



AI in Local Government: A Case Study on how Pleasant Ridge Adopted AI

▶ CALIFORNIA LEAGUE OF CITIES | APRIL 2025

JASON PETERS, VC3

JIM BREUCKMAN, CITY OF PLEASANT RIDGE

Session Agenda

- Icebreaker Activity — 15 min (complete)
- Meet Pleasant Ridge — 5 min
- The Problem & Chatbot Solution + Live Demo — 15 min
- Build & Implementation Challenges — 10 min
- Other AI Use Cases (Jim) — 5 min
- Live Demos: Synthreo Suite — 20 min
- How to Implement — 10 min



Meet Pleasant Ridge



▶ WHO THEY ARE

Pleasant Ridge: Who They Are



- ▶ Population: 2,500
- ▶ 3 City Hall Staff
- ▶ Small team. Big vision.
- ▶ Among the first cities in Michigan to adopt an AI chatbot for resident services.
- ▶ Visionary leadership proving that city size is no barrier to innovation.

The Problem & The Solution



▶ JIM'S STORY

The Problem: Why Something Had to Change

JIM'S STORY

- ▶ 3-person city hall team fielding all City
- ▶ No after - hours coverage. Calls going to voicemail. Residents frustrated.
- ▶ Staff time consumed answering: “What time does city hall close?” “When is trash pickup?” “Can I pay my water bill online?”
- ▶ Every repetitive call was time NOT spent on complex issues that actually needed a human.



The Solution: AI Chatbot for City Websites



- ▶ Resident asks a question — 24/7, no hold music.
- ▶ Chatbot answers instantly or routes to the right staff.
- ▶ Powered by n8n, Google Gemini, OpenAI and Llama
- ▶ Deployed on Pleasant Ridge's city website.

LIVE DEMO



- ▶ **Let's see it in action.**
- ▶ 2–3 realistic resident questions · One suggestion from the audience

Build & Implementation Challenges



▶ WHAT WE LEARNED THE HARD WAY



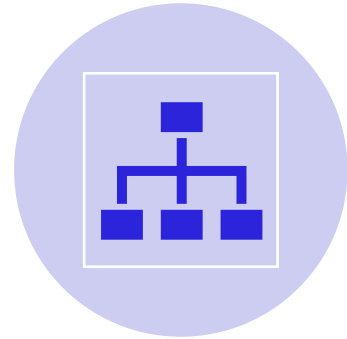
- ▶ **The First Version Had Accuracy Problems**
 - ▶ In testing, the chatbot delivered incorrect or incomplete answers.
 - ▶ Root cause: the AI had no way to identify WHAT TYPE of question it was answering before attempting a response.
 - ▶ The AI was operating based solely on probability matches and sometimes looked at the wrong data sources as a result
 - ▶ **That extra 5% makes all the difference**
 - ▶ Jim's biggest takeaway

How We Fixed It: Categorization First



- ▶ **Built a Categorization Layer**
 - ▶ AI now identifies the question type BEFORE attempting an answer.
 - ▶ Is it a meeting minutes question, website question, recreation programs, ordinances?
- ▶ **Mapped Every FAQ Category in Advance**
 - ▶ Defined every edge case and escalation path before deployment.
 - ▶ Result: Drove accuracy up to 99%

The AI Agent Lifecycle



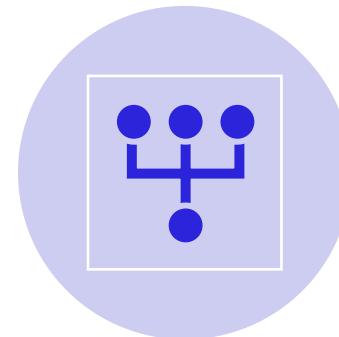
Intake: User inputs their question or request



Planning: Agent reviews request and determines where to pull the info and how to respond



Execution: AI executes the plan and passes the results to the validation stage



Revise: AI reviews the results, identifies issues and revises approach and/or response



Respond: Resident gets an accurate answer or is routed to the right staff.



Repeat: Every interaction makes the system smarter for the next one.

Other AIUse Cases: What Jim Built Next



▶ FROM STUDENT TO BUILDER

From Guidance to Action: Jim's AI Journey



- ▶ Automation with Claude Cowork
- ▶ Implemented AI Solution to automate payroll processes
- ▶ Writing optimization in his voice/tone
- ▶ Budget/bond planning
- ▶ Jim's AI Journey, where he started and how he grew his AI skills
- ▶ Jim's AI Goals/Projects



▶ Live Demos: AISystems

How to Implement AI in Your City



▶ YOUR NEXT FIVE STEPS

- ▶ **Risk:** AI integration exacerbates current security vulnerabilities in your file sharing permissions
 - ▶ **Example :** If sensitive data is saved in an obscure location where nobody will look but it's technically accessible, people are much more likely to find it with MS CoPilot.
- ▶ **Risk:** Improperly configured M365 environments may fail to protect your data
 - ▶ **Example :** CoPilot has an AI safety system that is enabled by default, which will flag a chat for review by a Microsoft employee if it's flagged as potentially malicious, harmful or otherwise violating policy. This could potentially expose confidential or otherwise protected data (such as PII or CJIS protected data) to an unauthorized Microsoft reviewer.
- ▶ **Mitigation:** An AI readiness assessment scan can identify these issues prior to rolling out integrated AI , protecting your organization from exposure of protected information.
- ▶ **Mitigation:** Remediate findings with a qualified IT provider with specialization in AI

- ▶ Establish AI Governance
- ▶ AI Governance Council
- ▶ AI Use Policy
- ▶ AI Governance Handbook
- ▶ AI Inventory
- ▶ Define risk tolerance
- ▶ AI Incident Response Plan

- ▶ Foundational AI Training
- ▶ Microsoft CoPilot
- ▶ Secure LLM Toolset (i.e. Synthreo)
- ▶ Tool specific use guides
- ▶ Tool specific training
- ▶ Approved use cases
- ▶ Prohibited use cases
- ▶ Performance Metrics and Monitoring

- ▶ Educate your workforce on how to use AI safely and effectively through foundational AI training (prompt engineering, identifying AI use cases, mitigating hallucination risks, etc.)
- ▶ Align training to your AI Use Policy and Governance handbook
- ▶ Align organization on AI data governance and ensure everyone understands how to protect data from disclosure through public AI models.
- ▶ Align organization on guiding principles of AI use
- ▶ Ensure everyone understands prohibited uses (and the why)
- ▶ Ensure everyone understands incident response procedures, especially their role in reporting potential issues
- ▶ Progress to role and tool specific AI training after you have completed foundational training

AI Governance



▶ HOW TO USE AI SAFELY

What are the risks?



- ▶ New York City's AI Chatbot told businesses to break the law, telling landlords they can reject tenants with housing vouchers
- ▶ A Fresno School Official Resigned after using quotes fabricated by ChatGPT, she moved too fast under pressure and didn't fact check - check the AI's output. The school had no AI use policy or guidance.
- ▶ Detroit police wrongfully arrested a man due to incorrect facial recognition, leading them to change their policy for use of facial recognition
- ▶ Seattle Police Reports drafted by AI (no AI Use policy) was found to undermine their credibility in court, in one test, an AI-drafted police report include an officer who wasn't there.

- ▶ **AI Use Policy Defined:** An AI Use Policy is a set of guidelines and regulations established by an organization to govern the development, deployment, and use of artificial intelligence systems.

- ▶ **Benefits include:**
 - ▶ **Risk Mitigation:** Helps identify and address potential risks associated with AI use, such as bias, privacy breaches, or unintended consequences.
 - ▶ **Ethical Alignment:** Ensures AI systems are developed and used in accordance with the municipality's values and ethical standards.
 - ▶ **Regulatory Compliance:** Aids in meeting legal and regulatory requirements related to AI, data protection, and privacy.
 - ▶ **Consistency and Standardization:** Provides a framework for consistent decision-making and implementation across the organization.
 - ▶ **Transparency and Trust:** Demonstrates commitment to responsible AI use, building trust with stakeholders and customers.
 - ▶ **Innovation and Exploration Guidance:** Encourages innovation within defined boundaries, balancing progress with responsible practices.

- ▶ Purpose & Scope
- ▶ Guiding Principles
- ▶ AI Governance & Roles
- ▶ AI Risk Management Process
- ▶ Acceptable Use
- ▶ Prohibited Uses
- ▶ Data Governance
- ▶ Procurement and Third -Party Management
- ▶ Enforcement, Monitoring and Review

- ▶ AI Governance Council
- ▶ Incident Response Procedure
- ▶ AI Training Roadmap
- ▶ Implementation Guide
- ▶ AI Governance Handbook
 - ▶ Define risk tolerance
 - ▶ AI Tool Inventory and Risk Assessment

▶ OECD AI Policy Principles

- ▶ **What are they?:** The OECD AI Principles promote use of AI that is innovative and trustworthy and that respects human rights and democratic values. Adopted in May 2019, they set standards for AI that are practical and flexible enough to stand the test of time.
- ▶ **What do they do?:** The Principles guide AI actors in their efforts to develop trustworthy AI and provide policymakers with recommendations for effective AI policies.
- ▶ **Who uses them?:** The EU, the Council of Europe, the US, the UN and other jurisdictions across the world.

▶ NIST AI Risk Management Framework

- ▶ **NIST:** The National Institute of Standards and Technology (NIST) is a federal agency under the U.S. Department of Commerce that develops standards, guidelines, and best practices across various fields.
- ▶ **NIST AI Risk Management Framework (RMF):** The NIST AI Risk Management Framework (AI RMF 1.0) is a voluntary guidance document that provides a structured approach for identifying, assessing, and managing AI-related risks.
- ▶ **NIST AI Risk Management Framework (RMF) Playbook:** The NIST AI Risk Management Framework (RMF) Playbook is a companion implementation guide that translates the AI RMF's high-level principles into actionable steps.

▶ San Jose Gov AI Policy Coalition

- ▶ Coalition of local governments focused on the responsible use of AI, providing AI policy guidance, templates and collaboration opportunities

Thank You.



JASON PETERS, VC3

JIM BREUCKMAN, CITY OF PLEASANT RIDGE



JASON.PETERS@VC3.COM

CITYMANAGER@CITYOFPLEASANTRIDGE.ORG