

Guide Pratique

OpenClaw + Bases de donnees PostgreSQL, RAG et recherche semantique

Integration base de donnees pour agents IA en production

Mars 2026 — Serie OpenClaw B10/10

Sommaire

1. Au-dela des fichiers : pourquoi connecter a une vraie base de donnees
2. PostgreSQL + OpenClaw : connexion directe via scripts
3. Cas d'usage concrets (timesheet, CRM, logs structures)
4. RAG (Retrieval-Augmented Generation) : pourquoi ca change tout
5. Implementer un RAG basique avec OpenClaw
6. Recherche semantique sur vos donnees
7. Limites et realite terrain
8. Architecture recommandee : quand utiliser quoi
9. Checklist integration DB (10 regles)
10. Conclusion de serie + index B1-B10

Les fichiers plats ont leur utilité. Mais quand un réseau d'agents doit lire, écrire et raisonner sur des données structurées à l'échelle — les bases de données changent la donne. Ce guide clos la série technique OpenClaw sur une note concrète : PostgreSQL, RAG et recherche sémantique.

1. Au-delà des fichiers : pourquoi connecter à une vraie BDD

Le workspace Git couvre 80% des scénarios d'automatisation. Mais certaines situations nécessitent plus :

- Volumes importants — 50 000 transactions ne se lisent pas en Markdown
- Requetes dynamiques — filtrer par date, statut, revenu : SQL > grep
- Concurrence — plusieurs agents écrivant dans le même fichier = race conditions
- Recherche sémantique — par sens, pas par mot-clé exact

■ *Gartner 2025 : 61% des déploiements agents IA en production ont besoin d'une couche persistance structurée dans les 6 premiers mois.*

2. PostgreSQL + OpenClaw : connexion directe

L'intégration la plus directe : un script Python dans le workspace. L'agent l'exécute via son skill exec, le script se connecte via psycopg2, retourne les résultats dans un fichier JSON.

```
conn = psycopg2.connect(  
    host=os.environ['PG_HOST'],  
    dbname=os.environ['PG_DB'],  
    user=os.environ['PG_USER'],  
    password=os.environ['PG_PASS']  
)  
  
cur.execute("SELECT client, SUM(hours) FROM timesheets  
WHERE month=%s GROUP BY client", ('2026-03',))
```

Règle absolue : les credentials ne transitent jamais dans le workspace Git. Variables d'environnement injectées par le vault uniquement.

3. Cas d'usage concrets

Timesheet automatique

L'agent LEDGER exécute chaque vendredi un script qui agrège les timesheets depuis PostgreSQL et génère un rapport Markdown. Zéro intervention humaine.

CRM mis à jour par un agent

Après chaque appel client logué en messagerie, l'agent NEXUS extrait les infos clés et met à jour la ligne correspondante dans la table contacts. Le CRM se tient à jour en temps réel.

Logs structures d'operations

Chaque action significative des agents est loggée dans agent_logs (timestamp, agent source, type, statut).
Audit complet — impossible à faire de façon fiable avec des fichiers texte.

4. RAG : pourquoi ça change tout

Un LLM a une fenêtre de contexte limitée. Impossible d'injecter 10 000 pages. Le RAG résout ça : plutôt que tout injecter, on récupère seulement les passages les plus pertinents à la question posée.

Architecture simple :

1. Indexation (une fois)

Documents -> Embeddings (vecteurs 1536 dims) -> pgvector

2. Requête (chaque question)

Question -> Embedding -> Similarité -> Top-K passages -> LLM

"La pertinence est calculée par similarité sémantique — proximité de sens dans l'espace vectoriel, pas correspondance exacte de mots-clés."

5. Implémenter un RAG basique

Générer des embeddings

```
def get_embedding(text: str) -> list[float]:
    response = client.embeddings.create(
        model='text-embedding-3-small',
        input=text
    )
    return response.data[0].embedding
```

pgvector (production)

```
CREATE EXTENSION IF NOT EXISTS vector;

CREATE TABLE documents (
    id SERIAL PRIMARY KEY,
    content TEXT,
    embedding vector(1536)
);

CREATE INDEX ON documents USING hnsw (embedding vector_cosine_ops);
```

Requête de similarité

```
SELECT content, 1-(embedding <=> %s::vector) AS score
FROM documents ORDER BY embedding <=> %s::vector LIMIT 5;
```

■ Pour les volumes < 10 000 docs : JSON + cosine en Python suffisent. Pour la production : pgvector.

6. Recherche sémantique sur vos données

Cas d'usage BOTUM :

- Retrouver des décisions passées : 'quelle était la décision sur le backup ?' -> passage exact dans les notes de réunion
- Trouver les billets similaires avant publication

- Chercher dans un CRM : 'quels clients ont mentionne des problemes de securite ?'
- Knowledge base interne : les agents trouvent la procedure standard applicable

7. Limites et realite terrain

Qualite des embeddings

Un corpus de notes cryptiques produit des embeddings inutiles. La qualite du corpus determine la qualite des resultats.

Cout

text-embedding-3-small : \$0.02/million tokens. 10 000 docs x 500 tokens = ~\$0.10 pour l'indexation initiale.

Latence

200-500ms par requete (API embedding + recherche vectorielle). Negligeable pour un agent interactif.

Quand le RAG est inutile

Donnees < 50 000 tokens : injecter directement dans le contexte. Claude 3.5 gere 200K tokens.

8. Architecture recommandee

Situation	Solution	Pourquoi
Donnees < 50K tokens	Injection directe	Zero latence
Donnees structurees	PostgreSQL psycopg2	SQL > grep
< 10K docs, proto	JSON + cosine Python	Simple, 0 dep.
10K-1M docs, Postgres	pgvector	Naturel, performant
> 1M docs	Qdrant/Weaviate	Vector stores dedies
Logs operations	PostgreSQL agent_logs	Audit, dashboards

9. Checklist integration DB

- Credentials via vault uniquement — jamais dans le workspace Git
- Utilisateur PostgreSQL dedie — droits minimaux (SELECT/INSERT sur tables necessaires)
- Pool de connexions — pgBouncer ou pool applicatif en production
- Timeout sur les requetes — statement_timeout obligatoire
- Requetes parametrees uniquement — jamais de f-string dans un SELECT
- Resultats ecrits dans data/ — l'agent lit un fichier, pas la connexion DB
- Index sur les colonnes frequentes — eviter les full scans
- Seuil de similarite RAG — filtrer avec score > 0.65-0.7
- Reindexation en batch — pas en temps reel a chaque ecriture
- Tests de regression — verifier les requetes de reference apres chaque modif du corpus

10. Conclusion de la serie

Ce dixième billet marque la fin de la série technique OpenClaw. En dix épisodes, BOTUM a documenté un déploiement réel : de l'installation jusqu'à l'intégration base de données et la recherche sémantique.

B1-B2 : concept, installation. B3-B4 : sécurité, vault. B5 : réseau d'agents. B6-B7 : choix LLM. B8-B9 : automatisation, mémoire. B10 : persistance structurée.

La série suivante explorera les dimensions business : ROI, PME vs Enterprise, edge. Articles sur botum.ca/blog.

Article complet : blog.botum.ca/openclaw-bases-donnees-postgresql-rag-recherche-semantique

Site web : www.botum.ca • contact@botum.ca