

RYE AI Core

Open Source & Local Deployment Guide for Apple iMac (macOS)

Architecture: Autonomous Local Reasoning & Orchestration **Inference Engine:** Ollama / Local Open Weights

Cost: \$0 / Open Source (Air-Gapped)

This deployment manual details the step-by-step setup for running the **RYE AI Core** app logic on an Apple iMac. Designed for local intelligence orchestration, agent execution, and decentralized workflow management, this stack runs 100% offline without external cloud API keys, recurring subscriptions, or SaaS paywalls.

1. System & Hardware Prerequisites

Hardware & OS Requirements

- **Hardware:** Apple Silicon (M-Series recommended for fast local tensor processing) or Intel iMac
- **OS:** macOS Ventura 13.0 or higher
- **Storage:** Minimum 12 GB available SSD space for local agent memory stores and open model layers

Core Tooling Installation

Open Terminal on your iMac and execute the following command to install Homebrew, Python, Node.js, and Ollama:

```
# Install Homebrew if not already installed
/bin/bash -c "$(curl -fsSL https://raw.githubusercontent.com/Homebrew/install/HEAD/install.sh)"

# Install core runtime tools, git, and local LLM runtime
brew install python@3.11 node git ollama
```

2. Local Inference & Reasoning Engine Setup (Zero API Cost)

The RYE AI Core delegates natural language understanding, function calling, and workflow execution to local open-weights models served directly on your iMac's hardware.

```
# Start the local Ollama background engine
brew services start ollama

# Pull local reasoning model layer (e.g., Gemma 2 2B or 9B)
ollama pull gemma2:2b
```

100% Offline & Private AI Orchestration

By routing RYE AI Core requests locally to `http://localhost:11434`, all prompt workflows, agent states, and local execution pipelines remain air-gapped on your physical iMac.

3. Local Directory & Environment Configuration

Workspace Setup

```
# Create project folder
mkdir -p ~/rye-ai-core
cd ~/rye-ai-core

# Initialize isolated Python virtual environment
python3 -m venv venv
source venv/bin/activate
```

Open Source Environment Configuration

Create a local `.env` file configured strictly for local offline execution:

```
cat << 'EOF' > .env
# RYE AI Core Open Source Configuration
RYE_NODE_ID="imac-rye-core-01"
RYE_PORT=8100
OFFLINE_MODE=true

# Local Open-Source AI Endpoint (No API Keys Required)
LOCAL_LLM_PROVIDER="ollama"
LOCAL_LLM_ENDPOINT="http://localhost:11434"
LOCAL_LLM_MODEL="gemma2:2b"
EOF
```

4. Application Setup & Python Packages

```
# Upgrade pip package installer
pip install --upgrade pip

# Install local AI orchestration, async web, and vector memory dependencies
pip install ollama python-dotenv fastapi uvicorn streamlit requests chromadb
```

5. Executing RYE AI Core

Launch the RYE AI Core daemon or interactive dashboard on your iMac:

```
# Ensure virtual environment is active
source ~/rye-ai-core/venv/bin/activate

# Launch primary Python core engine
python3 main.py

# Alternatively, run Streamlit user interface
streamlit run app.py --server.port 8100
```

Access the local RYE AI interface in Safari or Chrome at: `http://localhost:8100`

6. macOS Auto-Start Daemon Setup (`launchd`)

Configure your iMac to automatically boot the RYE AI Core background engine on system startup:

```
cat << 'EOF' > ~/Library/LaunchAgents/com.rye.aicore.plist
<?xml version="1.0" encoding="UTF-8"?>
<!DOCTYPE PLIST PUBLIC "-//Apple//DTD PLIST 1.0//EN" "http://www.apple.com/DTDs/PropertyList-1.0.plist">
<plist version="1.0">
<dict>
  <key>Label</key>
  <string>com.rye.aicore</string>
  <key>ProgramArguments</key>
  <array>
    <string>/Users/SHARED_USER/rye-ai-core/venv/bin/python3</string>
    <string>/Users/SHARED_USER/rye-ai-core/main.py</string>
  </array>
  <key>RunAtLoad</key>
  <true/>
  <key>KeepAlive</key>
  <true/>
  <key>WorkingDirectory</key>
  <string>/Users/SHARED_USER/rye-ai-core</string>
</dict>
</plist>
EOF
```

Note: Replace `/Users/SHARED_USER/` with your actual home directory path (found via `echo $HOME`).

```
# Load daemon into macOS launch system
launchctl load ~/Library/LaunchAgents/com.rye.aicore.plist
```

7. System Health Verification

```
# Check local Ollama model execution status
ollama list

# Test local RYE model generation endpoint directly
curl http://localhost:11434/api/generate -d '{
  "model": "gemma2:2b",
  "prompt": "RYE AI Core health check",
  "stream": false
}'
```