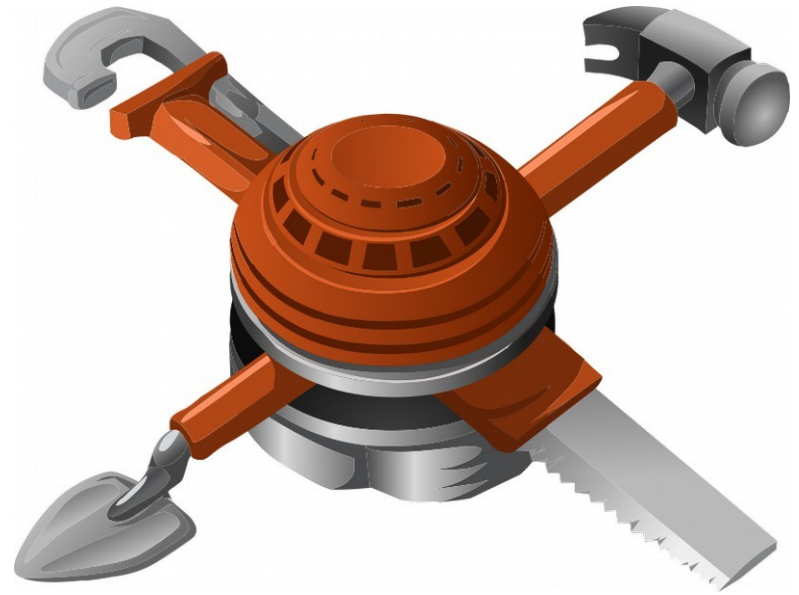


# CS 470 Spring 2026

Mike Lam, Professor



## Performance Tools

# Software Tools

- **Software tool**: computer program used by developers to create, debug, maintain or support other programs

# Traditional Software Tools

- Text editors
- Version control
- Debuggers
- Profilers
- Test automation frameworks
- Deployment frameworks
- Integrated development environments (IDEs)

# Traditional Software Tools

- **Debuggers**
  - Purpose: finding and removing software defects
  - Often done via a process monitoring interface
- **Profilers**
  - Purpose: detecting performance characteristics and identifying **bottlenecks**
  - Often done via instrumentation (added code that tracks the program's execution)
- Both of these are difficult in parallel and distributed systems

# Traditional Debugging

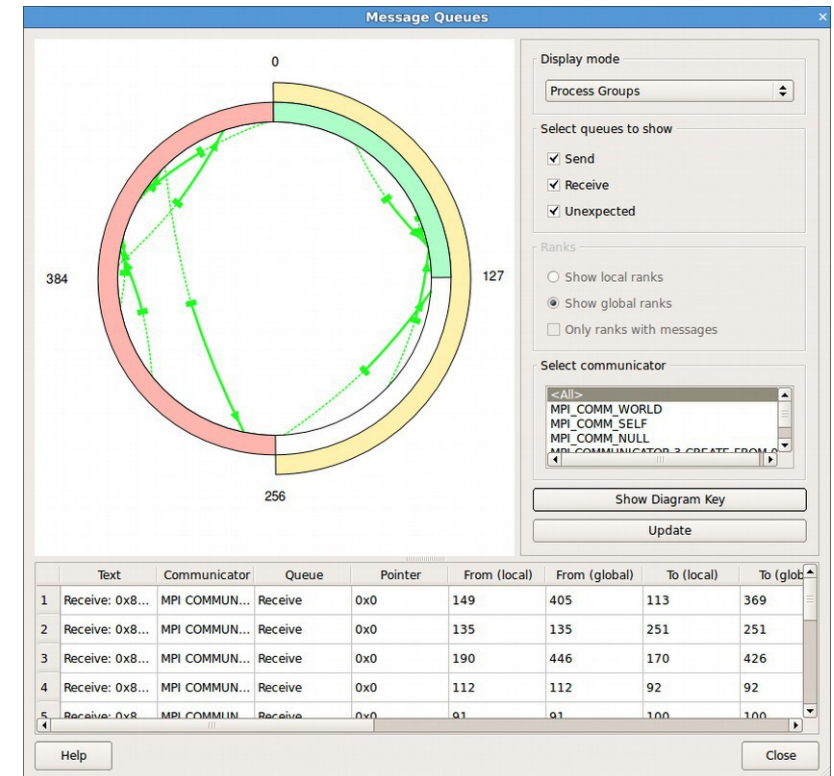
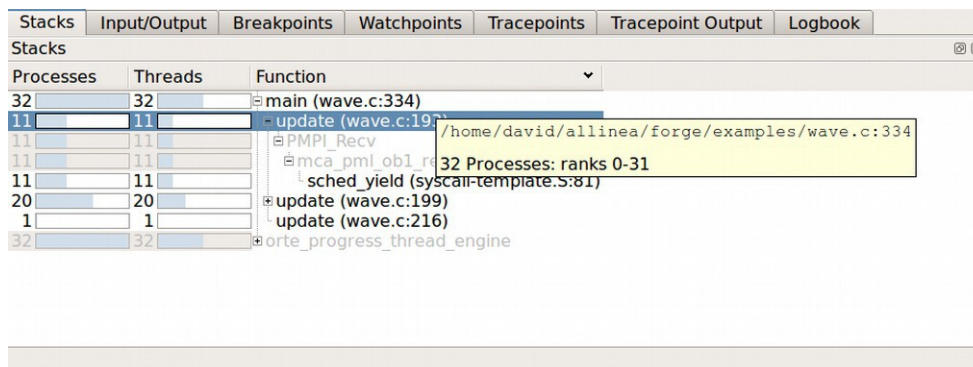
- Mechanisms
  - `ptrace`: system call that allows one process to control another
  - Simulation: slower, but safer
- Common features
  - Breakpoints and watchpoints
  - Single-stepping (by instruction or line of code)
  - Variable examination and modification
  - In newer debuggers: reverse-stepping
- Free debuggers: [gdb](#), [lldb](#), [Eclipse](#), [Valgrind \(Memcheck\)](#)

# Parallel Debugging

- Multithreaded debugging can be difficult
  - Must attach to the correct thread
  - Must control other threads as well
  - Nondeterminism means unpredictability
  - GDB does include support for multithreading:
    - <http://sourceware.org/gdb/current/onlinedocs/gdb/Threads.html>
  - Valgrind also provides the **Helgrind** error and race detector
- Distributed debugging is even harder
  - Hundreds or thousands of nodes; millions of processes
  - Enormous launch overhead
  - Control and visualization issues
  - **tmpi**: OSS tmux hack (<https://github.com/Azrael3000/tmpi>)

# Commercial debuggers

- Microsoft Visual Studio
- Intel Debugger
- Rational Purify
- RogueWave TotalView
- Allinea DDT



# Performance Profiling

- Goal: gain insights concerning a program's performance characteristics
- Common **metrics**
  - Wall or CPU time
  - Memory use, page faults, and cache misses
  - Network traffic and saturation
  - Power and energy use
- Common **scopes**
  - Function
  - Basic block
  - Instruction
  - Source code line



# Measurement

- **Instrumentation**: inserting analysis code
  - Binary vs. source
  - Static vs. dynamic
  - Best for event-based monitoring (e.g., function calls)
- **Sampling**: polling an analysis source
  - Hardware counters
    - Performance Application Programming Interface (**PAPI**)
  - Randomized vs. periodic
  - Averaging vs. min/max
  - Best for continuous monitoring (e.g., memory usage)

# Measurement

- **Context**
  - Flat vs. call graph
  - Partial vs. full context
- **Profiling** vs. **tracing** (latter builds time-series)
- **Issues**
  - **Overhead**: added run time due to profiling software
  - **Perturbation**: skewing of behavior due to profiling software
  - **Skid**: execution may not stop immediately on sample
- **Tradeoff: better information vs. lower overhead**
  - “Better” could mean more accurate, more detailed, more granular, etc.
  - Instrumentation: more instrumentation points
  - Sampling: higher frequency or less aggregation

# GNU Profiler (gprof)

- Compile with “-pg” flag
- Disable optimization with “-O0” flag
- Run as usual; generates “gmon.out” file
- View results with “gprof” utility
  - “gprof <executable>”
- See <https://sourceware.org/binutils/docs/gprof/> for more documentation
- Google also has a multi-threaded profiler:
  - <https://github.com/gperftools/gperftools>

# Sample gprof output

Each sample counts as 0.01 seconds.

| %<br>time | cumulative<br>seconds | self<br>seconds | calls | self<br>s/call | total<br>s/call | name                  |
|-----------|-----------------------|-----------------|-------|----------------|-----------------|-----------------------|
| 99.47     | 1.20                  | 1.20            | 1     | 1.20           | 1.20            | gaussian_elimination  |
| 0.83      | 1.21                  | 0.01            | 1     | 0.01           | 0.01            | rand_system           |
| 0.00      | 1.21                  | 0.00            | 1     | 0.00           | 0.00            | back_substitution_row |
| 0.00      | 1.21                  | 0.00            | 1     | 0.00           | 0.00            | find_max_error        |

# Callgrind/Cachegrind

- Run with Valgrind
  - Callgrind: `valgrind --tool=callgrind <executable>`
  - Cachegrind: `valgrind --tool=cachegrind <executable>`
    - `--cache-sim=yes` or `--branch-sim=yes` to enable cache/CPU simulations
  - This will produce a `*.out.xxxx` file with raw results (could be large!)
  - Remember to call `mpirun` first if it's an MPI program
    - (And use `cg_merge` to merge multiple Cachegrind output files)
- Post-process results
  - Callgrind: `callgrind_annotate <output-files>`
    - GUI alternative: `kcachegrind` (or `qcachegrind` on Mac OS X)
  - Cachegrind: `cg_annotate <output-file>` (`--auto=yes` for code)
    - Dx = data cache (level X)      Ix = instruction cache (level X)
    - 1 = L1 cache      L/LL = lowest level (on the cluster, this is L3)
    - r = read      w = write      m = miss      Ir = Instructions read
- See <http://valgrind.org/docs/manual> for more documentation

# Sample callgrind output

```
.          void gaussian_elimination()
5 ( 0.00%) {
    4,005 ( 0.00%)   for (int pivot = 0; pivot < n; pivot++) {
    2,005,000 ( 0.02%)       for (int row = pivot+1; row < n; row++) {
    10,989,000 ( 0.09%)           REAL coeff = A[row*n + pivot] / A[pivot*n + pivot];
    5,494,500 ( 0.05%)           A[row*n + pivot] = 0.0;
    1,334,830,500 (11.04%)       for (int col = pivot+1; col < n; col++) {
    10,650,672,000 (88.06%)           A[row*n + col] -= A[pivot*n + col] * coeff;
    .                               }
    9,990,000 ( 0.08%)           b[row] -= b[pivot] * coeff;
    .                               }
    .                               }
3 ( 0.00%) }
```

# Sample cachegrind output

```
==4077641== I refs:          12,095,390,722
==4077641== I1  misses:          1,655
==4077641== LLi misses:          1,580
==4077641== I1  miss rate:         0.00%
==4077641== LLi miss rate:         0.00%
==4077641==
==4077641== D refs:          6,377,487,537 (6,038,108,427 rd + 339,379,110 wr)
==4077641== D1  misses:          42,700,623 ( 42,574,295 rd + 126,328 wr)
==4077641== LLd misses:          128,126 (      2,148 rd + 125,978 wr)
==4077641== D1  miss rate:         0.7% (      0.7% + 0.0% )
==4077641== LLd miss rate:         0.0% (      0.0% + 0.0% )
==4077641==
==4077641== LL refs:          42,702,278 ( 42,575,950 rd + 126,328 wr)
==4077641== LL misses:          129,706 (      3,728 rd + 125,978 wr)
==4077641== LL miss rate:         0.0% (      0.0% + 0.0% )
```

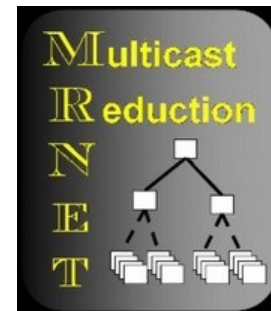
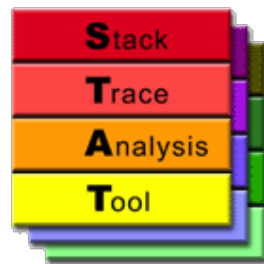
# Perf\_events

- Sample-based performance profiler
  - Kernel module reads performance counters
    - More lightweight than Valgrind-based analysis
    - Can sample many different events
  - User space utility `perf` to interface with kernel
    - `perf record -F 49 <command>`
      - Generates `perf.data` file
    - `perf report -n [--stdio]`
    - `perf annotate [--stdio]`
  - Cheat sheet link on resources page



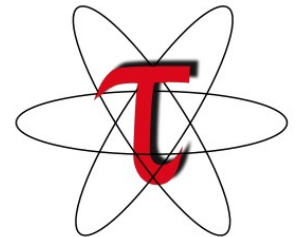
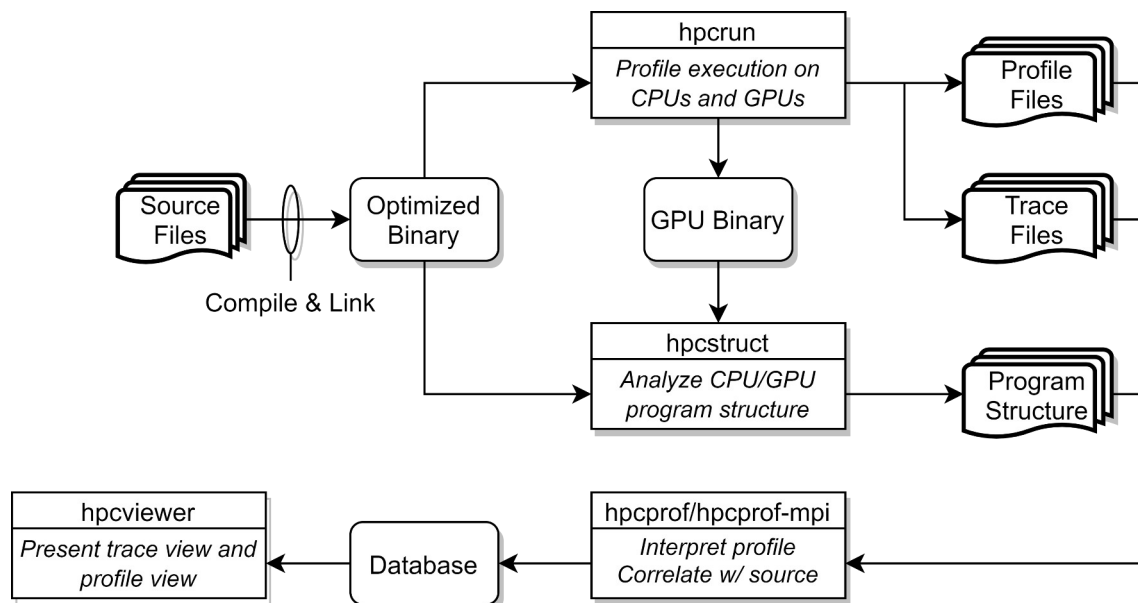
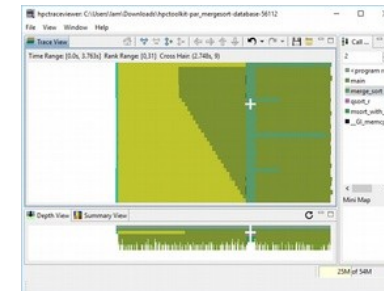
# Distributed Analysis

- Lots of data!
  - Collect at each rank but only store compressed or aggregated data
  - Aggregate using a tree-based reduction structure to reduce communication overhead
    - Research projects: [STAT](#) and [MRNet](#)

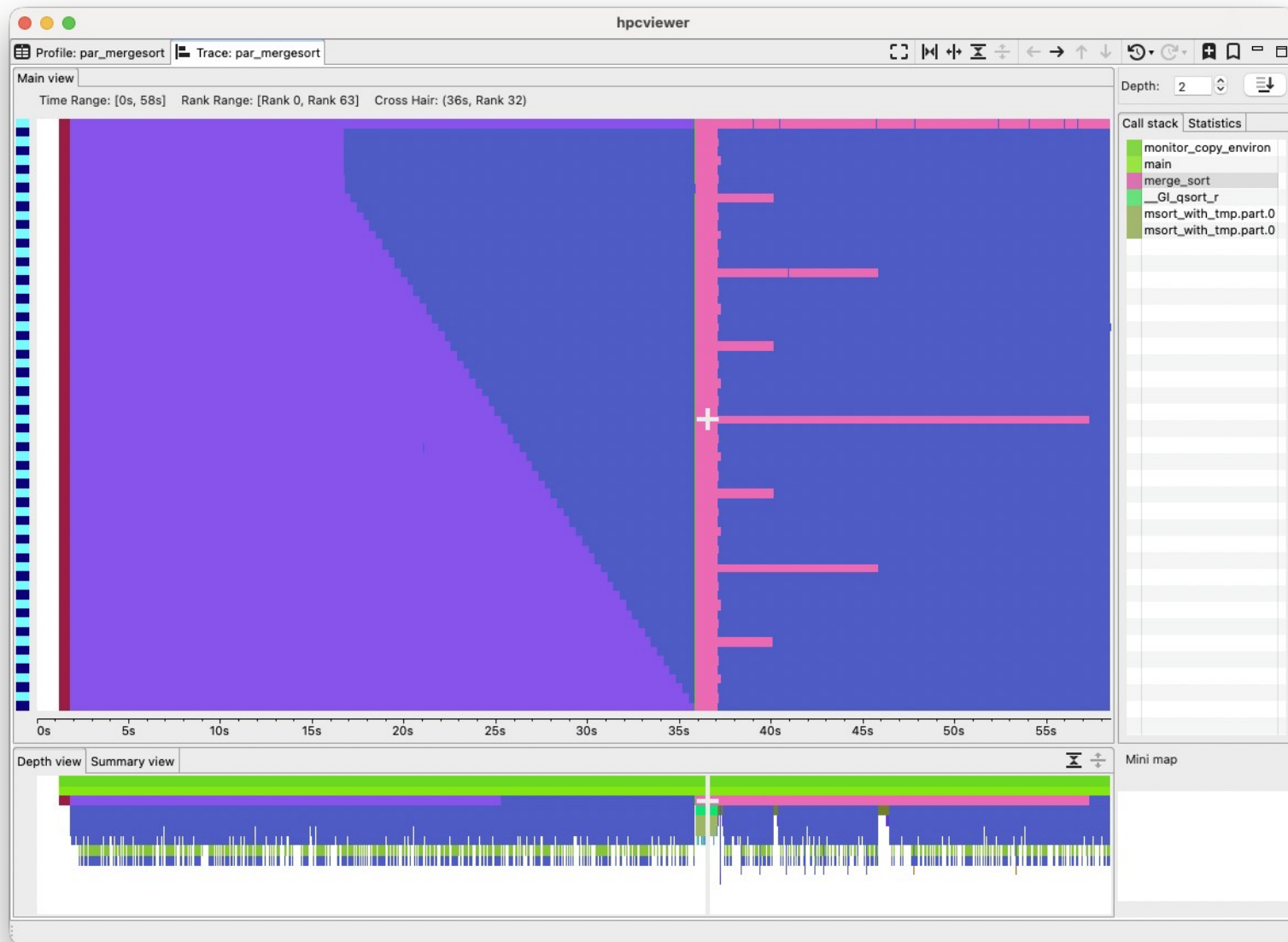


# Other HPC analysis tools

- [HPCToolkit](#) – Rice University
- [Tuning and Analysis Utilities \(TAU\)](#) – University of Oregon
- [Open|SpeedShop](#) - Krell Institute
- Scalasca
- Paraver



# Sample HPCToolkit results



# Tool frameworks

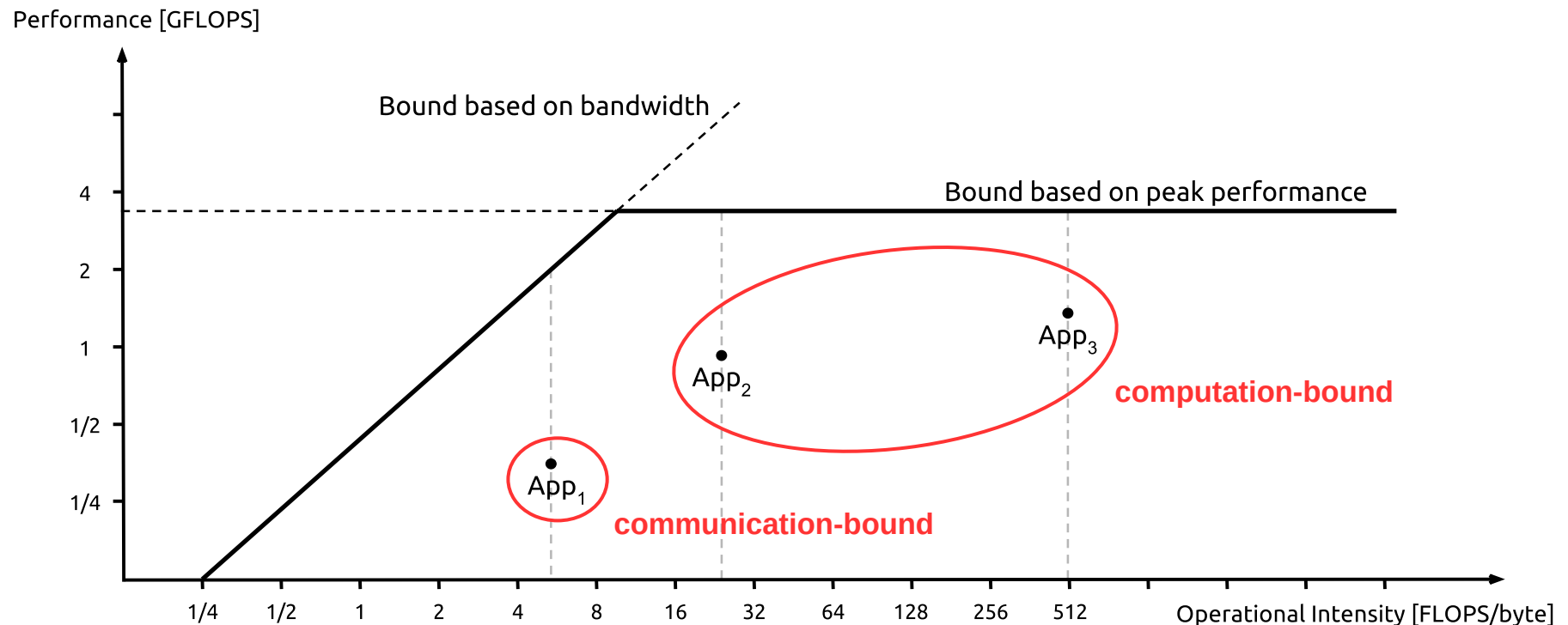
- Many analysis tools need similar functionality
  - E.g., source/binary parsing or instrumentation
- **Tool framework**: a library that provides common functionality upon which custom tools can be written
  - **Rose** (source-based compiler framework)
  - **LLVM** (binary-based compiler framework)
  - **Intel Pin** (insert just-in-time binary instrumentation)
  - **Dyninst** (insert binary instrumentation)
  - **Valgrind** (track memory accesses)
  - **CRAFT** (instrument floating-point arithmetic)

# Modeling and autotuning

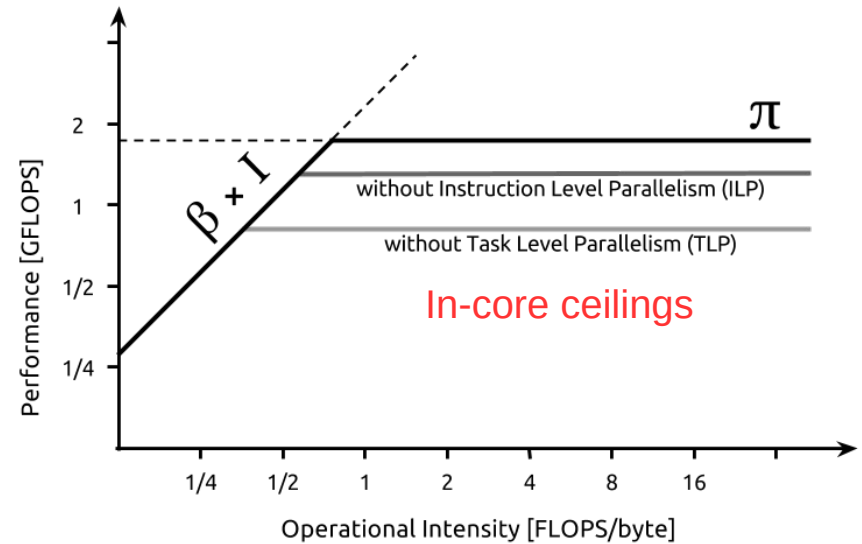
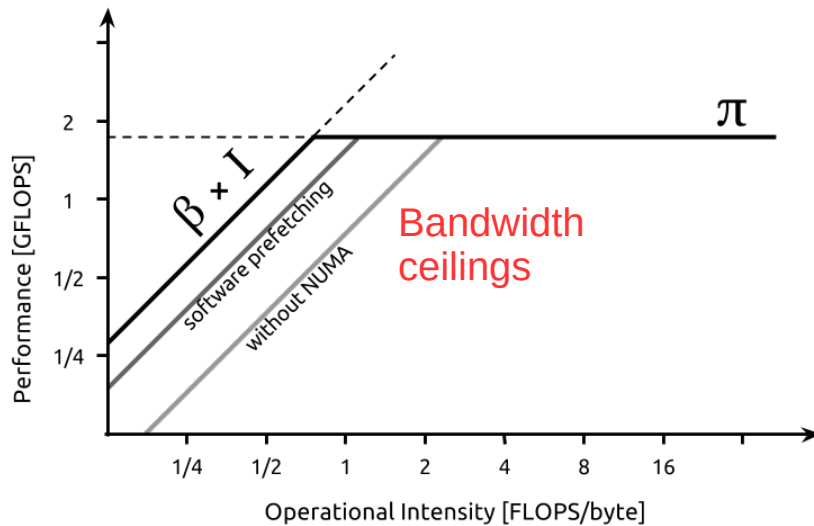
- Observation: modern systems have a lot of knobs
  - Message size, block size, # of threads, # of processes
  - Many of these factors influence each other
  - Different runs could require different “optimal” settings
- Idea #1: **autotune** the system at runtime
  - Managed by middleware (the autotuner)
  - Overhead could be expensive
  - Optimizing across multiple dimensions can be difficult
- Idea #2: build a **model** of these interactions
  - Requires training data
  - Leverages the significant recent advances in AI and ML

# Performance models

- One popular model: **roofline** model
  - Shows theoretical limits on performance
  - Based on computation and communication bounds



# Performance models



$$P = \min \left\{ \begin{array}{l} \pi \\ \beta \times I \end{array} \right.$$

$\pi$  = peak performance

$\beta$  = peak bandwidth

$I$  = arithmetic intensity

$P$  = attainable performance

