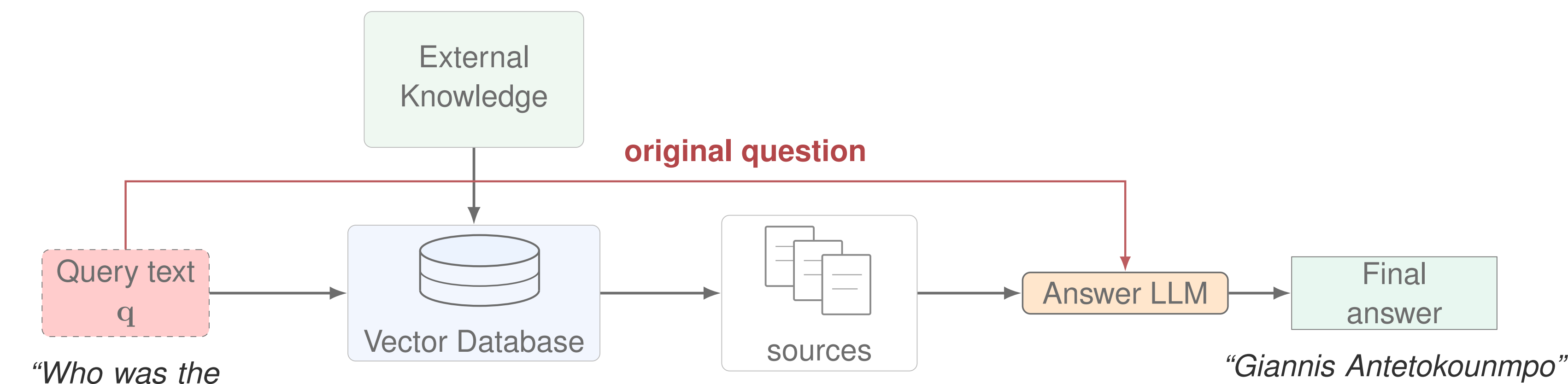


What is Retrieval-Augmented Generation (RAG)?



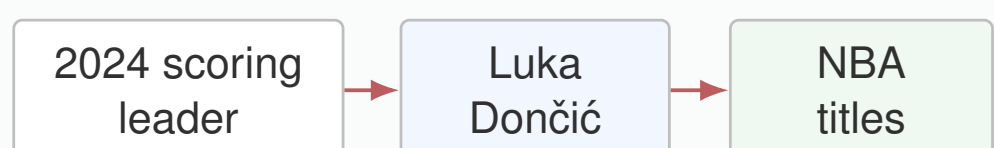
In standard deployments, the same knowledge index and retriever setting serve every query, making retrieval a fixed bottleneck when the question type changes.

Query needs are heterogeneous

Multi-hop

Needs a reasoning chain.

"How many NBA titles has the 2023–24 scoring leader won?"



Best suited:
Graph / context-expansion retrieval

Diverse evidence

Needs coverage, not redundancy.

"What materials are sports balls made of?"



Best suited:
Diversity-aware retrieval

Specialized vocabulary

Needs domain semantics.

"In *NSCLC*, which *EGFR* mutation is associated with resistance to first-generation TKIs?"

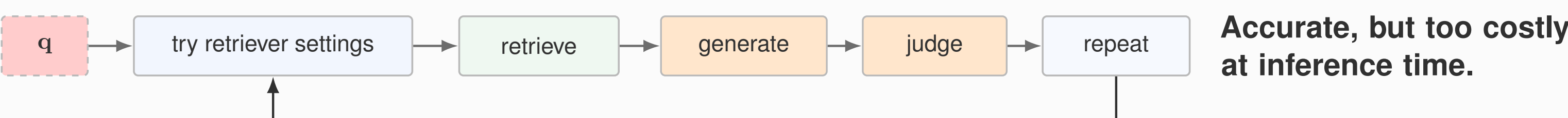
Best suited:
Biomedical embedding retriever

No single retriever is uniformly best across these needs.

Can we adapt retrieval efficiently?

Different queries need different retrievers; the challenge is adapting cheaply.

Online adaptive search per query



Accurate, but too costly at inference time.

Retriever portfolio

Offline: choose k retrievers with complementary behaviors.
Online: route to a few of them, retrieve in parallel, and answer.

Question: Which k retrievers should we choose?

How do we measure the quality of a portfolio?

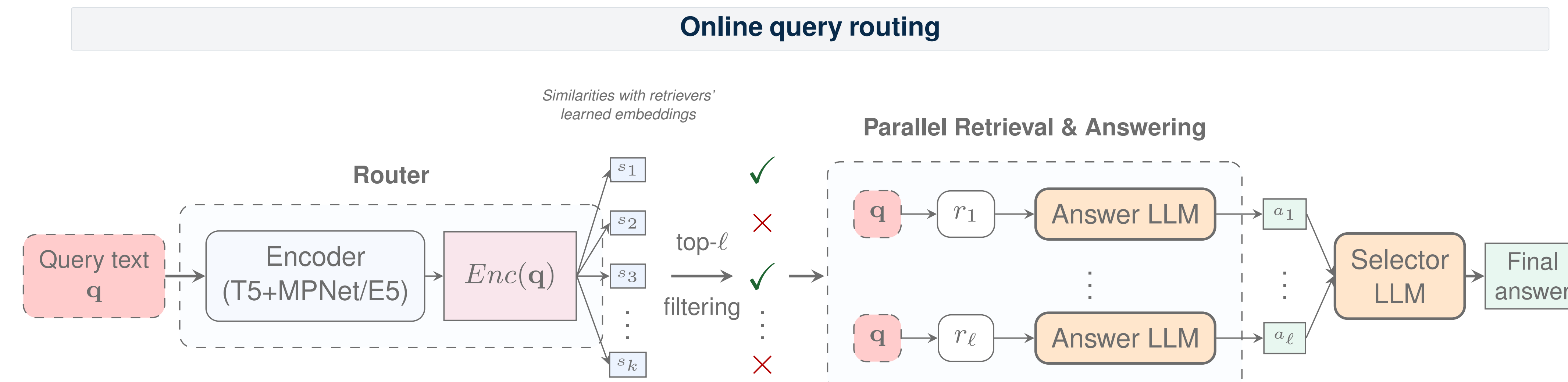
$$F(S) = \mathbb{E}_{q \sim \mathcal{D}} \left[\max_{r \in S} \text{score}(q, r) \right]$$

Promotes complementarity by definition: for each query, at least one member should be strong.

Submodular objective $\Rightarrow$ greedy efficiently computes a near-optimal portfolio.

The problem is not about finding good retrievers; it's about finding complementary ones.

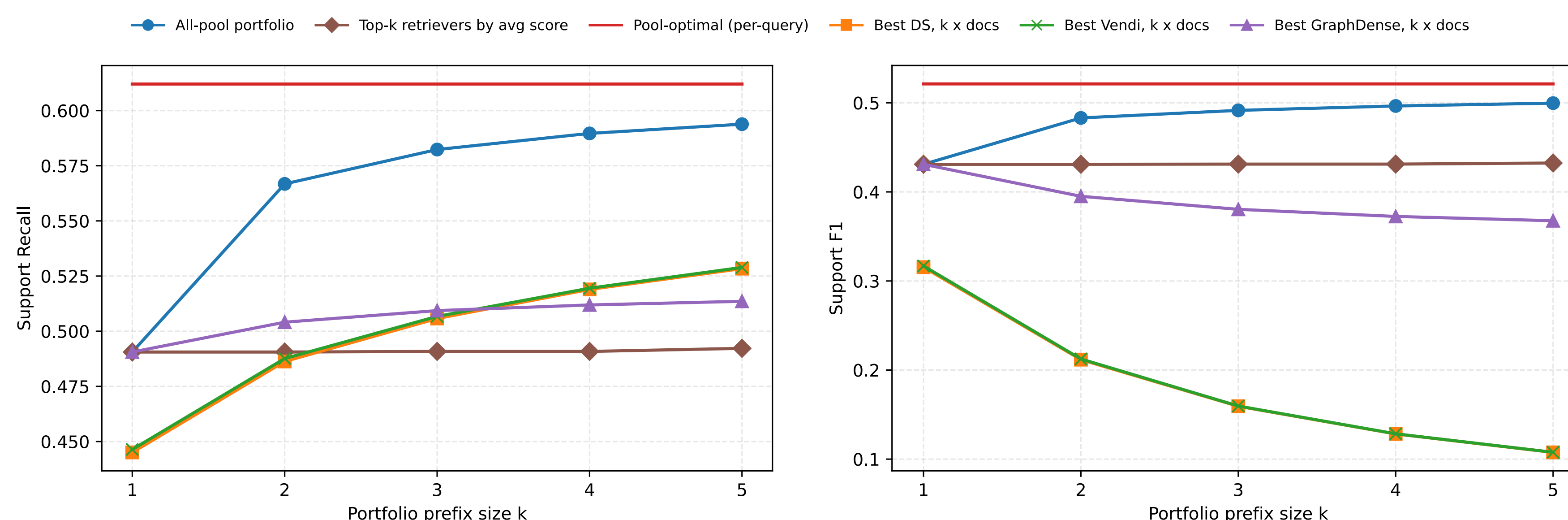
Method



Experimental setup



Portfolios improve retrieval coverage

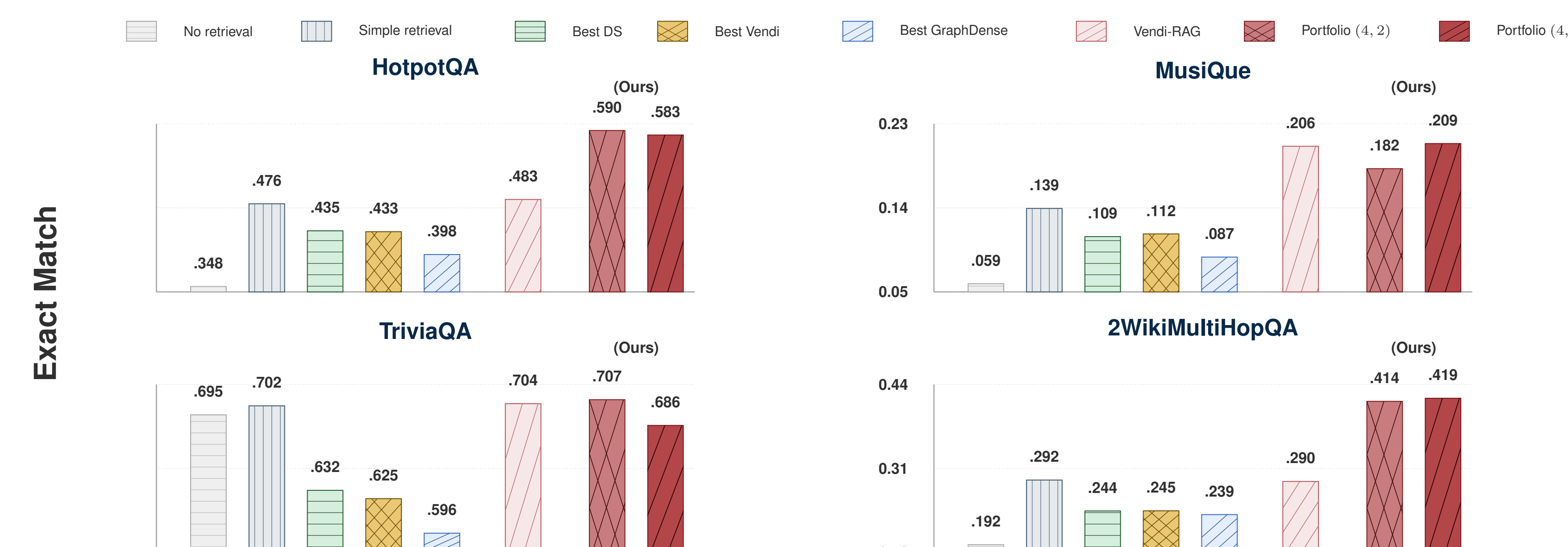


Takeaway. Portfolios outperform all single-retriever baselines, even when those baselines retrieve more documents.

Key point. The gains come from complementarity: top- k retrievers by average score are often redundant, while a learned portfolio of 5 members nearly matches pool-optimal adaptation.

Retrieval gains translate to QA accuracy

Answer Exact Match is calculated for Llama-3.1-70B-Instruct according to the labeled answers of each dataset. All baselines use the same answer model.

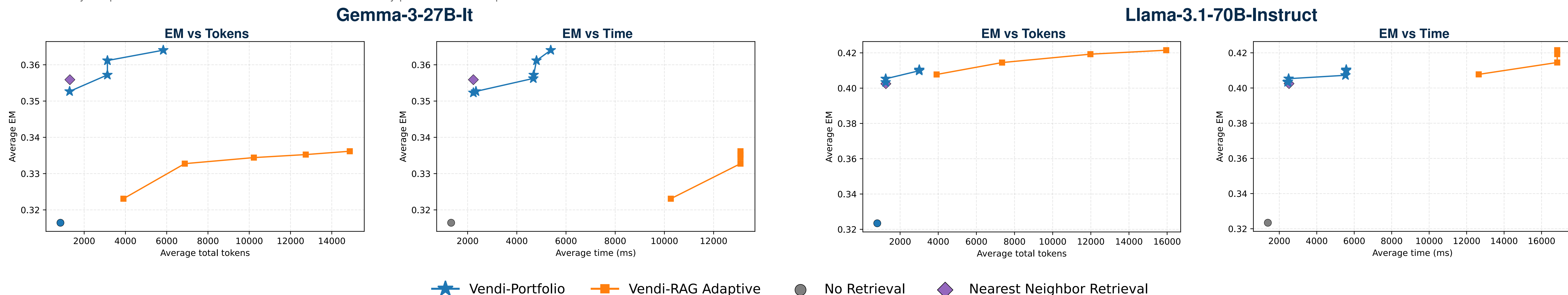


Takeaway. Routed portfolios give the highest Exact Match in all four datasets.

Key point. Complementary retrieval improves final answers, outperforming static baselines and online Vendi-RAG.

Offline adaptation is more efficient

Controlled Vendi-only comparison: same MPNet backbone and same Vendi diversity-parameter search space.



Takeaway. Portfolios avoid expensive per-query search while retaining strong Exact Match.