



Microsoft-approved Fabric Workload · also available as a cloud service

## Take Copilot and Data Agents to Production

BI Pixie assesses the AI Readiness of your semantic models, optimizes them, and benchmarks Copilot and data agents to prove they stay correct.

### Resolve Ambiguity Before It Derails AI Attention

**Your Power BI semantic models already carry the business context AI needs, but it was written for people, not for agents.** It can take time and many iterations until your data analysts can scope down context and resolve the ambiguity, so AI attention stays on what answers the question. How do you scale this process? When do you know your data is ready for AI?

### Three Problems Between You and Production

#### How do you get AI Readiness at scale?

Optimizing one semantic model takes many iterations and tests. You need a way to scale that across the mission-critical ones.

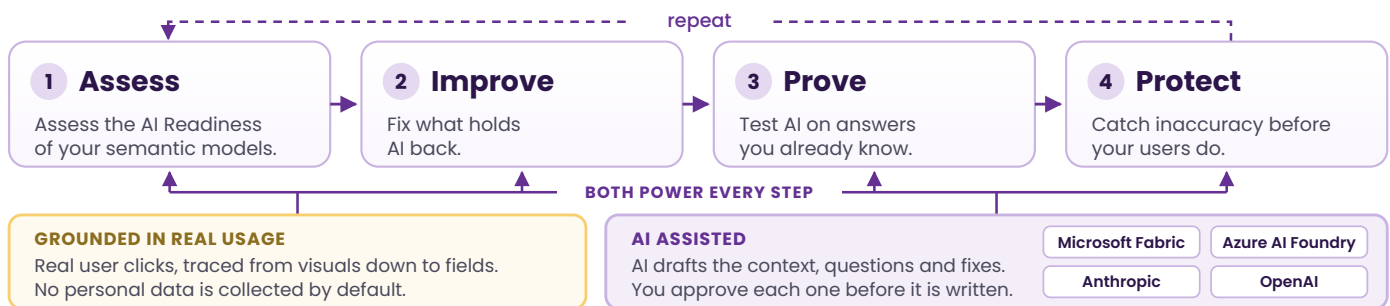
#### A wrong AI answer looks like a right one

An agent that lacks context does not report an error. It returns a confident answer that may be wrong, and you find out in an executive review.

#### Stay correct as semantic models change

A new measure or a deleted description can break an agent that answered correctly last month, and nothing flags it until a user reports it.

### The BI Pixie AI Readiness Loop



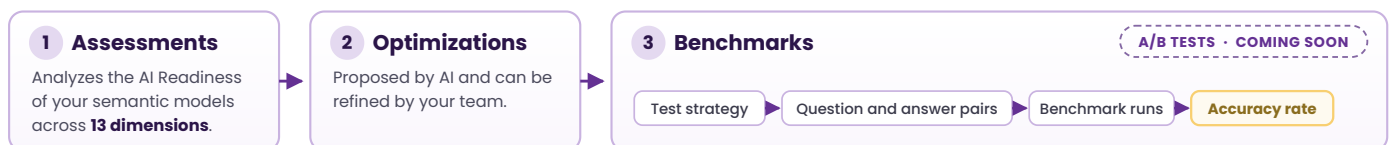
### Prove the Answers Are Right Before Rollout

- ◆ **You choose which AI to test.** It can be your own Fabric data agent, the query engine behind Copilot in Power BI, or a generic AI agent when you need every wrong answer traceable.
- ◆ **Every answer is checked against the expected one.** BI Pixie gets the expected answer by running its own DAX query against your semantic model, then compares what the agent said with it.
- ◆ **Built from what your people actually ask.** BI Pixie proposes the questions from your business domains and the report visuals your people use, mixed by complexity, and you edit any before the run.

### Catch Inaccuracy Before Executives Lose Trust

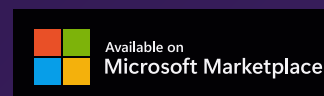
- ◆ **BI Pixie detects a regression in AI accuracy before it reaches your users.** You run the same benchmark after every change to your semantic models, and BI Pixie compares it with the run before.
- ◆ **BI Pixie keeps a repeated benchmark resilient to changes in your data.** The questions cover a time period that has already ended, so the expected answers are less likely to move, and you can have them recalculated from your current data before any run.
- ◆ **Drift in the context is caught early.** On the Enterprise plan, BI Pixie runs AI Readiness assessments on a schedule you set.

### Key Features



## Prove Your AI Is Ready for the Business

- Free assessment of up to 500 semantic models
- Fixed plans, no usage billing
- Available on Microsoft Marketplace



app.bipixie.com