

Lab. 08 (Parte 2): Colocando o Modelo em Produção com FastAPI e Docker

Aléssio Miranda Júnior

Setembro - 2026/2

Lab. 08 (Parte 2): Deploy com FastAPI e Docker

★ Objetivo da Sessão 2

No ecossistema corporativo de engenharia de software, um modelo de Machine Learning que permanece restrito a um notebook ou a um script isolado possui **impacto operacional nulo**. O valor de negócio se concretiza quando o modelo opera como um microsserviço resiliente, validado e acessível via rede por outros sistemas corporativos em poucos milissegundos.

Dinâmica Operacional da Sessão de Hoje:

- **Alunos com a Parte 1 em andamento:** Utilizem a primeira parte desta sessão para concluir o treinamento do classificador desbalanceado, gerar a matriz de confusão, serializar o modelo (.keras) e o escalonador (.joblib), assegurando a aprovação nos testes locais e no pipeline do GitLab CI (pytest tests/test_avaliacao.py).
- **Alunos com a Parte 1 concluída (ou assim que concluírem):** Iniciem imediatamente a Parte 2. Vocês transformarão os artefatos de IA em uma **API REST assíncrona** usando **FastAPI** e **Pydantic**, testarão a interface interativa com o **Swagger UI** e encapsularão o serviço em um container isolado com **Docker**.

1 Parte 1: A Arquitetura de Inferência em Produção (10 min)

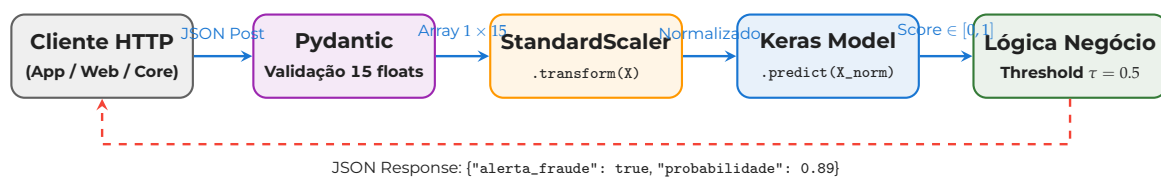
Na Sessão 1, você gerou dois artefatos fundamentais:

- 1 `modelo_fraude.keras`: Contém a topologia da rede neural e os tensores de pesos (W e b) otimizados.

- 2 `scaler_fraude.joblib`: Contém a média (μ) e o desvio padrão (σ) de cada uma das 15 características calculadas no conjunto de treino.

△ Importante: Atenção de Engenharia: A Armadilha do Escalonador Esquecido

Uma falha recorrente em pipelines amadores de deploy é expor o modelo diretamente sem o escalonador. Lembre-se: o modelo foi treinado em dados z-score ($\mu \approx 0, \sigma \approx 1$). Se o cliente HTTP enviar um payload com o valor de transação bruto (ex: R\$ 1.500,00), a rede neural receberá uma magnitude centenas de vezes superior à faixa de ativação de treino, resultando em previsões erráticas. **O escalonador faz parte indissociável da inferência.**



2 Parte 2: Construindo a API REST com FastAPI (15 min)

O framework **FastAPI** estabeleceu-se como referência para serviços de Inteligência Artificial devido à sua alta performance assíncrona (baseada em Starlette/ASGI), validação declarativa de dados com **Pydantic** e geração automática de especificações OpenAPI.

Crie o arquivo `src/api.py` no repositório do projeto:

```

1  import os
2  import joblib
3  import numpy as np
4  from fastapi import FastAPI, HTTPException
5  from pydantic import BaseModel, Field
6  from tensorflow.keras.models import load_model
7
8  # 1. Instanciacao da aplicacao com metadados
9  app = FastAPI(
10     title="NeoBank - Sentinela de Fraudes API",
11     description="Microsservico de inferencia em tempo real para
12         deteccao de transacoes fraudulentas.",
13     version="1.0.0"
14 )
15
16 # 2. Variaveis globais de persistencia
17 MODEL_PATH = "modelo_fraude.keras"
18 SCALER_PATH = "scaler_fraude.joblib"
  
```

```

19 modelo = None
20 scaler = None
21
22 # Carregamento sob demanda ou inicializacao
23 @app.on_event("startup")
24 def carregar_artefatos():
25     global modelo, scaler
26     if not os.path.exists(MODEL_PATH) or not os.path.exists(
27         SCALER_PATH):
28         raise RuntimeError(
29             f"Artefatos nao encontrados! Execute o treinamento
30             primeiro. "
31             f"Arquivos necessarios: {MODEL_PATH} e {SCALER_PATH}"
32         )
33     modelo = load_model(MODEL_PATH)
34     scaler = joblib.load(SCALER_PATH)
35     print("Artefatos de inferencia (modelo e scaler) carregados com
36         sucesso.")
37
38 # 3. Contratos de Dados (Schemas Pydantic)
39 class TransacaoInput(BaseModel):
40     features: list[float] = Field(
41         ...,
42         min_items=15,
43         max_items=15,
44         description="Vetor com exatamente 15 caracteristicas
45             numericas da transacao.",
46         example=[1250.0, 2.5, 0.8, -1.2, 0.05, 3.1, -0.4, 1.8,
47                 0.9, -0.2, 0.1, 0.4, -0.05, 0.12, -0.3]
48     )
49
50 class PredicaoOutput(BaseModel):
51     probabilidade: float
52     classe: int
53     alerta_fraude: bool
54     status: str
55
56 # 4. Endpoints da API
57 @app.get("/health", tags=["Monitoramento"])
58 def health_check():
59     """Endpoint de telemetria e sondagem de saude (Liveness/Readiness)"""
60     return {
61         "status": "healthy",
62         "service": "neobank-sentinela",
63         "modelo_carregado": modelo is not None,
64         "scaler_carregado": scaler is not None
65     }
66
67 @app.post("/predict", response_model=PredicaoOutput, tags=["Inferencia"])
68 def predict(transacao: TransacaoInput):

```

3 PARTE 3: EXECUÇÃO LOCAL E VALIDAÇÃO COM SWAGGER (10 MIN) 4

```

65     """Executa a predicao de fraude para uma transacao recebida."""
66     try:
67         # Converter para array bidimensional (1 amostra, 15 colunas)
68         X_raw = np.array(transacao.features).reshape(1, -1)
69
70         # Aplicar a normalizacao com os parametros de treino
71         X_norm = scaler.transform(X_raw)
72
73         # Executar inferencia na rede neural
74         score = float(modelo.predict(X_norm, verbose=0)[0][0])
75         classe = 1 if score > 0.5 else 0
76
77         return PredicaoOutput(
78             probabilidade=round(score, 4),
79             classe=classe,
80             alerta_fraude=bool(classe == 1),
81             status="BLOQUEADA" if classe == 1 else "APROVADA"
82         )
83     except Exception as e:
84         raise HTTPException(status_code=500, detail=f"Falha de
            inferencia: {str(e)}")

```

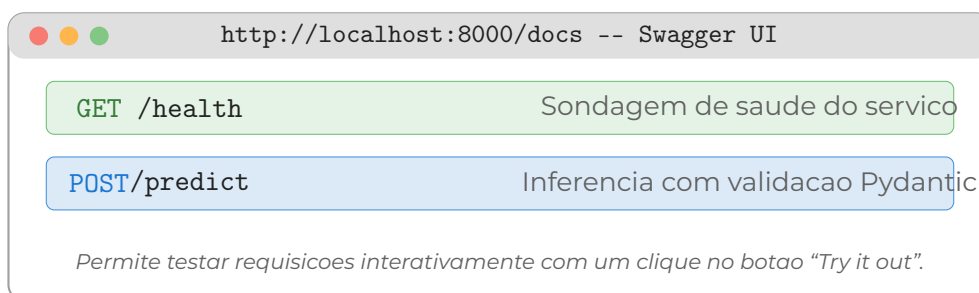
3 Parte 3: Execução Local e Validação com Swagger (10 min)

Para executar a API em ambiente de desenvolvimento, utilize o servidor ASGI de alta performance **Uvicorn**:

```
1 uvicorn src.api:app --host 0.0.0.0 --port 8000 --reload
```

Abra o navegador e acesse a URL:

<http://localhost:8000/docs>



Exercício de Verificação de Integridade:

- 1 Clique em **POST /predict** → **Try it out**.
- 2 Envie o payload padrão com 15 números e clique em **Execute**. Observe a resposta HTTP 200 com a chave "probabilidade".

- 3 Agora teste a robustez do Pydantic: **remova um número** (enviando 14 características). Clique em **Execute**.
- 4 Verifique que a API respondeu imediatamente com **HTTP 422 (Unprocessable Entity)**, informando que o vetor é inválido, sem lançar exceções não tratadas e sem derrubar o servidor.

4 Parte 4: Containerização com Docker (10 min)

Na engenharia de produção, a aplicação deve ser independente do sistema operacional hospedeiro. Para garantir isolamento de dependências e portabilidade completa, criaremos o arquivo Dockerfile na raiz do projeto.

```

1 # 1. Imagem base oficial enxuta com Python 3.12
2 FROM python:3.12-slim
3
4 # 2. Definir variaveis de ambiente para otimizar execucao Python
5 ENV PYTHONDONTWRITEBYTECODE=1
6 ENV PYTHONUNBUFFERED=1
7
8 # 3. Diretorio de trabalho da aplicacao
9 WORKDIR /app
10
11 # 4. Copiar arquivo de dependencias e instalar
12 COPY requirements.txt .
13 RUN pip install --no-cache-dir -r requirements.txt
14
15 # 5. Copiar codigo-fonte e os artefatos de modelo treinados
16 COPY src/ ./src/
17 COPY modelo_fraude.keras .
18 COPY scaler_fraude.joblib .
19
20 # 6. Declarar porta de rede exposta
21 EXPOSE 8000
22
23 # 7. Comando de inicializacao do servidor de producao
24 CMD ["uvicorn", "src.api:app", "--host", "0.0.0.0", "--port", "8000"]

```

Crie também o arquivo `.dockerignore` para evitar cópia de arquivos temporários desnecessários para a imagem:

```

1 __pycache__/
2 *.pyc
3 .git
4 .venv
5 tests/
6 .pytest_cache
7 *.png

```

Certifique-se de que o arquivo `requirements.txt` contenha as dependências mínimas necessárias:

```
1 fastapi>=0.110.0
2 uvicorn>=0.28.0
3 pydantic>=2.6.0
4 tensorflow-cpu>=2.16.0
5 scikit-learn>=1.4.0
6 joblib>=1.3.0
7 numpy>=1.26.0
```

5 Parte 5: Construção e Execução do Container (10 min)

No terminal do seu ambiente local, execute a montagem da imagem Docker:

```
1 # Construir a imagem com a tag neobank-fraude:v1
2 docker build -t neobank-fraude:v1 .
```

Após a compilação das camadas, inicialize o container mapeando a porta 8000 do host para a porta 8000 interna do container:

```
1 # Inicializar o container em background (daemon)
2 docker run -d -p 8000:8000 --name sentinela-container neobank-fraude:
  v1
```

Verifique o status do container em execução e inspecione os logs:

```
1 docker ps
2 docker logs sentinela-container
```

Se a saída exibir a mensagem confirmando a inicialização do Uvicorn na porta 8000, sua IA está oficialmente operando em um ambiente isolado de produção.

6 Parte 6: Auditoria de Inferência em Produção (5 min)

Para validar que o serviço responde a requisições externas como esperado em um sistema distribuído, execute uma requisição utilizando o utilitário curl diretamente no terminal:

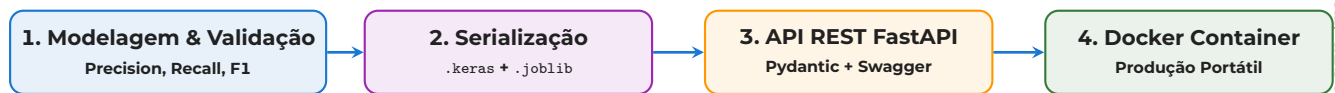
```
1 curl -X POST "http://localhost:8000/predict" \
2   -H "Content-Type: application/json" \
3   -d '{
4     "features": [
5       1250.0, 2.5, 0.8, -1.2, 0.05, 3.1, -0.4, 1.8,
6       0.9, -0.2, 0.1, 0.4, -0.05, 0.12, -0.3
7     ]
8   }'
```

A resposta retornará no formato JSON estruturado:

```
1 {  
2   "probabilidade": 0.0412,  
3   "classe": 0,  
4   "alerta_fraude": false,  
5   "status": "APROVADA"  
6 }
```

Conexão Estratégica: A Ponte para o TP1

Parabéns! Você acaba de construir o fluxo corporativo completo que conecta a Ciência de Dados à Engenharia de Software.



◇ Informação

O que você fará no Trabalho Prático 1 (TP1): No TP1, você e sua equipe desenvolverão uma aplicação completa de Inteligência Artificial — seja implementando a especificação oficial definida pelo professor (como o Classificador de Habilidades do ecossistema WOKDEX) ou selecionando um desafio real em plataformas abertas como o **Kaggle**, conforme a orientação da disciplina. O ciclo que você executou hoje (Labs 08 Parte 1 e Parte 2) serve como template arquitetural obrigatório:

- Coleta, limpeza e exploração do dataset real.
- Treinamento de uma arquitetura profunda com prevenção de sobreajuste (Dropout/EarlyStopping).
- Avaliação crítica através de métricas adequadas ao desbalanceamento do problema.
- Exposição da inferência através de uma API FastAPI documentada e empacotada em container Docker.

✓ Takeaways

- **Artefatos Complementares:** O modelo (.keras) armazena pesos e topologia, enquanto o escalonador (.joblib) garante que os dados em tempo de inferência recebam a mesma transformação estocástica do treino.
- **Contrato Estrito de Dados:** O Pydantic valida os tipos e a cardinalidade das entradas, rejeitando payloads malformados na fronteira da rede e protegendo o núcleo de IA contra erros de execução.
- **Independência de Ambiente:** O Docker elimina inconsistências entre sistemas operacionais, encapsulando código, runtime Python e dependências binárias em uma imagem imutável.
- **Fechamento do Módulo 1:** Você dominou a trilha que vai do cálculo manual do Perceptron solitário até o deploy containerizado de microsserviços de IA.