


```
LOCAL_MODELS=llama3.3

# 3. (Optional) run with NO Anthropic key — make
DEFAULT_MODEL=llama3.3
DEEP_REASONING_MODEL=llama3.3
ROUTING_MODEL=llama3.3
# ...and leave ANTHROPIC_API_KEY unset
```

The listed slugs appear in the **Council UI** model dropdown, so you can also run a hybrid setup — keep the Executive on Claude while flipping individual specialists to a local model per-agent.

Caveats. Server-side web search (`ENABLE_WEB_SEARCH`) and Anthropic prompt caching / extended thinking have no local equivalent and are automatically disabled for local models. Multi-agent routing leans heavily on tool use, so pick a model that's strong at it (e.g. Llama 3.3 70B, Qwen2.5) — small models may route poorly. `LOCAL_API_KEY` is only needed if your server (vLLM, or a gateway) requires a bearer token; Ollama and LM Studio need none.

Adding a New Specialist Agent

- Create `packages/core/openexecutive/agents/your_agent.py` extending `BaseAgent`
- Add a system prompt constant in `packages/core/openexecutive/prompts/domain_prompts.py`
- Register in `packages/core/openexecutive/orchestrator/router.py` — add to `SPECIALIST_REGISTRY` and the `specialist` enum in `SPECIALIST_TOOLS`
- Add domain alias to `DOMAIN_ALIASES` in `packages/core/openexecutive/knowledge/retriever.py`
- Add knowledge docs to `knowledge/builtin/your_domain/`
- Add at least 2 eval scenarios to `evals/scenarios/`
- Submit a PR — CI requires all of the above

Development

```
make dev          # Start FastAPI + Next.js
make test         # Run Python tests
make eval         # Run eval suite
make lint         # Run ruff + mypy
make docker       # Build and run Docker stack

# Unit tests only (no API calls required)
pytest packages/core/tests/unit/ -v
```

Evaluation System

`evals/` contains 29 scenarios covering all 8 domains, scored by `claude-opus-4-7` as an LLM-as-judge. Each scenario defines a query, simulated company context, expected topics, required specialist routing, and a domain-specific rubric. Five scoring dimensions (persona coherence, domain accuracy, company context utilization, routing quality, actionability) are each rated 1–5. The CI gate requires $\geq 3.5/5$ average; any dimension dropping $> 10\%$ vs `main` fails the PR.

Privacy

Everything in `company/` is gitignored — the profile YAML, uploaded documents, and the ChromaDB vector store. None of this leaves your local machine (or your own Fly volume in cloud deployments) except as part of prompts sent to the Anthropic API. Anthropic does not train on API data.

Contributing

See [github/CONTRIBUTING.md](#). All PRs must include:

- Working implementation (no stubs)
- Tests for new behavior
- Eval scenarios for new agents or prompt changes

License

Apache 2.0 — free to use commercially, requires attribution.

