

800 MILLONES DE USUARIOS

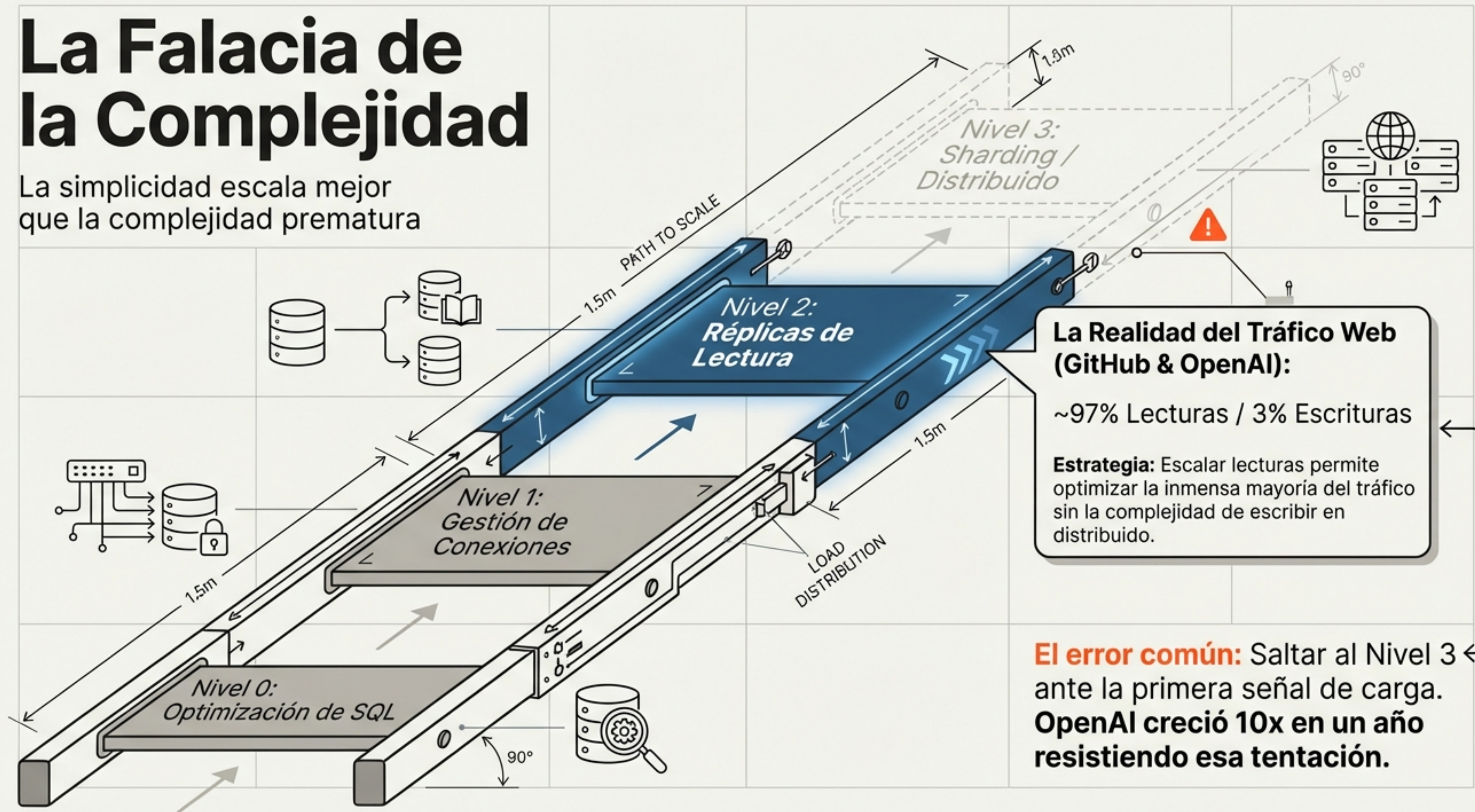


Una Base de Datos Primaria

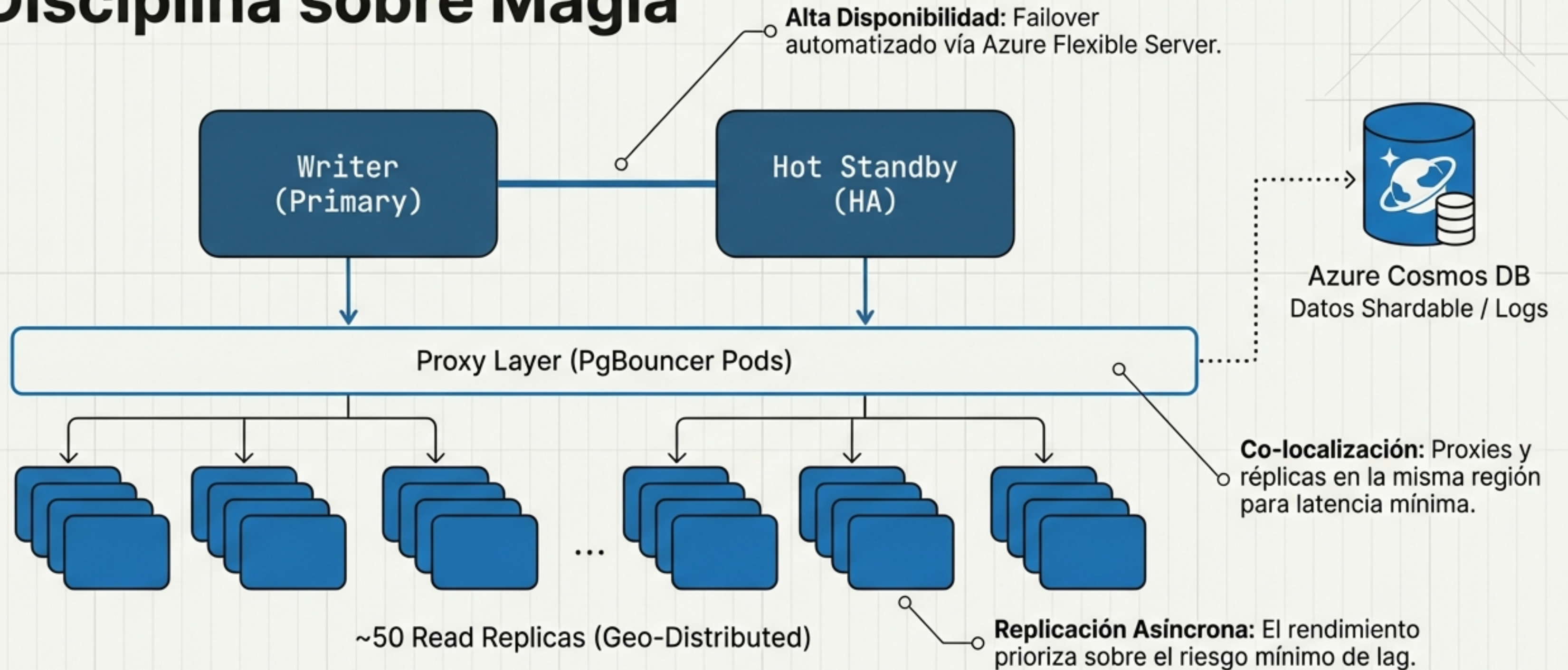
La arquitectura detrás de ChatGPT no es magia distribuida. Es dominio de los fundamentos. Sin sharding en el núcleo. Sin microservicios complejos. Solo PostgreSQL bien ejecutado.

La Falacia de la Complejidad

La simplicidad escala mejor que la complejidad prematura



Topología de Infraestructura: Disciplina sobre Magia



Gestión de Conexiones: La Capa Invisible

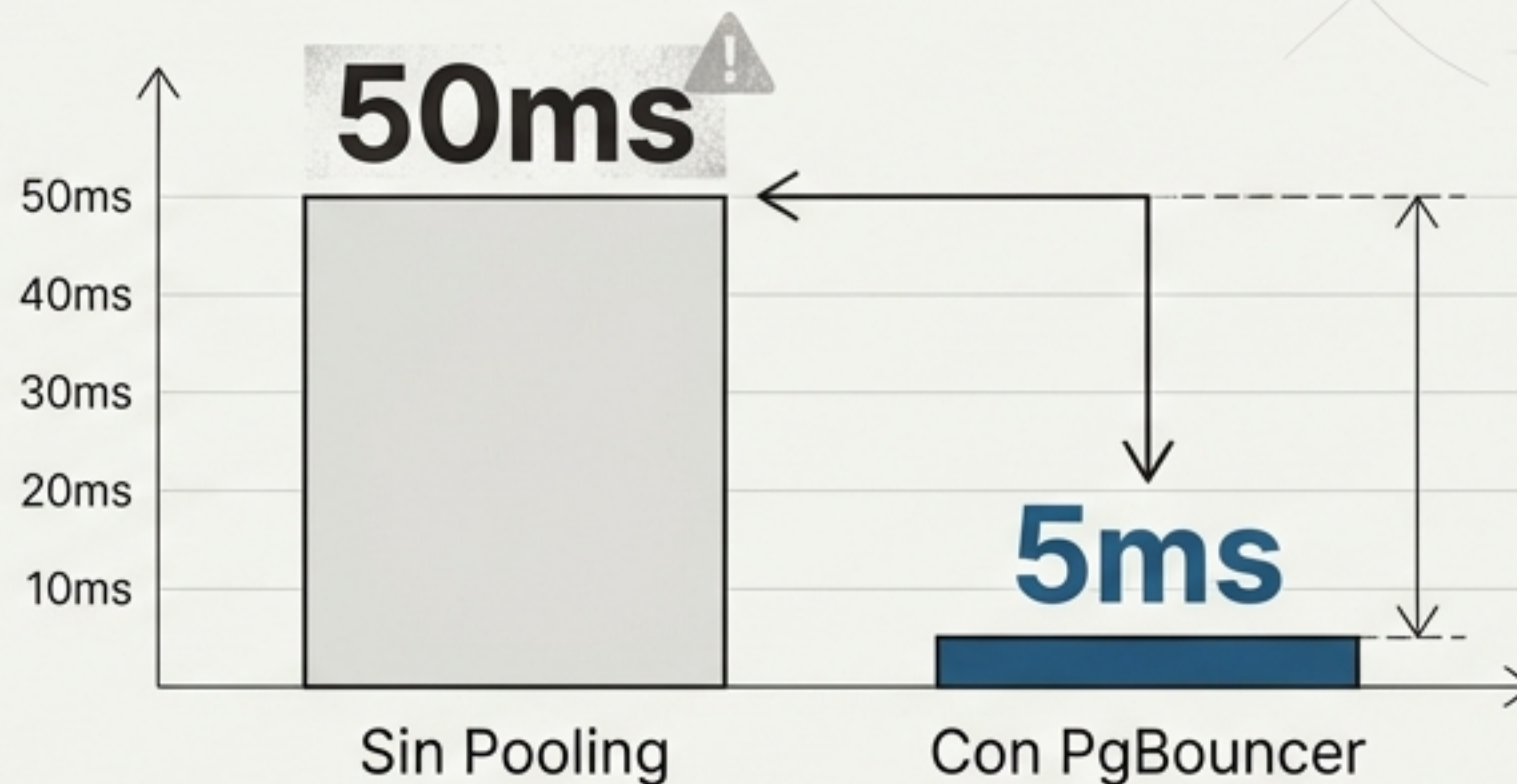
El Problema:

Postgres tiene un límite duro de conexiones (~5,000). Las “tormentas de conexiones” matan el rendimiento antes de tocar la CPU.

La Solución:

PgBouncer en modo “Transaction Pooling”.

Latencia de Establecimiento de Conexión



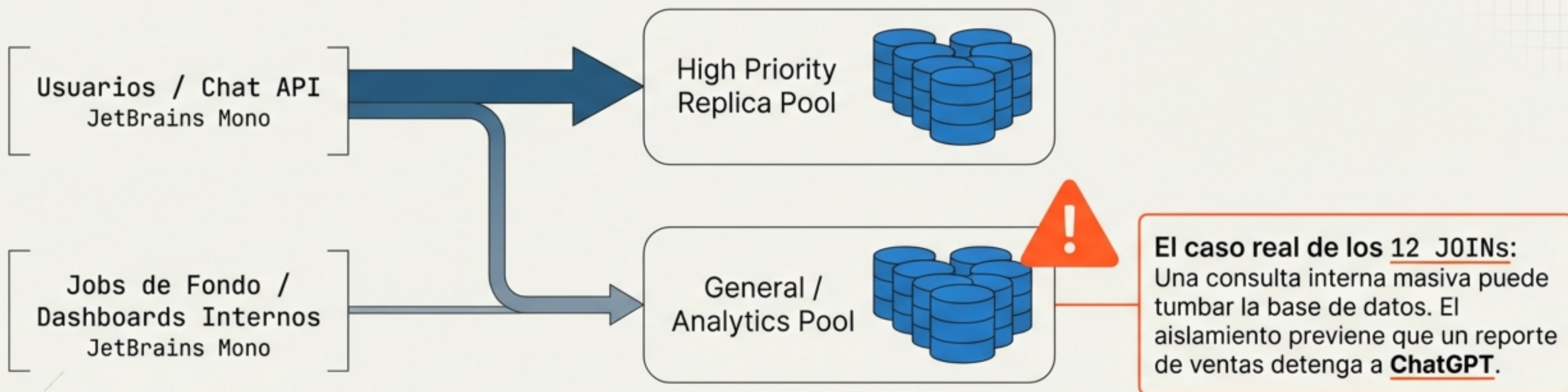
CONFIG_CRÍTICA:

```
pool_mode = transaction  
idle_transaction_timeout = 30s  
max_client_conn = 10000+
```

Engineering Editorial in Inter Tight

Aislamiento de Cargas (Noisy Neighbors)

Protegiendo el tráfico crítico de la analítica interna

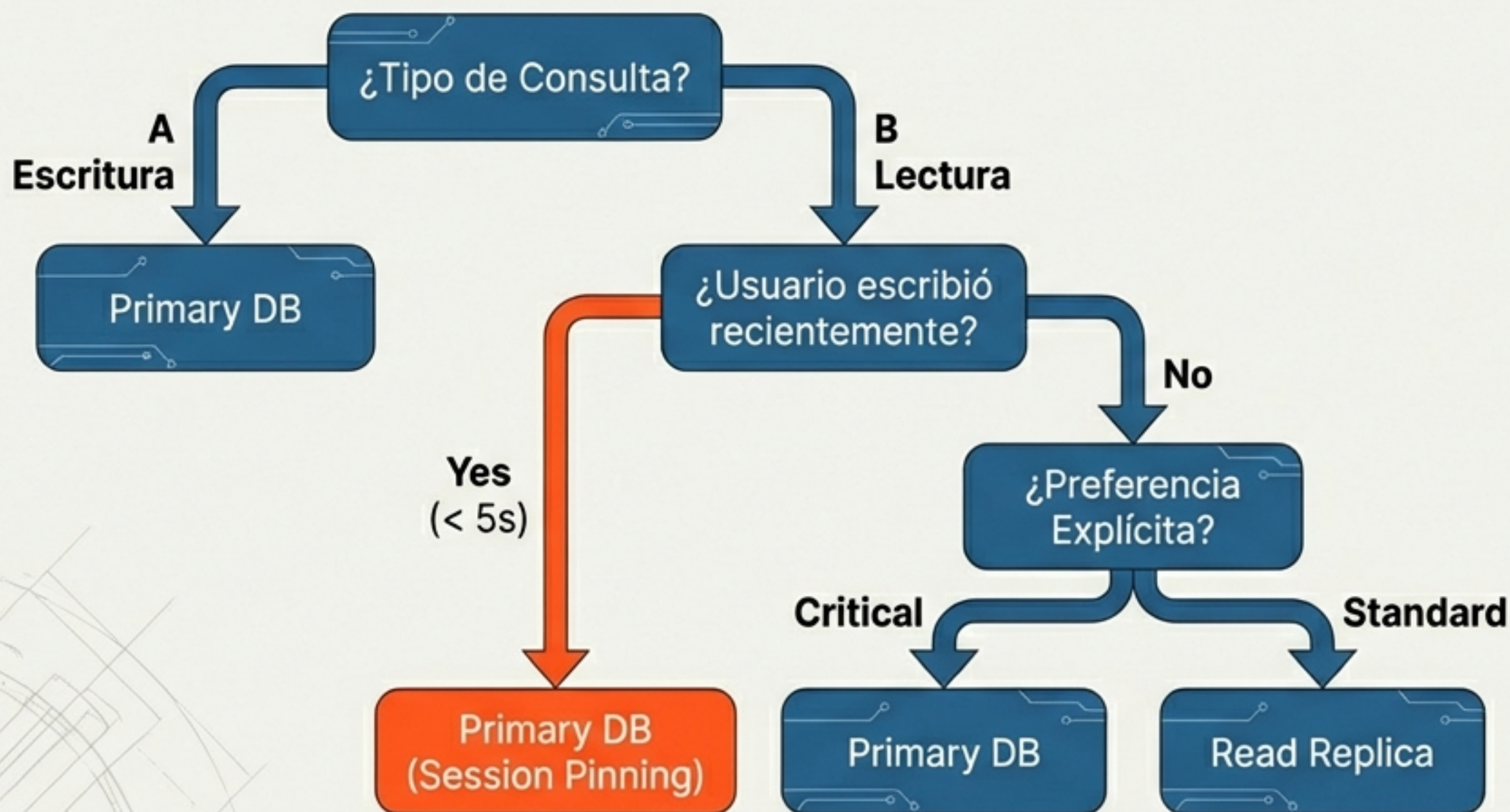


KEY TAKEAWAYS

1. Separación estricta de réplicas según SLA.
2. **Prioridad Alta:** Pagos, Auth, Chat en tiempo real.
3. **Prioridad Baja:** Analítica, reportes, tareas batch.

Estrategias de Enrutamiento y Consistencia

Resolviendo el problema de "Read-after-Write"



El Desafío:

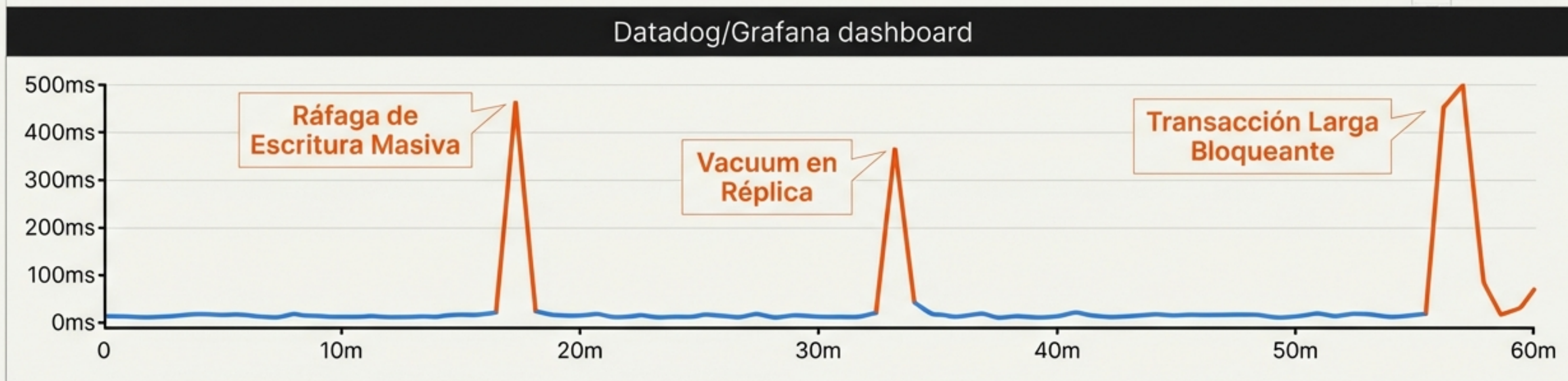
La replicación asíncrona crea "Lag". Un usuario puede no ver su propio post inmediatamente.

Estrategias:

- 1. Enrutamiento por Sesión:** Si escribiste hace poco, lees del primario.
- 2. Consciencia de Lag:** Si Réplica **Lag > 100ms**, redirigir al primario.

La Verdad sobre el "Replication Lag"

Entendiendo los picos y el monitoreo crítico en sistemas distribuidos



Mecánica

El Primario envía **WAL** (Write-Ahead Log) a 50 réplicas. El ancho de banda es finito.

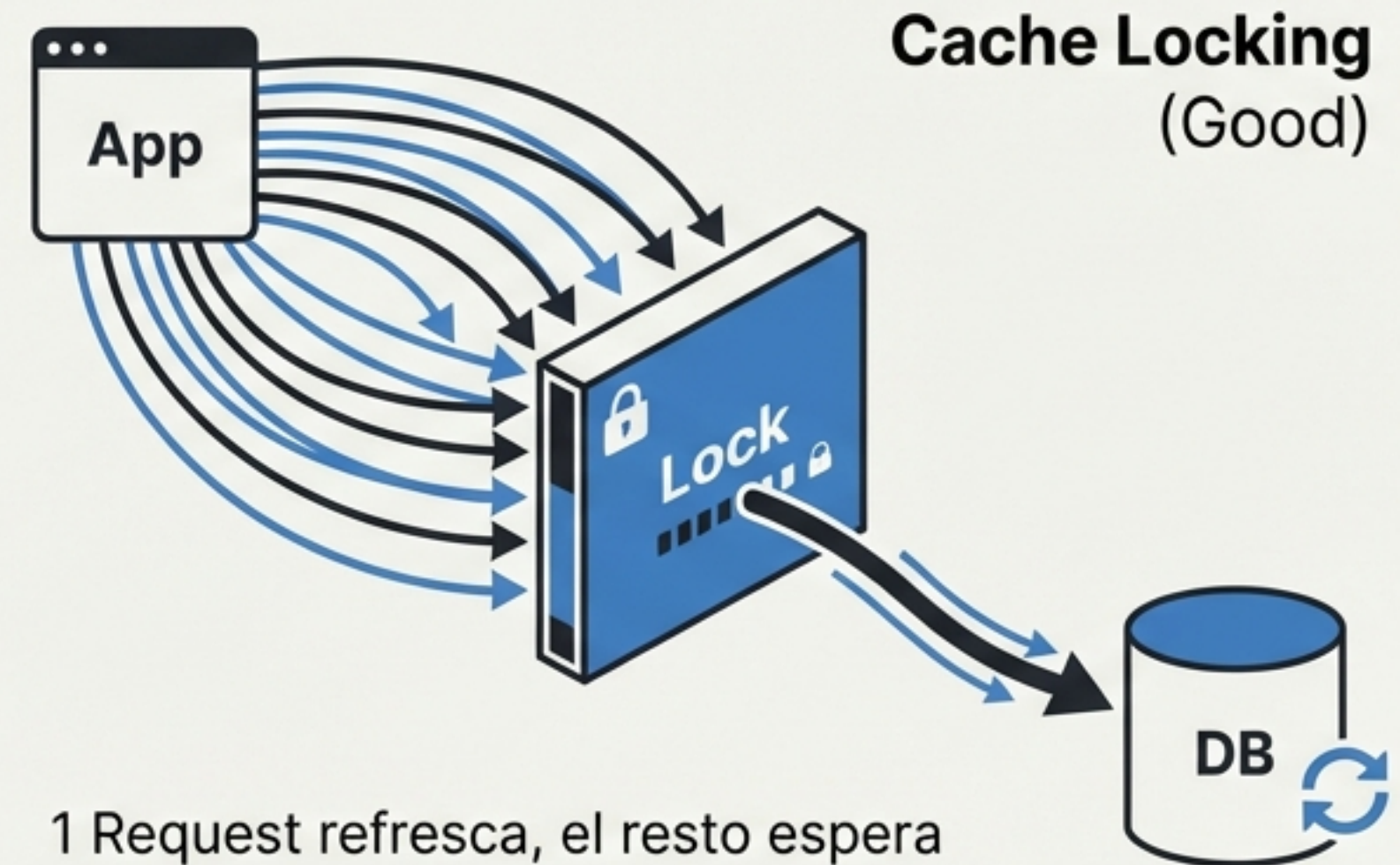
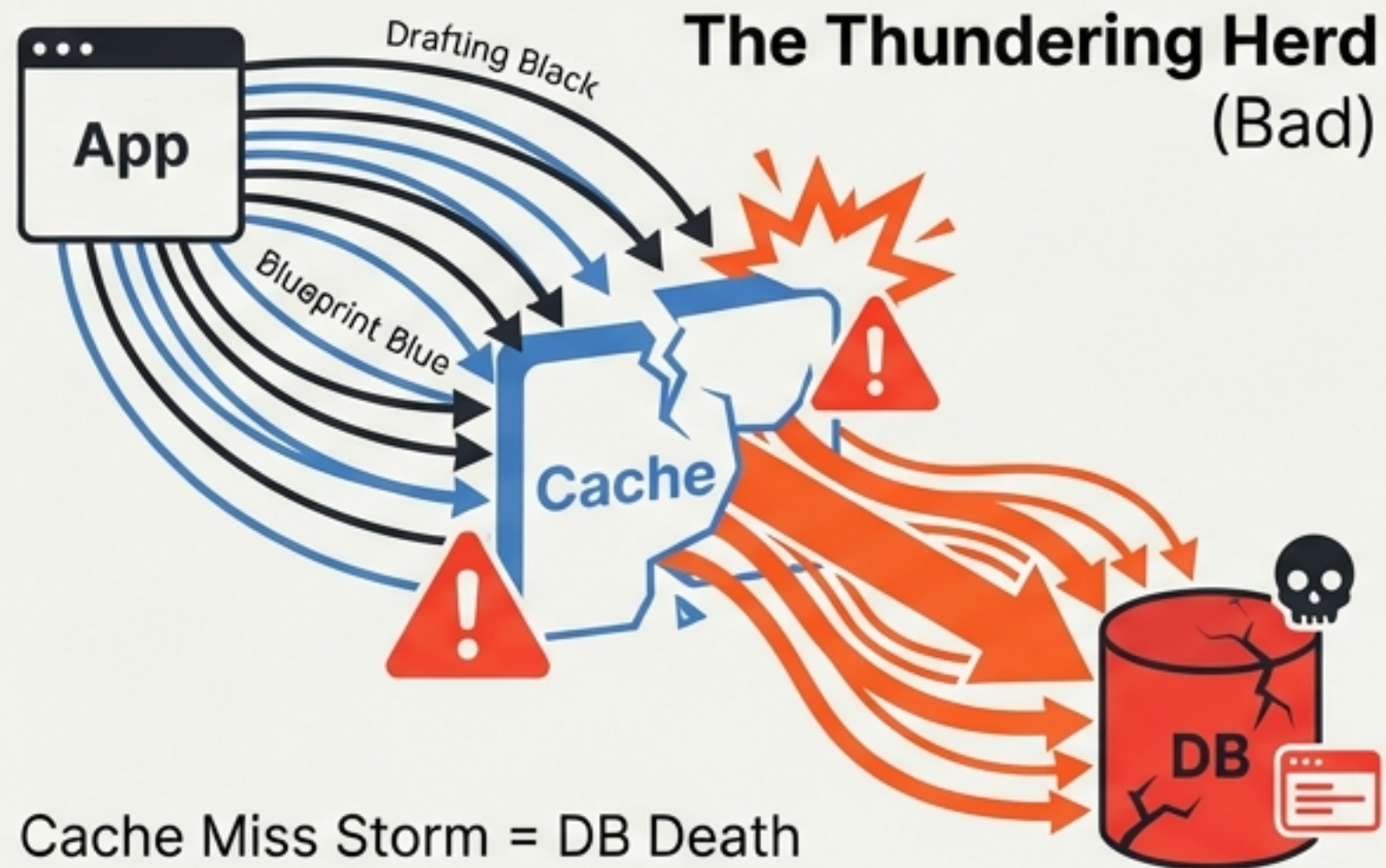
Monitoreo Crítico (pg_stat_replication):

No medir el promedio. Medir el dolor.

- Umbral Usuarios: >100ms (Alerta)
- Umbral Analítica: >30s (Advertencia)

Optimización de Consultas y Defensa de Caché

Estrategias para mitigar la sobrecarga y mejorar la eficiencia.



Higiene de SQL

- ⚙️ Eliminar JOINS masivos (caso OpenAI: 12 tablas -> Lógica en App).
- 🔍 Auditar consultas ORM (N+1, índices faltantes).
- 🕒 Lazy Writes: Diferir escrituras no críticas.

Reglas de Producción I: Infraestructura

Minas terrestres operacionales



Regla 1: Nunca confíe en el DNS durante un Failover.

Los TTLs mienten. Durante un failover, la app seguirá contactando al nodo muerto.

Solución: Usar proxies con health-checks activos o descubrimiento vía API (Patroni).

JetBrains Mono, Helvetica Now



Regla 2: PgBouncer y Prepared Statements son enemigos.

En modo 'transaction pooling', la conexión cambia entre sentencias. Los prepared statements fallarán.

Solución: Desactivar en el ORM o usar proxies avanzados (Odyssey).

JetBrains Mono, Helvetica Now

Reglas de Producción II: Operaciones Peligrosas

Operaciones de alto riesgo y sus mitigaciones.



Regla 3: Migraciones de esquema (DDL) bloquean replicación.

Un ALTER TABLE detiene el replay del WAL. El lag se dispara instantáneamente.

Solución: Usar CONCURRENTLY para índices. Timeouts estrictos (5s) para bloqueos.

JetBrains Mono, Helvetica Now



Regla 4: 'hot_standby_feedback' es un arma de doble filo.

- **ON:** Evita cancelar consultas largas de lectura, pero causa 'bloat' (basura) masiva en el Primario.
- **OFF:** Mantiene el Primario limpio, pero mata consultas largas en réplicas.

Consejo: OFF por defecto. Usar pool dedicado si es necesario.

JetBrains Mono, Helvetica Now

Reglas de Producción III: Lógica de Aplicación

Patrones de interacción y sus riesgos



Regla 5: Transacciones de lectura implícitas

Los ORMs a menudo abren transacciones para simples lecturas. Si quedan abiertas ("Idle in Transaction"), pausan la limpieza del WAL en réplicas.

Solución: Configurar
``idle_in_transaction_session_timeout``
agresivo (30s).

JetBrains Mono, Helvetica Now



Regla 6: Lógica de Retry Inteligente

Si falla una lectura en réplica, no reintentar en la misma (está saturada).

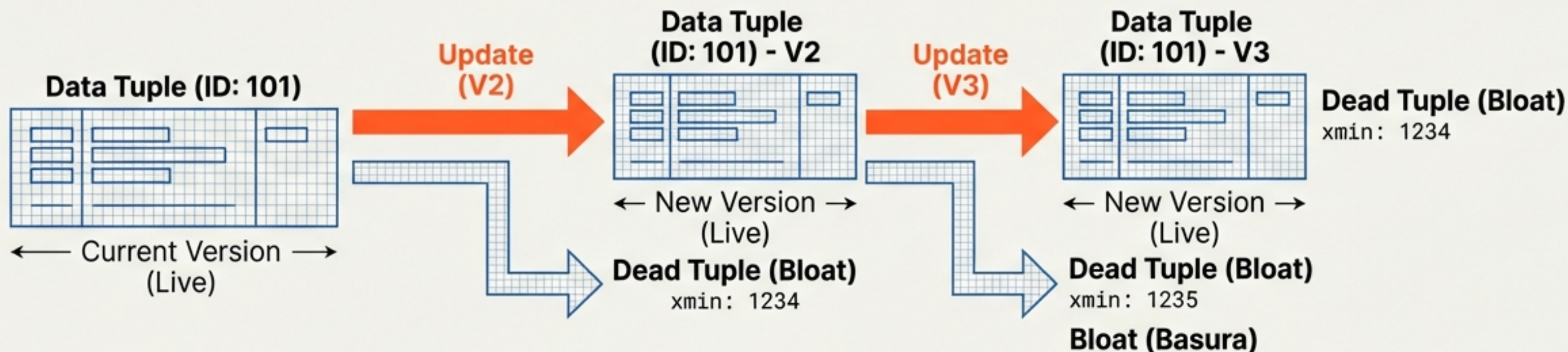
No reintentar inmediatamente en el Primario (riesgo de cascada).

Patrón: Probar otra réplica -> Fallback a Primario solo como último recurso.

JetBrains Mono, Helvetica Now

El Límite del Nivel 2: Cuello de Botella de Escritura

El Muro de Postgres: Las réplicas no escalan escrituras. Todo pasa por el Primario.



Amplificación de Escritura (MVCC): Cada update crea una nueva versión de la fila. Cargas masivas de escritura generan basura (bloat) y carga de Vacuum.

Solution

Estrategia de OpenAI: Migrar cargas 'Write-Heavy' y 'Shardable' (**Logs, Historial**) a **Azure Cosmos DB**. El núcleo transaccional permanece en **Postgres**.

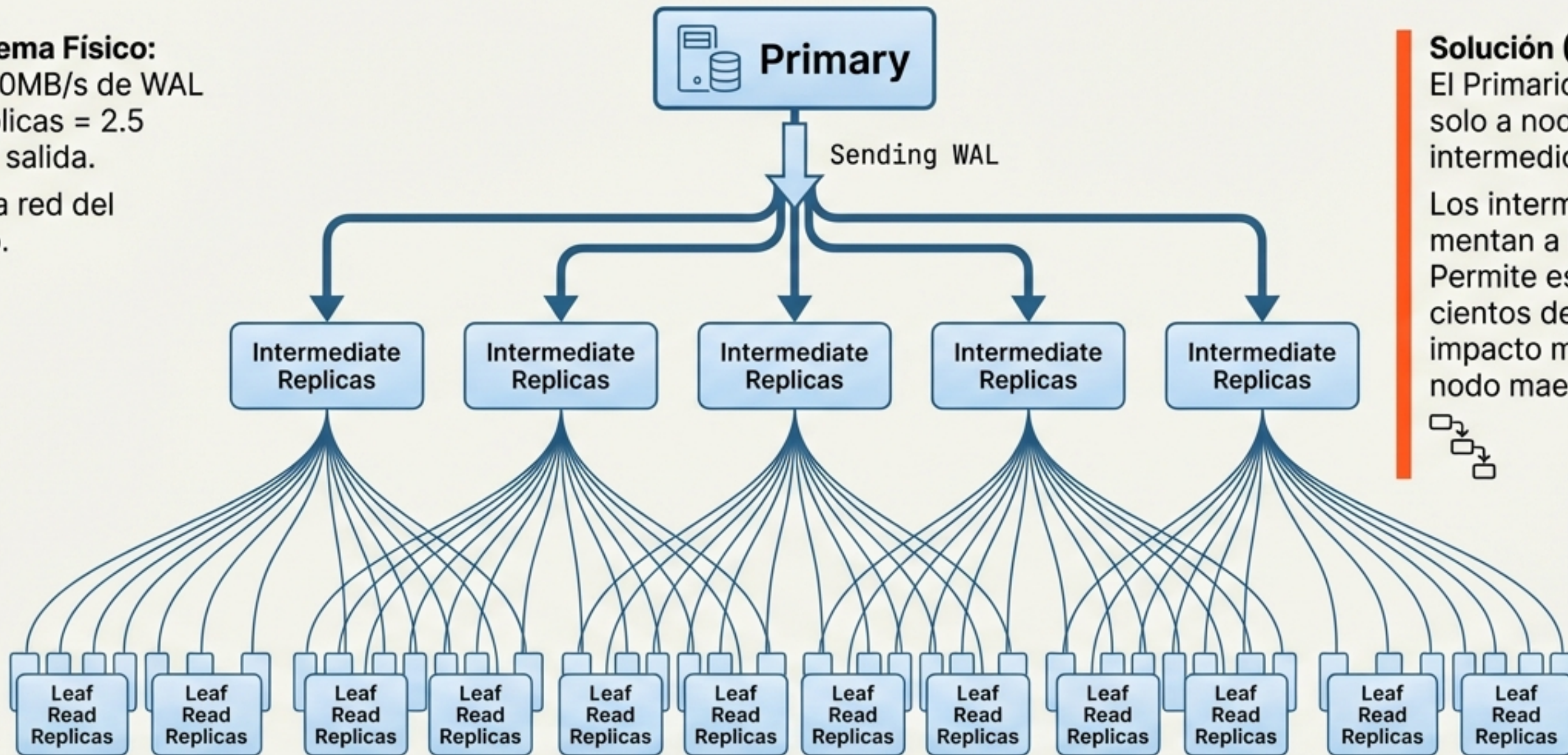
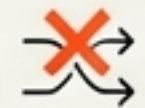
El Futuro: Replicación en Cascada

Superando el límite de ancho de banda del Primario

El Problema Físico:

Enviar 50MB/s de WAL
a 50 réplicas = 2.5
GB/s de salida.

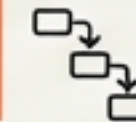
Satura la red del
Primario.



Solución (En Pruebas):

El Primario alimenta
solo a nodos
intermedios.

Los intermedios ali-
mentan a las hojas.
Permite escalar a
cientos de réplicas con
impacto mínimo en el
nodo maestro.



Lista de Verificación (Checklist)

Lista de Verificación (Checklist)

Antes de Escalar	Implementación	Operaciones
 Verificar ratio de lectura (>70% para ser efectivo)	 Configurar PgBouncer (Transaction Mode)	 Monitorear lag con alertas de umbral
 Optimizar consultas (Eliminar ILIKE, reducir Joins)	 Estrategia de Enrutamiento (Session-based)	 Auditar transacciones (Idle timeouts)
 Definir consultas 'Replica-Safe'	 Aislar tráfico por prioridad (Pools dedicados)	 Validar que los retries no maten al primario

MAESTRÍA SOBRE **MAGIA**

OpenAI no usa una sola base de datos primaria porque les falte talento para sistemas distribuidos.

Es una elección deliberada.

La lección no es 'nunca usar sharding'. La lección es:

**Domine cada nivel de la escalera
antes de subir al siguiente.**

CONNECTION POOLING • QUERY OPTIMIZATION • READ REPLICAS • DISCIPLINA