



INSIGHTS · INTELIGÊNCIA ARTIFICIAL

**Um modelo que não explica
decisões é um relatório
que ninguém confia.**

Desliza →

Por que explicabilidade importa

Scores são operadores de decisão: sem explicações surgem hesitação e intervenções manuais. Em casos reais, explicações reduziram o tempo de investigação em ~60%.

Requisitos mensuráveis

Defina latência, granularidade e auditabilidade antes de escolher técnica ou arquitetura. Ex.: p95 <150 ms ou explicações assíncronas em minutos.

Arquitetura prática

Treino no Lakehouse/Notebooks; artefactos versionados (ONNX, MLflow) na Lakehouse; serving online ou batch para integração com Power BI.

Abordagem híbrida

Scoring online leve e explicações full SHAP apenas para top-N (ex.: 5%) ou para casos críticos. Redução de custos de compute ~65% em testes.

Técnicas aplicáveis

SHAP para árvores; DeepSHAP/Integrated Gradients para redes; regras legíveis por humanos em paralelo para mensagens e prioridade.

Monitorização accionável

Monitore estabilidade dos SHAP, PSI/KS das features e alerts. Alertas anteciparam 3 regressões e evitaram perdas estimadas em ~€120k.

Mini-caso: e-commerce

Empresa com 80 pessoas e 200k compras/mês;
latência alvo 200 ms; solução híbrida com
latência média 60 ms e explicações imediatas
para 20% dos casos.

CONTINUA NO BLOG

Se a sua equipa não percebe por que um cliente recebeu 0.8 em vez de 0.4, como espera que confiem nas recomendações do modelo?

www.bconcepts.pt/insights