este é o erro #1 em entrevistas de Machine Learning.

E. Tabela Numérica: Antes e Depois da Normalização

Para entender o efeito do `StandardScaler`, observe os valores reais (a tabela abaixo usa dados ilustrativos):

Feature	Antes (Bruto)		Depois (Normalizado)	
	Média	Escala	Média	Escala
Idade	≈ 44	[18,70]	≈ 0.0	$[-1.8, +1.7]$
Salário	$\approx 176k$	[2k,350k]	≈ 0.0	$[-1.7, +1.7]$

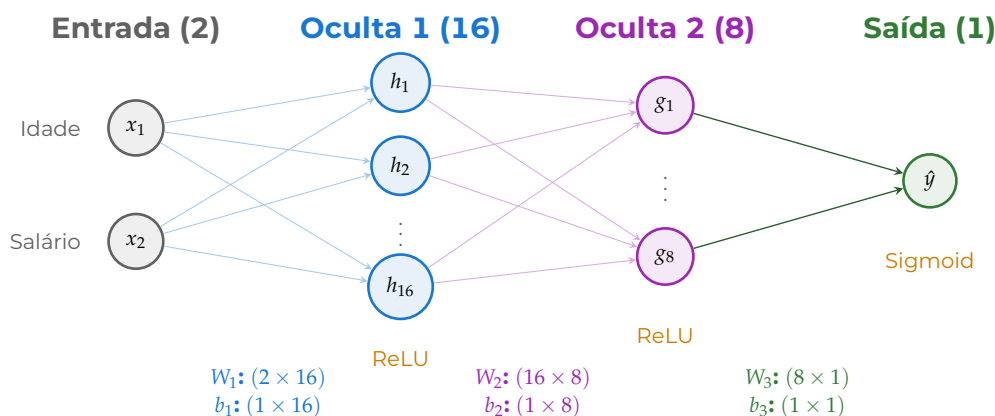
Sem normalização, o salário (escala de milhares) dominaria completamente a idade (escala de dezenas). O `StandardScaler` coloca ambas as features na mesma escala, com média ≈ 0 e desvio padrão ≈ 1 .

Adicione a verificação de sanidade:

```
1 print(f"Antes: Idade media={X_train[:,0].mean():.1f}, "
2       f"Salario media={X_train[:,1].mean():.0f}")
3 print(f"Depois: Idade media={X_train_escalonado[:,0].mean():.4f}, "
4       f"Salario media={X_train_escalonado[:,1].mean():.4f}")
5 # Depois da normalização, ambas as médias devem ser ~0.0
```

4 Parte 4: Construindo o Modelo Keras (12 min)

Agora que os tensores estão digeríveis, é hora de construir a rede. Observe o diagrama da arquitetura:



A rede possui **3 camadas**: duas ocultas com ReLU (16 e 8 neurônios) e uma de saída com Sigmoid. O total de parâmetros treináveis é $(2 \times 16 + 16) + (16 \times 8 + 8) + (8 \times 1 + 1) = 48 + 136 + 9 = 193$ — minúsculo para o Keras, mas suficiente para o nosso problema.

Adicione as funções ao seu arquivo:

```
1 def criar_modelo():
2     """Cria uma MLP com Keras Sequential."""
3     model = Sequential()
4     model.add(Input(shape=(2,)))
5     # Camada Oculta 1: 16 neurônios + ReLU
6     model.add(Dense(16, activation='relu'))
7     # Camada Oculta 2: 8 neurônios + ReLU
8     model.add(Dense(8, activation='relu'))
```

```

9     # Saída: Sigmoid para classificação binária
10    model.add(Dense(1, activation='sigmoid'))
11    return model
12
13    def compilar_e_treinar(model, X_t, y_t, epochs=50):
14        """Compila e treina o modelo."""
15        model.compile(
16            optimizer='adam',          # Gradient Descent inteligente
17            loss='binary_crossentropy', # Loss para classificação binária
18            metrics=['accuracy']       # Métrica legível
19        )
20        model.fit(X_t, y_t, epochs=epochs, verbose=0)
21        return model

```

• Teoria: Do NumPy ao Keras — A Tradução

Cada linha do Keras tem um equivalente artesanal que você já programou:

Keras (Industrial)	NumPy (Artesanal — Labs 01–03)
Dense(16, activation='relu')	CamadaDensa(2,16) + relu(Z)
loss='binary_crossentropy'	calcular_mse(y, y_pred)
optimizer='adam'	peso = peso - lr * gradiente
model.fit(X, y, epochs=80)	Loop for epoch in range(80)

Você **sabe** o que acontece por trás de cada abstração — essa é a vantagem competitiva do engenheiro que aprendeu na mão.

5 Parte 5: Avaliando e Analisando (8 min)

Adicione o bloco principal que treina e avalia o modelo:

```

1  if __name__ == "__main__":
2      # Criar e treinar
3      modelo = criar_modelo()
4      compilar_e_treinar(modelo, X_train_escalonado, y_train, epochs
5                          =80)
6
7      # Avaliar no conjunto de TESTE (dados nunca vistos)
8      perda, acuracia = modelo.evaluate(
9          X_test_escalonado, y_test, verbose=0)
10
11     print(f"\n{'='*40}")
12     print(f"    Loss no Teste:      {perda:.4f}")
13     print(f"    Acurácia no Teste: {acuracia:.2%}")
14     print(f"{'='*40}")
15
16     if acuracia > 0.65:

```

```

16     print("O modelo SUPEROU o baseline de 65%!")
17     else:
18         print("ALERTA: O modelo NÃO atingiu o baseline.")

```

Execute `python src/churn_keras.py` e verifique que a acurácia supera 65%. Se não superar, tente:

- Aumentar epochs para 100 ou 150.
- Adicionar mais neurônios na camada oculta.
- Verificar se a normalização está aplicada corretamente.

6 Parte 6: Garantia de Qualidade — Testes Automatizados (7 min)

Os testes já estão prontos no arquivo `tests/test_churn.py` do seu repositório. Abra-o e observe a estrutura:

```

1  import numpy as np
2  from src.churn_keras import (criar_modelo, compilar_e_treinar,
3                               X_train_escalonado, y_train,
4                               X_test_escalonado, y_test,
5                               scaler)
6
7  def test_modelo_supera_baseline():
8      """A acurácia no teste DEVE superar 65%."""
9      modelo = criar_modelo()
10     compilar_e_treinar(modelo, X_train_escalonado, y_train)
11     _, acc = modelo.evaluate(X_test_escalonado, y_test, verbose=0)
12     assert acc > 0.65, f"Falha! Acurácia = {acc:.2%} (< 65%)"
13
14  def test_normalizacao_aplicada():
15      """Garante que os dados foram normalizados corretamente."""
16      media = np.abs(X_train_escalonado.mean(axis=0))
17      # A média de cada coluna deve estar próxima de 0
18      assert np.all(media < 0.1), f"Dados não normalizados! Média={
19          media}"
20
21  def test_shapes_corretos():
22      """Garante que os shapes de treino/teste estão consistentes."""
23      assert X_train_escalonado.shape[1] == 2, "Features erradas"
24      assert X_test_escalonado.shape[1] == 2, "Features erradas"
25      assert len(y_train.shape) == 1, "Target deve ser 1D"

```

- 1 Rode os testes no terminal: `pytest tests/test_churn.py -v`.
- 2 Se a tela exibir **3 passed** em verde, sua implementação está correta.

7 Parte 7: Commit e Push (5 min)

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

```
1 git add src/churn_keras.py
2 git commit -m "Lab 04: Keras Sequential + StandardScaler + CI/CD"
3 git push origin main
```

O GitLab CI interceptará seu push, criará um ambiente em nuvem e executará o pytest automaticamente. O **Selo Verde** (Pipeline Passed) confirma sua pontuação.

△ Importante: Se o Pipeline Falhar

Leia o log de erro no GitLab (aba CI/CD > Pipelines > Jobs). Os erros mais comuns são: (1) esquecer o `os.environ['TF_CPP_MIN_LOG_LEVEL']` antes dos imports do TensorFlow, (2) usar `fit_transform` no conjunto de teste (Data Leakage), (3) nome de função diferente do esperado pelo teste. O CI roda os **mesmos testes** que você rodou localmente.

Conclusão: O que Você Construiu Hoje

✓ Takeaways

- Você construiu seu primeiro **pipeline completo de Machine Learning industrial**: dados → normalização → modelo → treino → avaliação.
- Usou **pandas** para estruturar dados tabulares e separar Features (X) do Target (y).
- Aplicou o **StandardScaler** corretamente (`fit_transform` no treino, `transform` no teste) e entendeu o risco do *Data Leakage*.
- Treinou uma rede **Keras Sequential** com 2 camadas ocultas ($16 \rightarrow 8 \rightarrow 1$), validando acurácia > 65%.
- Os 3 testes CI/CD garantem: baseline superado, dados normalizados e shapes consistentes.
- Cada abstração do Keras tem um equivalente artesanal que você programou nos Labs 01–03 — a diferença é **escala**, não **conceito**.

◇ Informação

Gancho para a Próxima Aula: Até agora, nosso target y é binário (0 ou 1: churn ou não-churn). Mas e se o problema tiver **3, 10 ou 100 classes**? E como saber se o modelo realmente aprendeu o padrão ou apenas **decorou** os dados de treino? Na próxima aula, entraremos na classificação **Multiclasse** com a função **Softmax** e enfrentaremos o vilão mais traiçoeiro do Machine Learning: o **Overfitting**. A pergunta será: *“meu modelo tira 99% no treino e 55% no teste — o que aconteceu?”*