

Revisiting RaBitQ and TurboQuant: A Symmetric Comparison of Methods, Theory, and Experiments

Jianyang Gao¹, Yutong Gou², Yuexuan Xu², Jifan Shi², Yongyi Yang³, Shuolin Li⁴, Raymond Chi-Wing Wong⁵, Cheng Long²

¹ETH Zurich ²Nanyang Technological University ³University of Michigan ⁴Tsinghua University ⁵Hong Kong University of Science and Technology

Methodology

Preprocessing: **random rotation** / **JL transformation** for both methods.

RaBitQ codebook: **uniform, shifted-integer grid** — pick rescaling factor t , round to nearest grid point, store 1 scalar.

TurboQuant codebook: **non-uniform, k-means centroids** (Lloyd–Max) - quantize to nearest centroid, apply QJL to residual for unbiased estimate.

IP estimate: RaBitQ uses **native unsigned-int arithmetic (no decoding)**; TurboQuant needs codebook lookup/decoding.

Code size: 1 scalar + D unsigned B-bit ints (TurboQuant unbiased mode: +1 scalar).

Theoretical Guarantees

RaBitQ — optimal. Matches the Alon–Klartag (FOCS 2017) space–distortion trade-off: a sub-Gaussian tail bound gives bit-width **$B = \Theta(\log \log(1/\delta))$** .

TurboQuant — suboptimal. Only a variance (MSE) bound is proved. Chebyshev's inequality => **$B = \Theta(\log(1/\delta))$** - exponentially worse in δ .

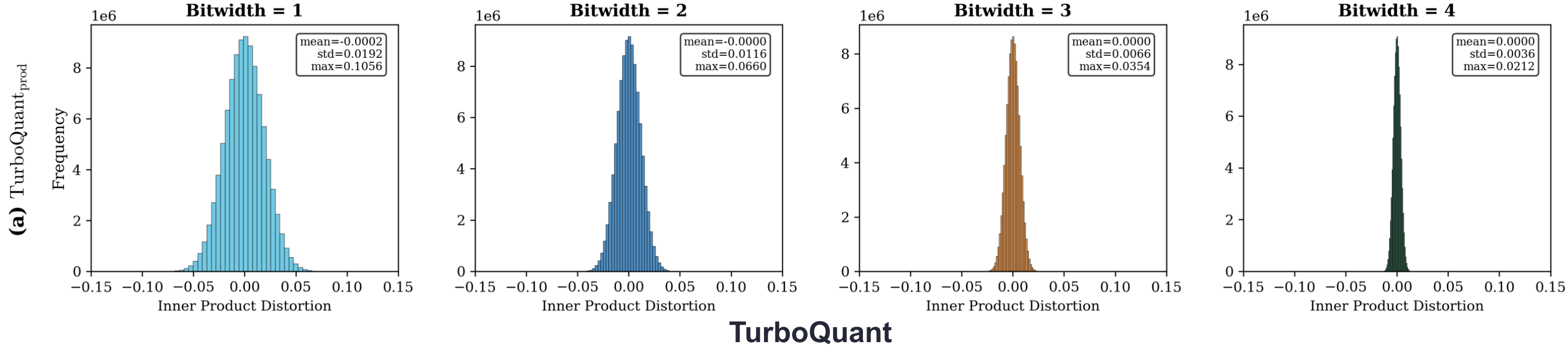
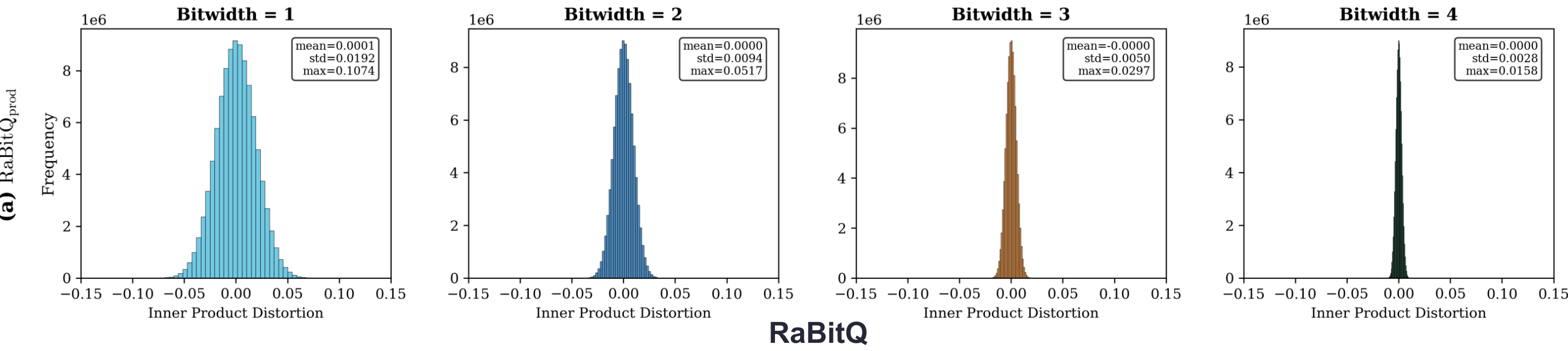
Experimental Setup

Datasets: DBpedia-1536, GloVe-200, OpenAI3-1536/3072.

Hardware: Nvidia A100 (80GB) + dual Xeon Gold 6418H (96 threads)

Code: github.com/VectorDB-NTU/rabitq-turboquant-comparison

Quantization Accuracy



TurboQuant shows no consistent advantage. For the unbiased variants, RaBitQ has smaller std. and max error across most tested bit-widths.

Quantization Efficiency (100K vectors, 4-bit, in seconds)

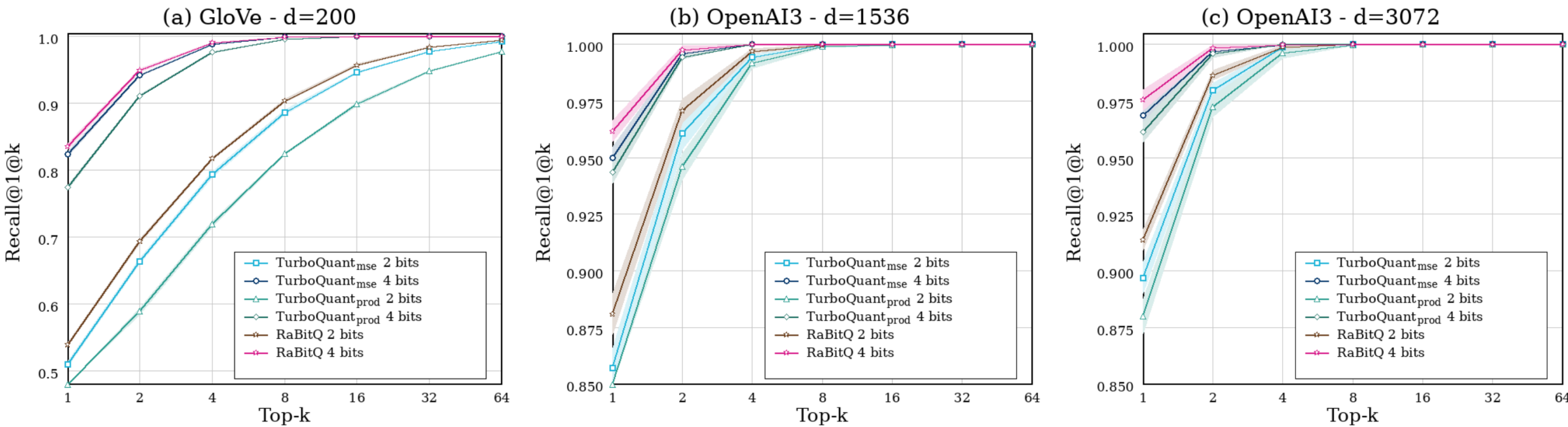
Approach	d=200	d=1536	d=3072
RaBitQ (CPU)	0.125	1.003	4.176
RaBitQ _{fastOn-FWHT} (CPU)	0.085	0.143	0.218
RaBitQ (GPU)	0.009	0.065	0.152
RaBitQ _{fastOn-FWHT} (GPU)	0.003	0.008	0.013
TurboQuant (GPU)	0.011	0.114	0.276

Results in TurboQuant Paper			
Approach	d=200	d=1536	d=3072
Product Quantization	37.04	239.75	494.42
RabitQ	597.25	2267.59	3957.19
TURBOQUANT	0.0007	0.0013	0.0021

- **RaBitQ on single-core CPU, multithreading disabled**
- **TurboQuant on A100 GPU**

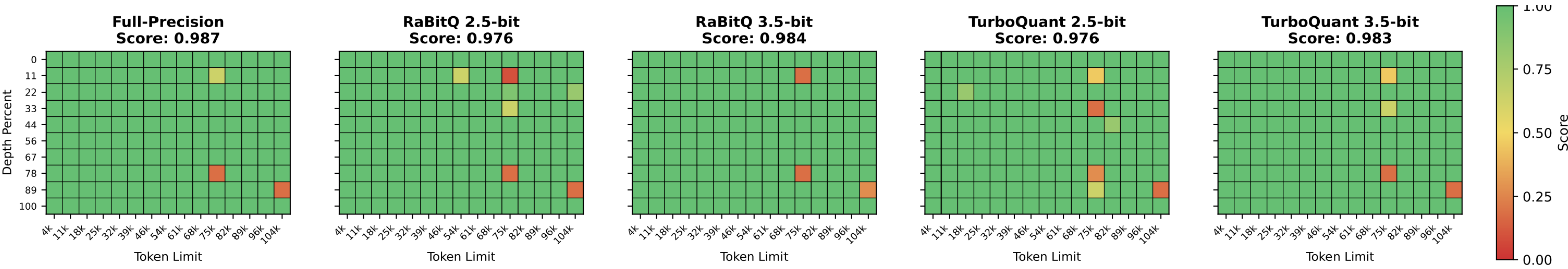
On identical GPU hardware, RaBitQ is 1.2–1.8× faster than TurboQuant across all tested dimensions.

Nearest Neighbor Search



RaBitQ consistently achieves higher recall than both TurboQuant variants across all datasets and B settings.

KV Cache Quantization (NIAH and LongBench-E)



RaBitQ and TurboQuant are matched - NIAH scores differ by ≤ 0.002 ; LongBench-E overall average differs by ≤ 0.58 pts across all tested configurations.

Reproducibility Issues

1. RaBitQ baseline in the TurboQuant paper was run single-core CPU, multithreading disabled, in Python — vs. TurboQuant on an A100 GPU. **Not disclosed.**
2. Reported TurboQuant quantization times are ~2 orders of magnitude off from what the released code reproduces on the same GPU.
3. Reported RaBitQ recall values fall below our measured 1-std band across 10 seeded runs for ANN search.

Conclusion

1. RaBitQ **matches or beats** TurboQuant in accuracy.
2. RaBitQ is **substantially faster** in quantization.
3. RaBitQ **wins consistently** on NN search recall.
4. The two are **comparable** in KV-cache quantization.

Paper



Code

