



SEARCHING FOR THE MAXIMUM ELEMENT IN A VECTOR USING BUILT-IN FUNCTION `max()` AND LOCATING THE MAXIMUM ELEMENT FROM THE VECTOR USING POPULAR SEARCHING TECHNIQUE BINARY SEARCH METHODOLOGY: A CASE STUDY

R. Daranish

6780

Class- XII , 2026-27 , Sainik School Amaravathinagar Post: Amaravathinagar, Udumalpet Taluka, Tirupur Dt, Tamilnadu State

ABSTRACT

This paper looks at a question that sounds almost too simple to deserve a research paper of its own: what is the fastest, most reliable way to find the largest value sitting inside a vector? Most programmers reach for Python's built-in `max()` function without a second thought, and for good reason - it is short, it is readable, and it almost always works. But once a vector is already sorted, a second option opens up. Because the maximum element of a sorted list sits at a fixed, predictable position, a binary-search style walk toward that position can, in principle, locate it without touching every element along the way.

We set out to test this idea properly rather than take it on faith. Three additional approaches were brought into the comparison alongside `max()` and the sorted binary-search method - a classical divide-and-conquer `max` finder, a tournament-style pairwise elimination method, and a simple threaded variant of `max()` intended to see whether parallel execution offers any real benefit on commodity hardware. Each method was implemented in Python, profiled across vector sizes ranging from one hundred elements up to one million elements, and measured for execution time, comparison count, and peak memory consumption.

The results, on the whole, confirm what the theory predicts but with a few wrinkles that are easy to miss on paper. `max()` remains the most practical default for everyday, unsorted data because of how tightly CPython's C implementation is written. The binary-search approach is genuinely fast once a vector is sorted, although the cost of sorting itself, when it has to be paid, erases nearly all of the advantage. Divide-and-conquer performs comparably to `max()` in terms of comparisons but suffers from recursion overhead in Python. The tournament method, somewhat surprisingly, was the most comparison-efficient of all, needing close to $n - 1$ comparisons in the best structured cases, though its real-world runtime did not always reflect that theoretical elegance because of interpreter-level overhead. We close with a discussion of where each technique earns its keep and offer some practical recommendations for choosing between them.

1. INTRODUCTION

Vectors, or one-dimensional arrays, are about as basic a data structure as computing offers, and yet a surprising amount of real engineering time is spent simply scanning them for an extreme value. Whether the task is finding the highest temperature recorded by a sensor over a day, the largest transaction in a batch of financial records, or the peak signal in an audio buffer, the underlying operation is the same: walk through a collection of numbers and report the biggest one.

Python makes this trivial with `max()`, and most developers never think twice about it. That convenience, however, can hide an interesting question. If the data happens to already be sorted - which is not as rare a situation as it might first appear, since many pipelines sort data for other reasons anyway - is there a smarter way to grab the maximum than scanning the whole thing again?

This paper takes that question seriously. We are not trying to argue that anyone should abandon `max()` in everyday code; that would be a strange thing to claim. Instead, the goal is to build a clear, evidence-backed picture of when alternative strategies genuinely pay off, when they merely look clever on paper, and

what the practical cost of "being clever" turns out to be once Python's interpreter overhead enters the picture.

2. RELATED WORK

The literature on searching and selection algorithms is enormous, and rather than try to summarise all of it we focus on the strands that are directly relevant here. Cormen, Leiserson, Rivest and Stein's treatment of selection algorithms remains the standard reference for understanding why finding a single extreme value in an unordered collection cannot, in the comparison model, be done in fewer than $n - 1$ comparisons. That lower bound is the yardstick against which every method in this paper is measured.

Knuth's earlier volumes on sorting and searching cover binary search in detail and note, almost in passing, that binary search is only meaningful when an ordering invariant already holds; this single observation turns out to be the crux of our entire investigation.

On the practical side, the Python documentation and the CPython source code show that `max()` is implemented as a tight



loop in C, with very little of the per-element overhead that a hand-written Python loop would carry. This explains, at least qualitatively, why naive "improvements" written in pure Python often lose to the built-in function even when they require fewer asymptotic operations.

Finally, work on tournament-based selection (sometimes called the "knockout" method) goes back to early analyses of parallel computation, where pairwise comparison trees were studied as a way of finding both the maximum and second-maximum element with a near-optimal number of comparisons. We borrow that idea here mainly to see how it behaves once it is forced through the comparatively slow path of the Python interpreter rather than analysed only in the abstract.

3. RESEARCH OBJECTIVES

The study was guided by five concrete objectives. First, to measure and compare the runtime behaviour of `max()`, sorted binary search, divide-and-conquer, the tournament method, and a threaded variant across a wide range of vector sizes. Second, to verify whether the theoretical complexity of each method ($O(n)$ versus $O(\log n)$ versus $O(n)$ with fewer comparisons) actually shows up in measured wall-clock time once Python-specific overheads are accounted for. Third, to quantify memory usage, since some of these methods - particularly the recursive ones - carry stack and object-creation costs that a simple Big-O analysis tends to ignore. Fourth, to test the common claim that "sorting first and then binary searching" is worthwhile, by explicitly including the cost of sorting in at least one experiment. And fifth, to translate all of the above into a short, usable set of recommendations for practitioners who are not interested in the theory but do want to make a sensible choice.

4. RESEARCH METHODOLOGY

Random integer vectors were generated using Python's random module, with sizes of 100, 1,000, 10,000, 50,000, 100,000, 500,000 and 1,000,000 elements. For the algorithms that require sorted input, vectors were pre-sorted using Python's built-in `Timsort` (`list.sort()`) and the sorting time was recorded separately from the search time so that the two costs could be examined both together and apart. Every timing figure reported in this paper is the mean of 30 independent runs, with the fastest and slowest run discarded to reduce the influence of background system noise, a precaution that turned out to matter more than expected on the shared machine used for testing. Python's `time.perf_counter()` was used for all timing because it has the best resolution available on the host operating system. Memory figures were captured with the `tracemalloc` module, sampled at the point of peak usage during each call.

5. ALGORITHM – BUILT-IN `max()`

`max()` is conceptually a linear scan: it keeps a running candidate for the largest value seen so far, compares each new element against that candidate, and replaces it whenever a larger element turns up. What makes it fast in practice is not the algorithm itself - any first-year programmer could write the same logic - but the fact that CPython implements the loop in C,

avoiding the per-iteration bytecode dispatch overhead that a pure-Python loop would otherwise incur.

6. SOURCE CODE – `max()`

```
def find_max_builtin(vector):
    """Return the maximum element using
    Python's built-in max()."""
    return max(vector)

# Sanity check
if __name__ == "__main__":
    sample = [12, 45, 7, 89, 3, 56]
    print("Max:", find_max_builtin(sample))
```

7. ALGORITHM – BINARY SEARCH METHOD (SORTED VECTORS ONLY)

When a vector is sorted in ascending order, its maximum element always sits at the final index, so there is technically no "searching" needed at all - a single index lookup, `vector[-1]`, would do the job in constant time. We nonetheless implement a binary-search-flavoured version that narrows a `[low, high]` window toward the rightmost boundary, both because it mirrors how a reader unfamiliar with the trick might first approach the problem, and because it lets us measure the overhead a loop-based approach adds compared with the trivial index lookup. Each iteration halves the search window, giving the method its $O(\log n)$ character on paper, although in this particular degenerate case every iteration simply pushes low up toward high without ever discarding the half containing the answer, which is itself an interesting illustration of how binary search behaves at the edge of its usual assumptions.

8. SOURCE CODE – BINARY SEARCH (RIGHTMOST ELEMENT)

```
def find_max_binary(vector):
    """Locate the maximum of a SORTED
    vector via boundary narrowing."""
    low = 0
    high = len(vector) - 1
    while low < high:
        mid = (low + high + 1) // 2
        low = mid
    return vector[low]

def find_max_trivial(sorted_vector):
    """For comparison: the honest  $O(1)$  way
    once data is sorted."""
    return sorted_vector[-1]
```

9. ALGORITHM – DIVIDE AND CONQUER MAXIMUM

A more classical alternative is the divide-and-conquer max finder, which splits the vector in half, recursively finds the maximum of each half, and then returns the larger of the two results. This is the same family of technique used in merge sort, and it has the same recurrence relation, $T(n) = 2T(n/2) + O(1)$, which resolves to $O(n)$ overall - the same asymptotic class as the



linear scan, but arrived at through a very different mechanism. The interest here is less about beating max() and more about whether recursion overhead in Python eats into the theoretical tie.

10. SOURCE CODE – DIVIDE AND CONQUER

```
import sys
sys.setrecursionlimit(2000)

def find_max_divide_conquer(vector, low=0,
high=None):
    if high is None:
        high = len(vector) - 1
    if low == high:
        return vector[low]
    mid = (low + high) // 2
    left_max =
find_max_divide_conquer(vector, low, mid)
    right_max =
find_max_divide_conquer(vector, mid + 1,
high)
    return left_max if left_max > right_max
else right_max
```

11. ALGORITHM – TOURNAMENT (KNOCKOUT) METHOD

The tournament method pairs elements up, compares each pair, and keeps only the winners, repeating the process until a single champion remains - much like a single-elimination sports bracket. For n elements this requires roughly n - 1 comparisons in total, which matches the theoretical lower bound for finding a maximum by comparison and is, in that narrow sense, optimal. The method is also attractive because the comparison tree it builds can, in principle, be reused to find the second-largest element with very few extra comparisons, although we did not pursue that extension here.

12. SOURCE CODE – TOURNAMENT METHOD

```
def find_max_tournament(vector):
    current_round = list(vector)
    comparisons = 0
    while len(current_round) > 1:
        next_round = []
        for i in range(0,
len(current_round) - 1, 2):
            comparisons += 1
            winner = current_round[i] if
current_round[i] >= current_round[i + 1]
else current_round[i + 1]
            next_round.append(winner)
        if len(current_round) % 2 == 1:
            next_round.append(current_round[-1])
        current_round = next_round
    return current_round[0], comparisons
```

13. ALGORITHM – THREADED max() VARIANT

To check whether parallelism offers any real benefit, the vector was split into four roughly equal chunks, each handed to its own worker thread which computed a local maximum using the built-in max(), and the four partial results were then combined with a final max() call. Because CPython's Global Interpreter Lock prevents pure-Python bytecode from running on more than one core at a time, we did not expect dramatic speed-ups, but the experiment was included anyway since "just use threads" is a suggestion that comes up often enough in informal discussions that it deserved a direct test rather than a dismissal.

14. SOURCE CODE – THREADED VARIANT

```
from concurrent.futures import
ThreadPoolExecutor

def _chunk(vector, n_chunks):
    size = len(vector) // n_chunks
    return [vector[i:i + size] for i in
range(0, len(vector), size)]

def find_max_threaded(vector, n_chunks=4):
    chunks = _chunk(vector, n_chunks)
    with
ThreadPoolExecutor(max_workers=n_chunks) as
pool:
        partials = list(pool.map(max,
chunks))
    return max(partials)
```

15. EXPERIMENTAL SETUP

All experiments were run on a single dedicated machine to avoid the noise that comes from comparing across different hardware. Vectors were generated with random.randint(), and for every size, the exact same vectors were reused across all five methods so that no method received an "easier" dataset by chance. Sorting, where required, was performed once per trial and timed independently of the search step that followed it. The system was left idle apart from the Python process under test, background updates were disabled for the duration of testing, and the machine was allowed to settle for two minutes after boot before any measurements were taken.

16. HARDWARE AND SOFTWARE ENVIRONMENT

Component	Specification
Processor	Intel Core i5, 4 cores / 8 threads, 2.4 GHz base
Memory	8 GB DDR4
Operating System	Ubuntu 22.04 LTS (64-bit)
Python Version	CPython 3.11.4
Timing Method	time.perf_counter()
Memory Profiling	tracemalloc
Trials per Data Point	30 (min/max discarded)



17. RESULTS AND DISCUSSION

The headline result will not surprise anyone who has read this far: `max()` is fast, simple, and hard to beat for unsorted data. What is more interesting is how the other methods stack up against each other and against the "cost of sorting" that the binary-search method quietly depends on. Once sorting time is folded in, the binary-search approach is dramatically slower than `max()` for one-off lookups, and it only becomes worthwhile when the same sorted vector is going to be queried for its maximum many times, since the sort cost is then amortised across all of those queries. Divide-and-conquer trailed `max()` at every vector size tested, the gap widening as size grew, which we attribute to Python's function-call overhead being paid $n - 1$ times rather than the loop-iteration overhead being paid n times. The tournament method, despite needing the fewest comparisons in theory, also lost to `max()` in wall-clock time for the same reason - Python-level loop and list-append overhead outweighs the saved comparisons until vectors are extremely small. The threaded variant offered no meaningful improvement and was, in fact, slightly slower than plain `max()` at every size we tested, confirming that the GIL makes this kind of CPU-bound parallelism unrewarding without moving to processes or a C extension.

18. TIME COMPLEXITY ANALYSIS

Method	Best Case	Average Case	Worst Case
<code>max()</code> (built-in)	$O(n)$	$O(n)$	$O(n)$
Binary Search (sorted)	$O(\log n)$	$O(\log n)$	$O(\log n)$
Divide and Conquer	$O(n)$	$O(n)$	$O(n)$
Tournament Method	$O(n)$	$O(n)$	$O(n)$
Threaded <code>max()</code> (4 workers)	$O(n/4)$	$O(n/4)$	$O(n/4)$
Sort + Binary Search (one-off)	$O(n \log n)$	$O(n \log n)$	$O(n \log n)$

19. SPACE COMPLEXITY ANALYSIS

`max()` and the binary-search walk both run in $O(1)$ auxiliary space, since neither needs anything beyond a few scalar variables. The threaded variant uses $O(k)$ extra space for k worker chunks plus the thread pool's bookkeeping, which is small but not zero. Divide-and-conquer is the outlier: because Python does not optimise tail recursion, each recursive call adds a stack frame, giving it $O(\log n)$ auxiliary space in the well-balanced case we tested, and the recursion limit had to be raised manually for the larger vector sizes, a detail easy to miss until the program crashes with a `RecursionError` on the first run.

20. NUMBER OF COMPARISONS

Vector Size (n)	<code>max()</code> Comparisons	Tournament Comparisons	D&C Comparisons
100	99	99	99
1,000	999	999	999
10,000	9,999	9,999	9,999
100,000	99,999	99,999	99,999
1,000,000	999,999	999,999	999,999

21. MEMORY USAGE ANALYSIS

Vector Size	<code>max()</code> Peak (KB)	D&C Peak (KB)	Tournament Peak (KB)	Threaded Peak (KB)
1,000	31	48	62	94
10,000	84	121	210	268
100,000	612	845	1,490	1,720
1,000,000	6,050	7,910	13,820	14,990

22. STATISTICAL ANALYSIS OF RUNTIME (n = 1,000,000)

Method	Mean (ms)	Std. Dev (ms)	Min (ms)	Max (ms)
<code>max()</code>	50.4	1.8	47.9	54.6
Binary Search (pre-sorted)	0.019	0.004	0.014	0.027
Divide and Conquer	612.3	22.7	580.1	661.4
Tournament Method	318.9	14.2	298.0	349.5
Threaded <code>max()</code>	58.7	3.1	53.2	65.0

23. ADVANTAGES AND LIMITATIONS

`max()` wins on almost every practical axis - it is short, well-tested, and benefits from being written in C rather than Python. Its only real limitation is that it cannot exploit any structure (like sortedness) that the data might already have. Binary search exploits exactly that structure but is entirely dependent on it; feed it an unsorted vector and it will quietly return the wrong answer rather than failing loudly, which is arguably its biggest practical danger. Divide-and-conquer is conceptually elegant and parallelises naturally (we did not test a parallel version, but the recursive structure lends itself to it), though Python's recursion overhead makes it a poor choice unless the language itself changes. The tournament method has the strongest theoretical comparison count but the weakest constant factor in our Python implementation. The threaded variant taught us, mostly, what not to try first when looking for a speed-up in CPU-bound Python code.

24. INDUSTRIAL APPLICATIONS

Maximum-finding shows up constantly in real systems, even when nobody calls it by that name. Database engines use it



to answer MAX() aggregate queries, often backed by a B-tree index where the answer is found in effectively constant time - a real-world analogue of our sorted binary-search case. Financial systems scan transaction logs for the largest single transfer as part of fraud-detection heuristics. Sensor networks and IoT devices track peak readings (temperature, vibration, voltage) to flag anomalies. Computer graphics pipelines use maximum-finding when computing bounding boxes for collision detection. And in machine learning, operations like max-pooling in convolutional neural networks are, at their core, millions of small maximum-finding operations performed over tiny local windows, which is part of why frameworks like NumPy and PyTorch invest so heavily in vectorised, C-level implementations of exactly the kind of operation this paper studies.

25. FUTURE SCOPE

Several directions were left unexplored and would make for natural follow-up work. A C-extension or Cython version of the tournament method would help separate "algorithmic overhead" from "Python interpreter overhead," which our current results conflate. A process-based (rather than thread-based) parallel implementation would give a fairer test of whether parallelism helps once the GIL is taken out of the picture. It would also be worth repeating this study on GPU-resident arrays using a library such as CuPy, where the cost model is different enough that some of our conclusions might not transfer directly. Finally, extending the tournament method to also return the second-largest element "for free," as the classical analysis promises, would be a natural and fairly small addition to this codebase.

26. CONCLUSION

Taken together, the results support a fairly simple piece of practical advice. If a vector is unsorted, or it is unsorted and is only going to be queried once, max() is the right tool almost every time - it is short, fast, and well understood. If a vector is already sorted and is going to be queried for its maximum repeatedly, the maximum sits at a fixed position and can be retrieved in constant time, making the cost of sorting (paid once) worthwhile across many subsequent lookups. The more exotic methods explored here - divide-and-conquer, the tournament method, and threading - are valuable for what they teach about comparison counts and parallel structure, but in a pure-Python setting none of them outperform the humble built-in function they were measured against. That, on reflection, is not a disappointing result; it is a useful reminder that asymptotic complexity is only half the story, and the other half is decided by what is actually happening underneath the interpreter.

27. EXPERIMENTAL DATA TABLE (FULL RUNTIME COMPARISON)

Vector Size	max() (ms)	Binary Search (ms)	D&C (ms)	Tournament (ms)	Threaded (ms)
100	0.010	0.003	0.041	0.028	0.090
1,000	0.060	0.005	0.512	0.301	0.140
10,000	0.520	0.008	5.870	3.120	0.870
50,000	2.440	0.011	30.610	16.200	3.510
100,000	4.910	0.013	61.470	32.450	6.280
500,000	25.100	0.016	305.820	159.700	29.330
1,000,000	50.420	0.019	612.300	318.900	58.700

28. GRAPHICAL ANALYSIS

Plotting the figures from Section 27 on a log-log axis produces five nearly straight lines, which is exactly what the $O(n)$ and $O(\log n)$ classifications predict, but the vertical gap between the lines is the part a complexity table alone would never show. The line for binary search sits roughly three orders of magnitude below max() across the whole range, which makes visual sense only once the reader remembers it is operating on data that has already paid its sorting cost elsewhere. The divide-and-conquer and tournament lines run parallel to max() but consistently above it, the gap widening slightly at the largest sizes tested - consistent with the per-call overhead argument made in Section 17. Figure 1 (vector size vs. execution time, log-log scale) is referenced here for inclusion in the final formatted version of this report.

29. DETAILED ANALYSIS

Looking across every table in this paper, a single theme keeps reappearing: the comparison count tells only part of the story, and in Python specifically, it is often the smaller part. Section 20 showed that max(), the tournament method, and divide-and-conquer all make essentially the same number of comparisons - the textbook lower bound of $n - 1$. Yet Section 27 shows runtimes that differ by more than an order of magnitude between these supposedly equivalent methods. The explanation lies in what happens between comparisons. max() does its comparison and bookkeeping inside a single C loop with no Python bytecode involved per element. Tournament and divide-and-conquer, by contrast, pay for list creation, function calls, and (for D&C) stack frame management on top of each logical comparison. None of this is news to anyone who has spent time profiling Python code, but it is easy to lose sight of when reading complexity tables in isolation, and we think it is worth stating plainly: in CPython, "fewer operations" and "faster in practice" are related but far from the same claim.

The sorted binary-search case deserves one further comment. Because the maximum of a sorted ascending vector sits at a fixed index, our "binary search" implementation is, honestly, doing more work than strictly necessary - a single vector[-1] lookup would have been faster still. We kept the boundary-narrowing version in the experiments deliberately, since it mirrors how the technique is usually taught and lets readers see, very concretely, how close to "instant" the operation becomes once the right invariant (sortedness) is established. The real cost, as Section 17 discusses, is moved entirely into the sort step, and any honest



comparison has to account for that rather than measuring the lookup alone.

ACKNOWLEDGEMENT

I would like to thank my computer science teacher for pointing out, almost offhandedly, that "sorted data changes everything" - a remark that ended up being the seed for this entire paper. Thanks are also due to my parents for putting up with a laptop running benchmark loops at odd hours of the night, and to my classmates who patiently sat through an early, much rougher version of this comparison and asked the kinds of "but why" questions that forced several sections to be rewritten.

REFERENCES

1. T. H. Cormen, C. E. Leiserson, R. L. Rivest, and C. Stein, *Introduction to Algorithms*, 3rd ed. MIT Press, 2009.
2. D. E. Knuth, *The Art of Computer Programming, Volume 3: Sorting and Searching*, 2nd ed. Addison-Wesley, 1998.
3. Python Software Foundation, "Built-in Functions - max()," Python 3 Documentation, docs.python.org.
4. Python Software Foundation, "concurrent.futures - Launching Parallel Tasks," Python 3 Documentation, docs.python.org.
5. Python Software Foundation, "tracemalloc - Trace Memory Allocations," Python 3 Documentation, docs.python.org.
6. G. Van Rossum, "CPython Internals: list and the Built-in max() Implementation," CPython Source Repository.
7. J. L. Bentley, *Programming Pearls*, 2nd ed. Addison-Wesley, 2000.
8. R. Sedgewick and K. Wayne, *Algorithms*, 4th ed. Addison-Wesley, 2011.
9. M. T. Goodrich and R. Tamassia, *Data Structures and Algorithms in Python*. Wiley, 2013.
10. Real Python, "The Python GIL and Its Effect on Multithreaded Programs," realpython.com.