

Iteratively Improving Spark Application Performance

William C. Benton
Red Hat, Inc.

Forecast

- **Background: Spark, RDDs, and Spark's execution model**
- **Case study overview**
- **Improving our prototype**

Background

Apache Spark

- **Introduced in 2009; donated to Apache in 2013; 1.0 release in 2014**
- **Based on a fundamental abstraction: the resilient distributed dataset**
- **Supports in-memory computing and a wide range of algorithms**

Spark is general

Spark core

Spark is general

Graph

Spark core

Spark is general

Graph

SQL

Spark core

Spark is general

Graph

SQL

ML

Spark core

Spark is general

Graph

SQL

ML

Streaming

Spark core

Spark is general

Graph

SQL

ML

Streaming

Spark core

ad hoc

Mesos

YARN

Spark is general

*APIS FOR SCALA, JAVA, PYTHON, AND R
(3RD-PARTY BINDINGS FOR CLOJURE ET AL.)*

Graph

SQL

ML

Streaming

Spark core

ad hoc

Mesos

YARN



Resilient distributed datasets

[illegible]



Resilient distributed datasets

[illegible]



Resilient distributed datasets

[illegible]

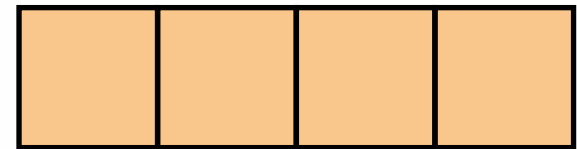
Resilient distributed datasets



Resilient distributed datasets



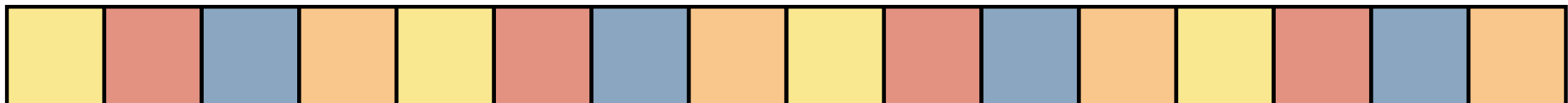
PARTITIONED ACROSS MACHINES BY RANGE...



Resilient distributed datasets



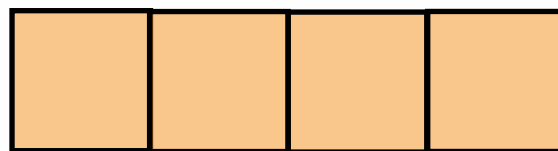
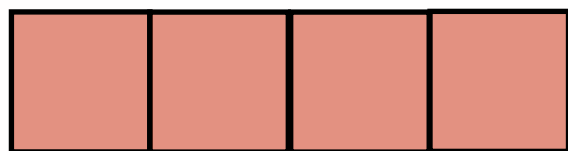
PARTITIONED ACROSS MACHINES BY RANGE...



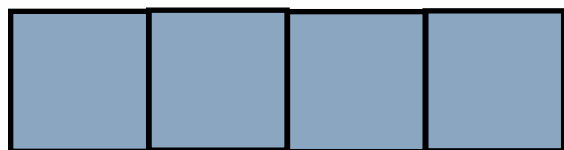
Resilient distributed datasets



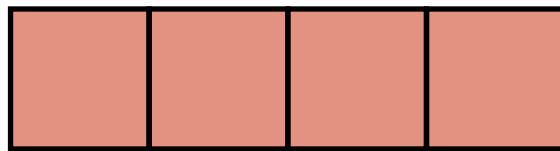
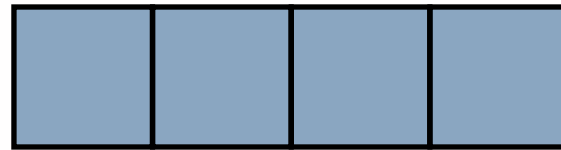
PARTITIONED ACROSS MACHINES BY RANGE...



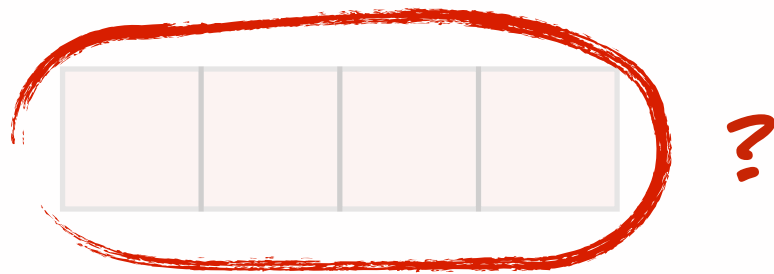
...OR BY HASH



Resilient distributed datasets

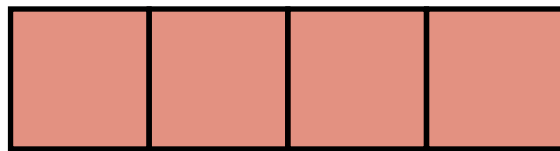
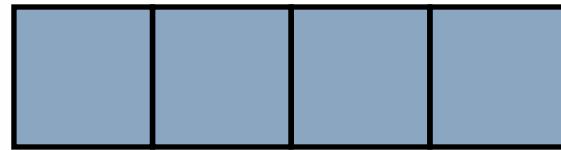


Resilient distributed datasets



FAILURES MEAN PARTITIONS CAN *DISAPPEAR...*

Resilient distributed datasets



FAILURES MEAN PARTITIONS CAN *DISAPPEAR...*
...BUT THEY CAN BE RECONSTRUCTED!

**RDDs are partitioned,
immutable, lazy collections**

RDDs are partitioned, **immutable, lazy** collections

*TRANSFORMATIONS CREATE NEW RDDS
THAT ENCODE A DEPENDENCY DAG*

*ACTIONS RESULT IN EXECUTING CLUSTER
JOBS & RETURN VALUES TO THE DRIVER*

Creating RDDs

- **From a collection:** `parallelize()`
- **...a local or remote file:** `textFile()`
- **...or HDFS:** `hadoopFile();`
`sequenceFile(); objectFile()`
- **(These all act lazily)**

RDD[T] transformations

- `map(f: T=>U): RDD[U]`
- `flatMap(f: T=>Seq[U]): RDD[U]`
- `filter(f: T=>Boolean): RDD[T]`
- `distinct(): RDD[T]`
- `keyBy(f: T=>K): RDD[(K, T)]`

RDD[(K,V)] transformations

- `sortByKey(): RDD[(K,V)]`
- `groupByKey(): RDD[(K, Seq[V])]`
- `reduceByKey(f: (V,V)=>V): RDD[(K,V)]`
- `join(other: RDD[(K,W)]): RDD[(K, (V,W))]`

RDD[(K,V)] transformations

- `sortByKey(): RDD[(K,V)]`
- `groupByKey(): RDD[(K, Seq[V])]`
- `reduceByKey(f: (V,V)=>V): RDD[(K,V)]`
- `join(other: RDD[(K,W)]): RDD[(K, (V,W))]`

---AND MANY OTHERS!

RDD[T] actions

- `collect(): Array[T]`
- `count(): Long`
- `reduce(f: (T,T)=>T): T`
- `saveAsTextFile(path)`
- `saveAsSequenceFile(path)`

RDD[T] actions

*(REMEMBER: ALL
ACTIONS ARE EAGER)*

- `collect(): Array[T]`
- `count(): Long`
- `reduce(f: (T,T)=>T): T`
- `saveAsTextFile(path)`
- `saveAsSequenceFile(path)`

---AND MANY OTHERS!

Example: word count in Spark

```
val file = spark.textFile("hdfs://...")

val counts = file.flatMap(line =>
    line.split(" "))
    .map(word => (word, 1))
    .reduceByKey(_ + _)

counts.saveAsTextFile("hdfs://...")
```

Example: word count in Spark

```
val file = spark.textFile("hdfs://...")

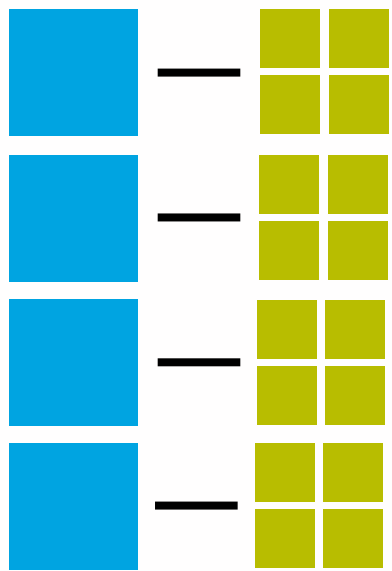
val counts = file.flatMap(line =>
    line.split(" "))
    .map(word => (word, 1))
    .reduceByKey(_ + _)

counts.saveAsTextFile("hdfs://...")
```

Example: word count in Spark

```
val file = spark.textFile("hdfs://...")
```

```
val counts = file.flatMap(line =>  
    line.split(" "))  
    .map(word => (word, 1))  
    .reduceByKey(_ + _)
```

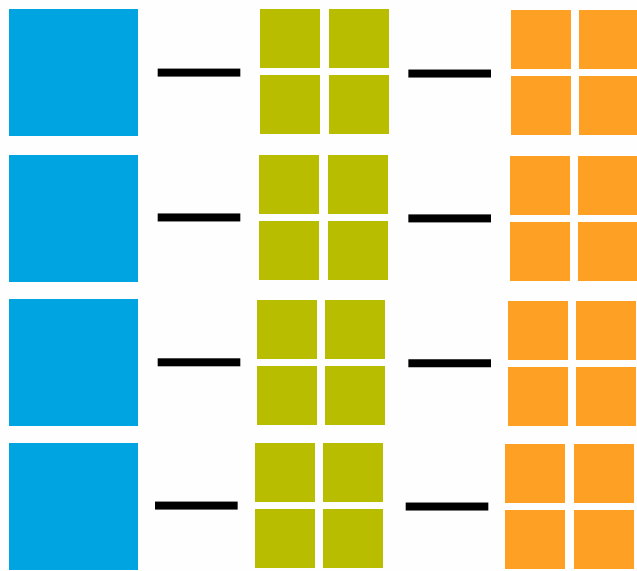


```
counts.saveAsTextFile("hdfs://...")
```

Example: word count in Spark

```
val file = spark.textFile("hdfs://...")
```

```
val counts = file.flatMap(line =>  
    line.split(" "))  
    .map(word => (word, 1))  
    .reduceByKey(_ + _)
```



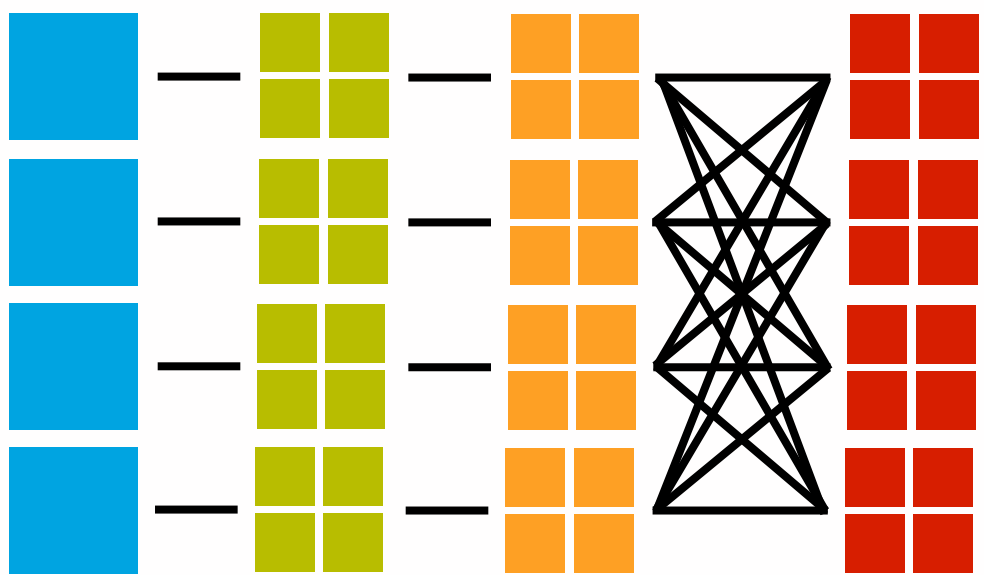
```
counts.saveAsTextFile("hdfs://...")
```

Example: word count in Spark

```
val file = spark.textFile("hdfs://...")

val counts = file.flatMap(line =>
    line.split(" "))
    .map(word => (word, 1))
    .reduceByKey(_ + _)

counts.saveAsTextFile("hdfs://...")
```

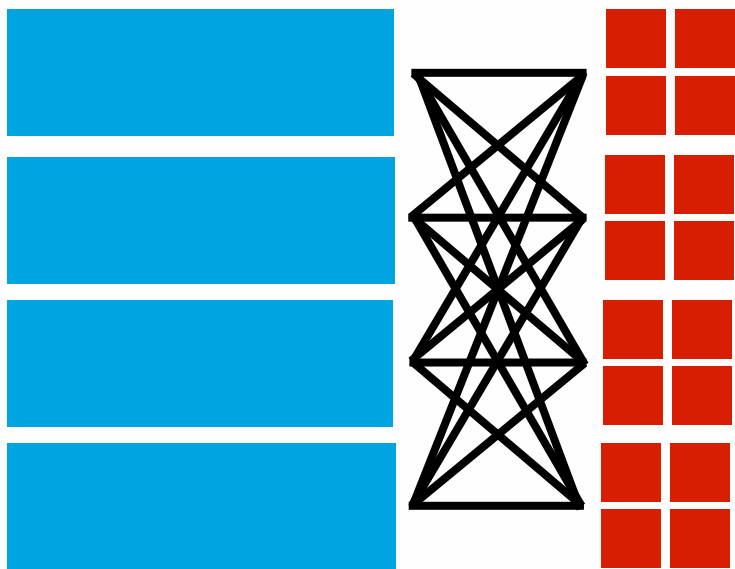


Example: word count in Spark

```
val file = spark.textFile("hdfs://...")

val counts = file.flatMap(line =>
    line.split(" "))
    .map(word => (word, 1))
    .reduceByKey(_ + _)

counts.saveAsTextFile("hdfs://...")
```



Case study: bicycling data

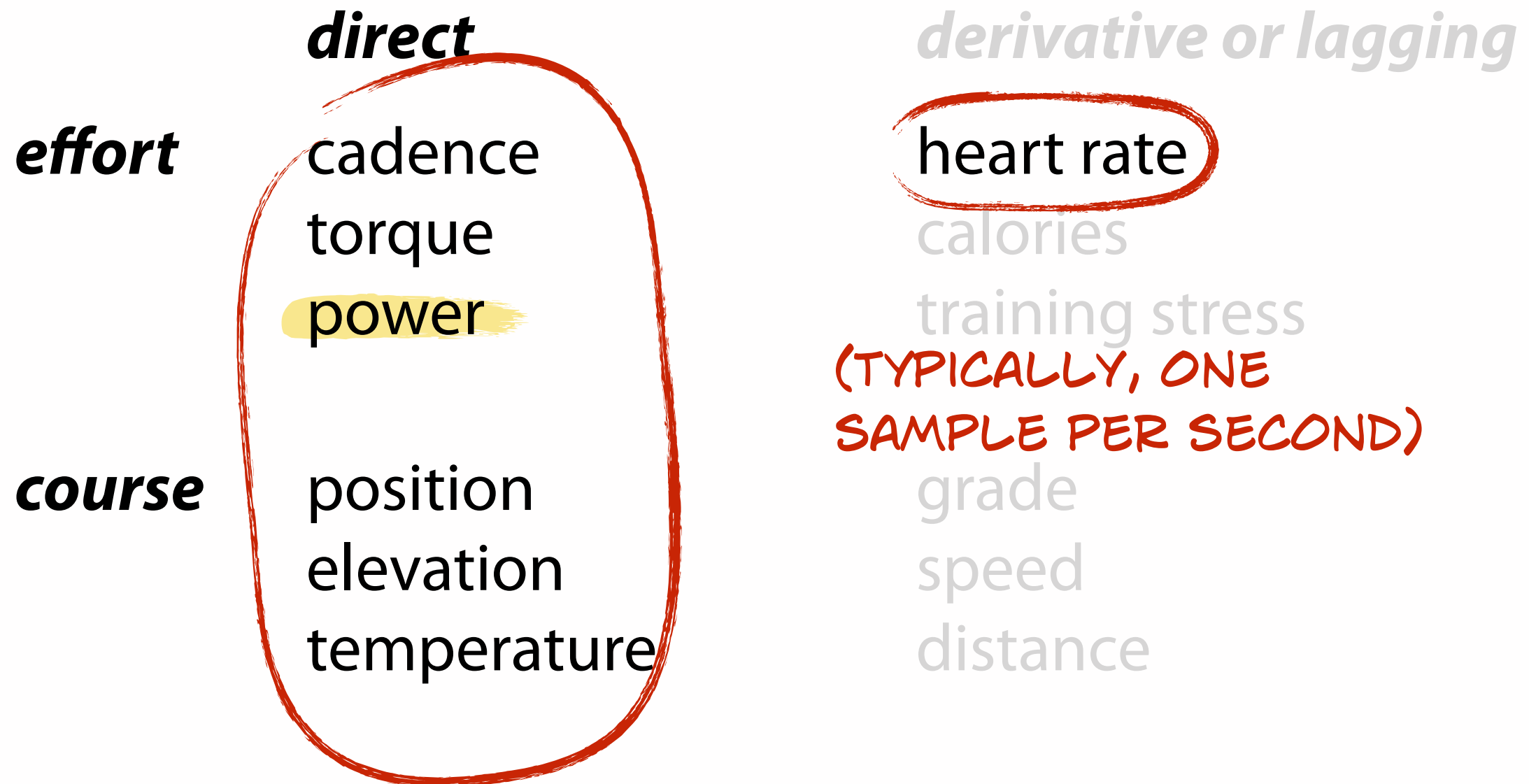
Metrics available to cyclists

	<i>direct</i>	<i>derivative or lagging</i>
<i>effort</i>	cadence torque power	heart rate calories training stress
<i>course</i>	position elevation temperature	grade speed distance

Metrics available to cyclists

	<i>direct</i>	<i>derivative or lagging</i>
<i>effort</i>	cadence torque power	<u>heart rate</u> calories training stress
<i>course</i>	position elevation temperature	(TYPICALLY, ONE SAMPLE PER SECOND) grade speed distance

Metrics available to cyclists



***“Where should I
do intervals?”***

MT. DIABLO (FROM NORTH GATE) 3,970'
NEAR WALNUT CREEK, CA 10.5 MILES

405'

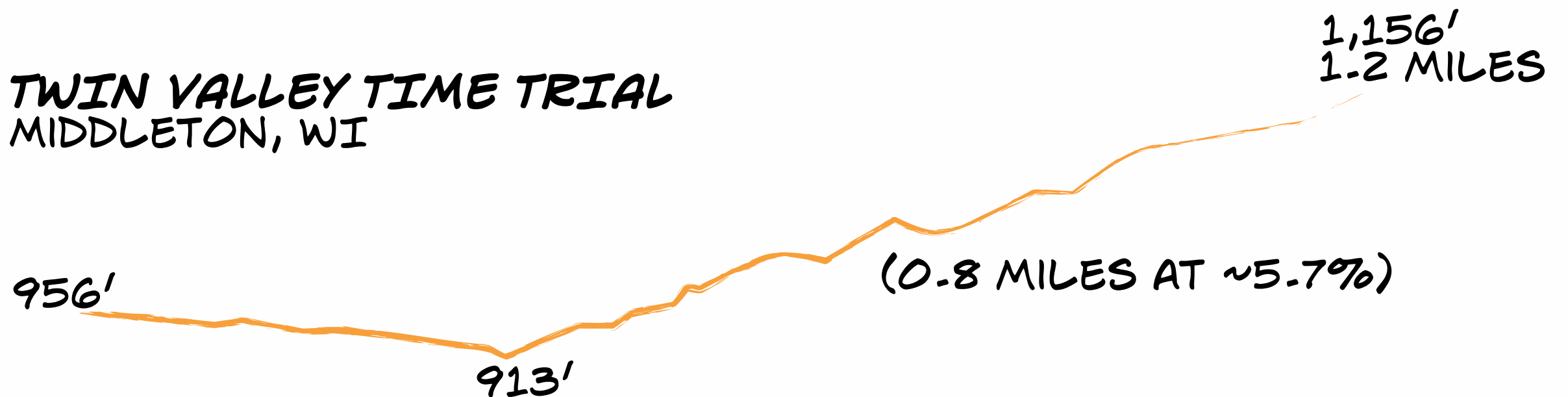
"BUDDHA'S PALM"
NEAR CROSS PLAINS, WI 1,214'

890'

16.8 MILE

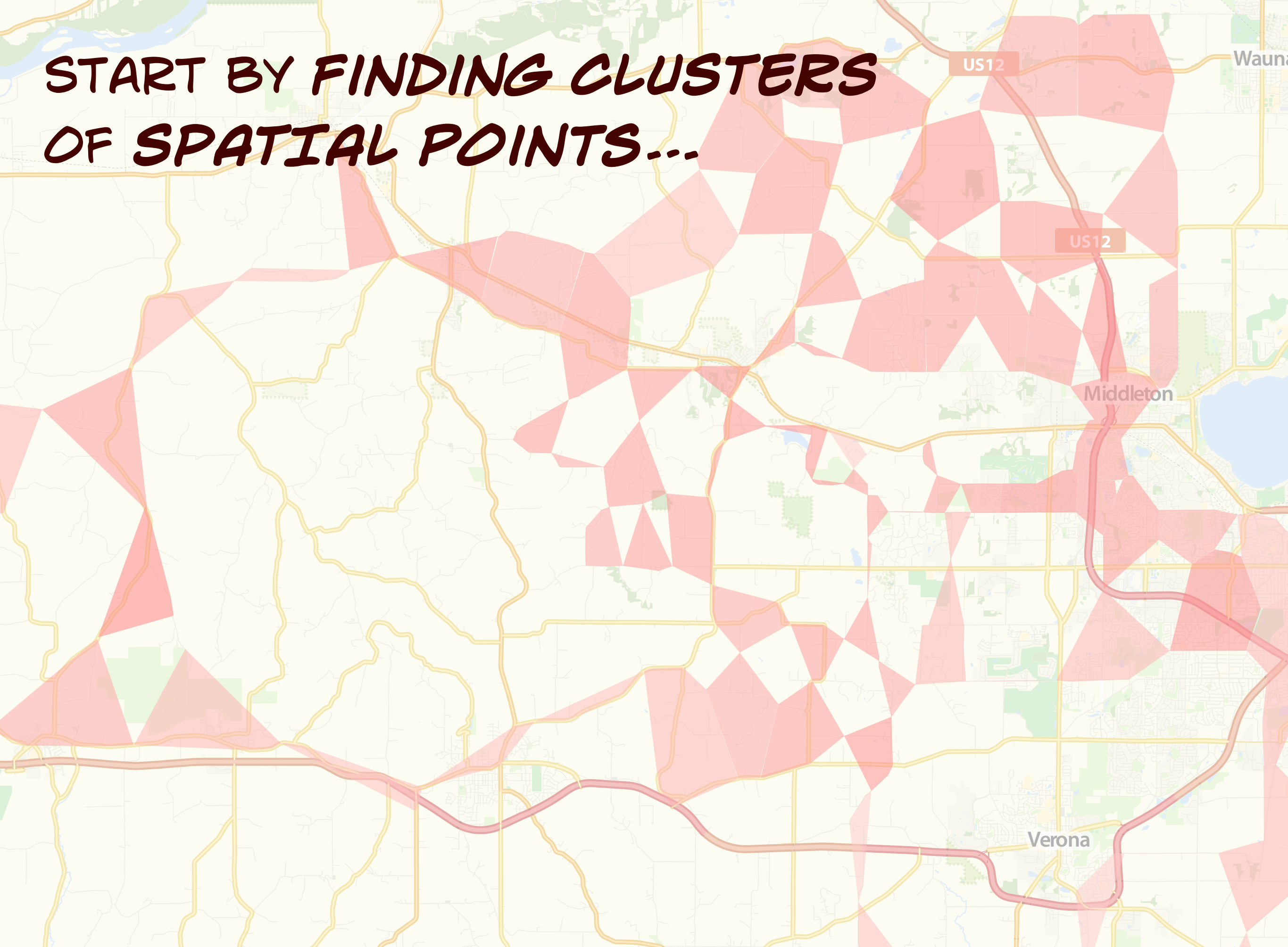
Finding best efforts

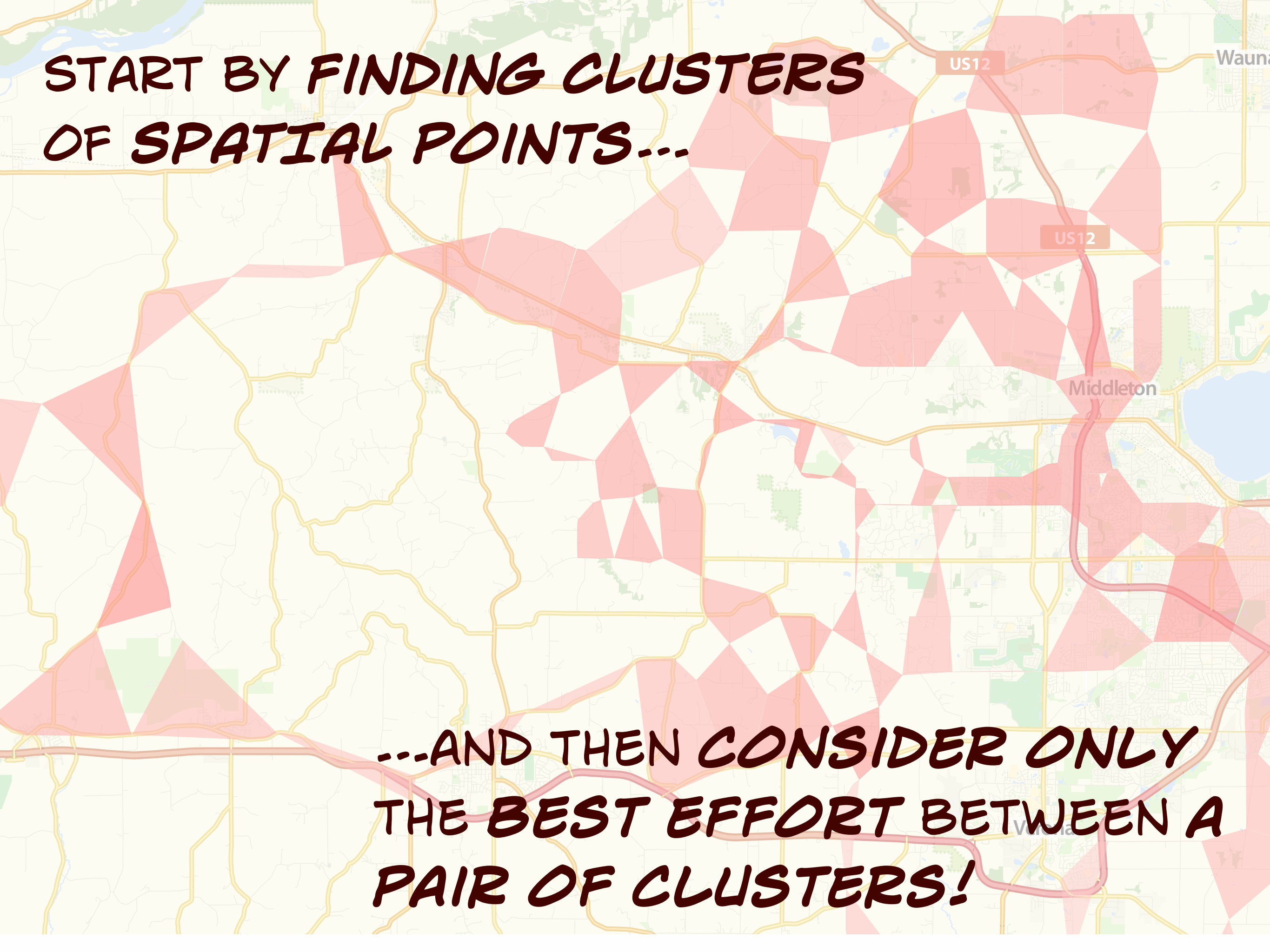
*TWIN VALLEY TIME TRIAL
MIDDLETON, WI*



**A FOUR-MINUTE ALL-OUT EFFORT INCLUDES
SIXTY STRONG THREE-MINUTE EFFORTS.**

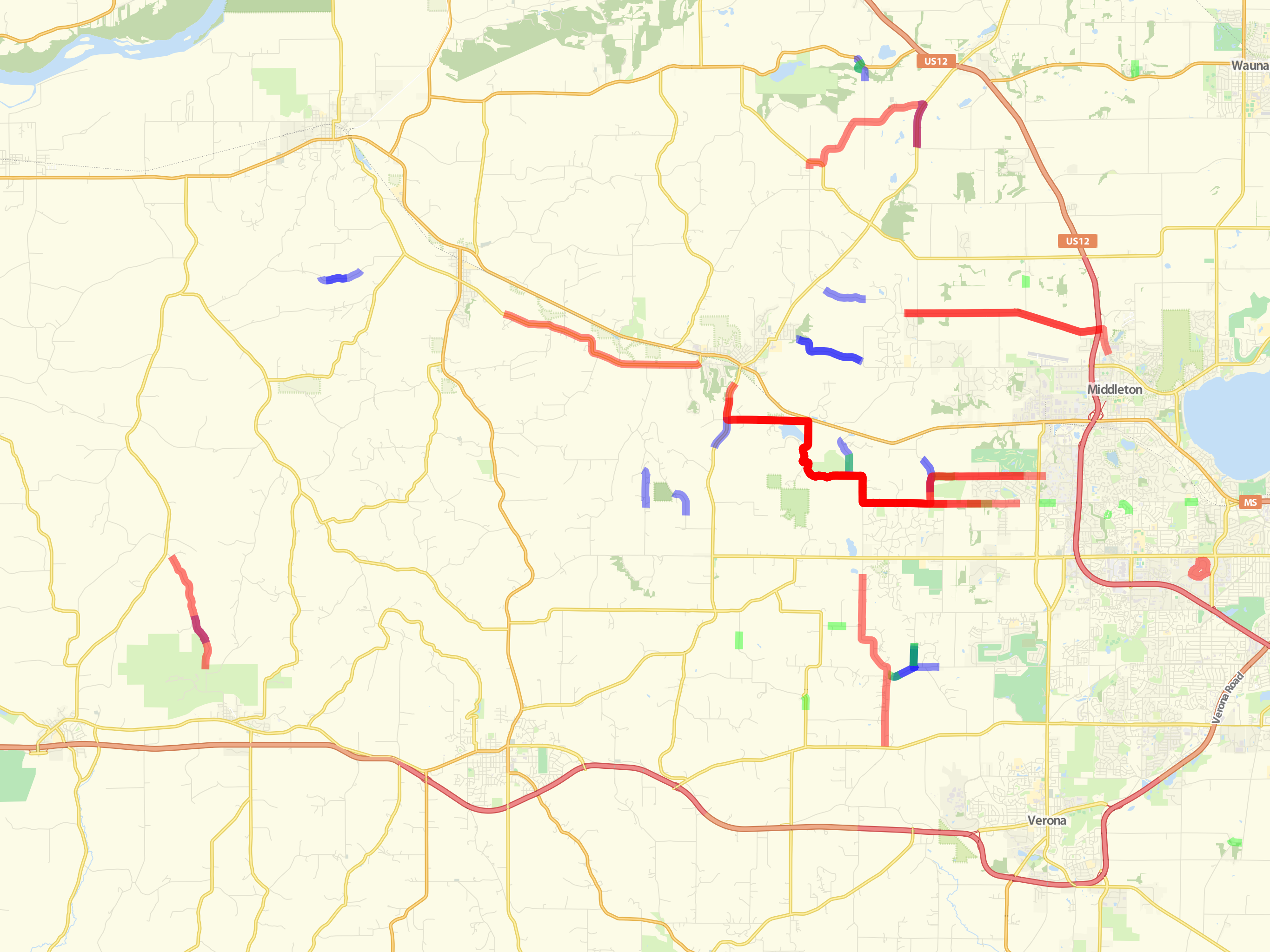
**START BY FINDING CLUSTERS
OF SPATIAL POINTS...**

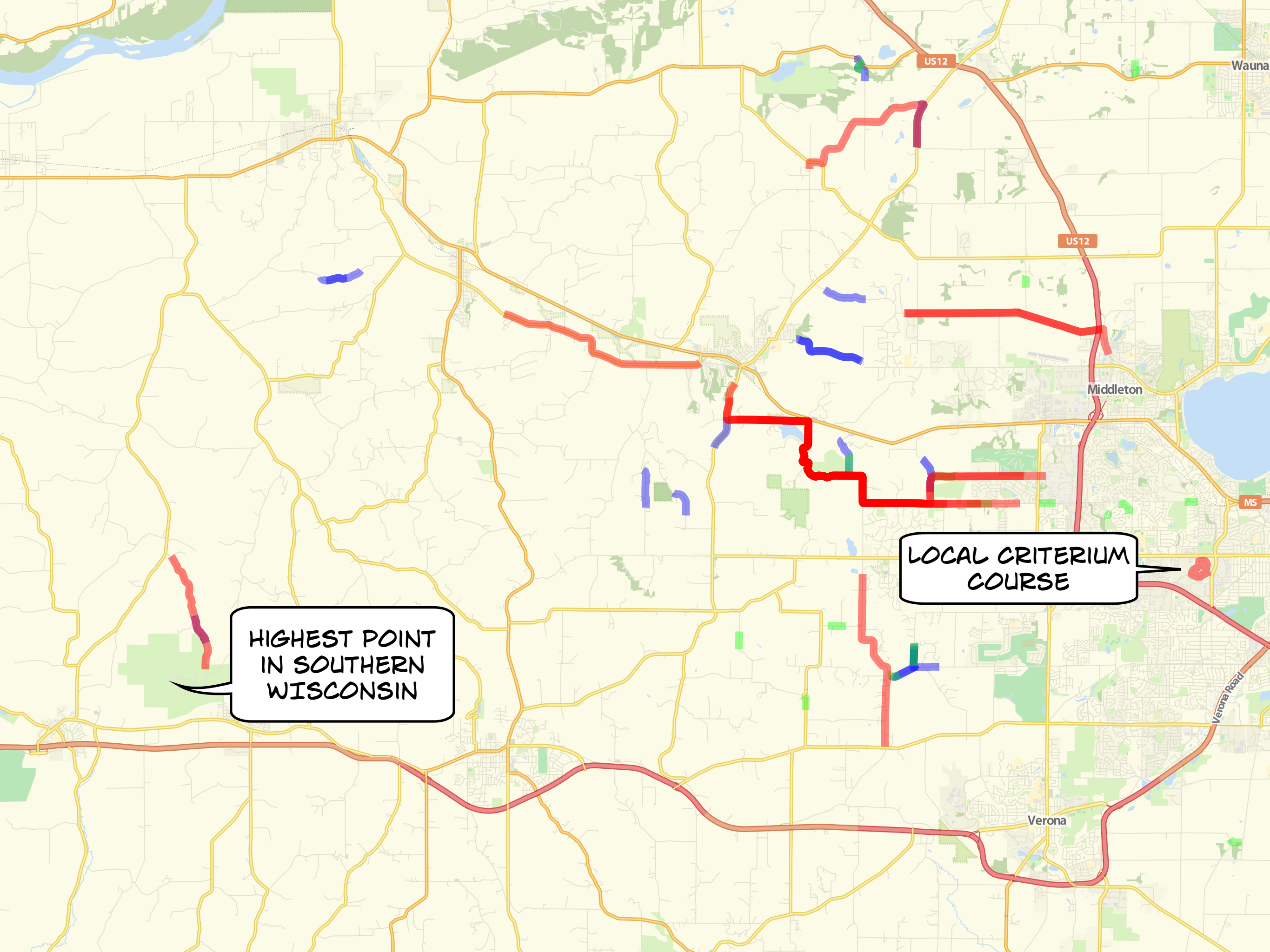




**START BY FINDING CLUSTERS
OF SPATIAL POINTS---**

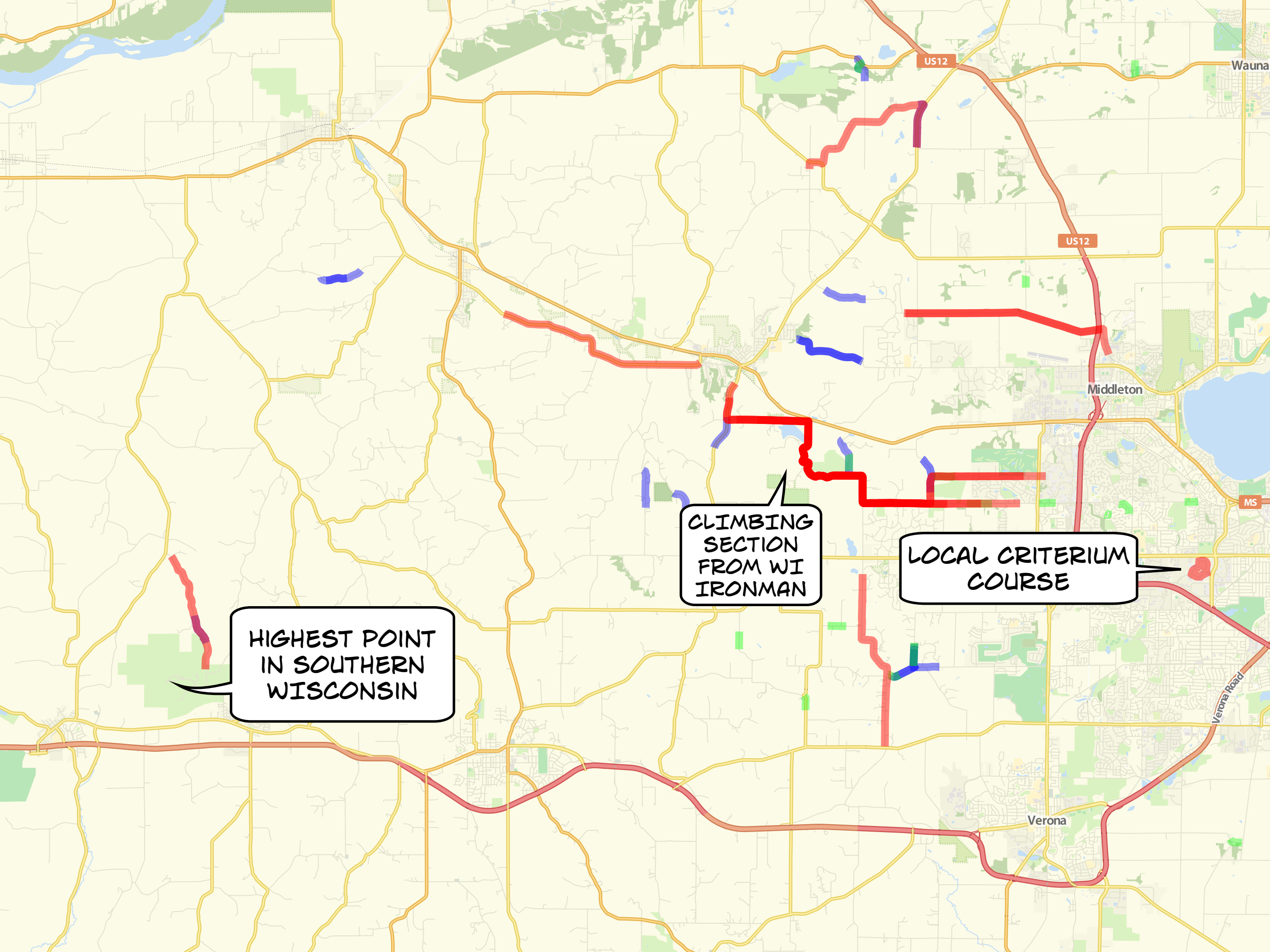
**---AND THEN CONSIDER ONLY
THE BEST EFFORT BETWEEN A
PAIR OF CLUSTERS!**





HIGHEST POINT
IN SOUTHERN
WISCONSIN

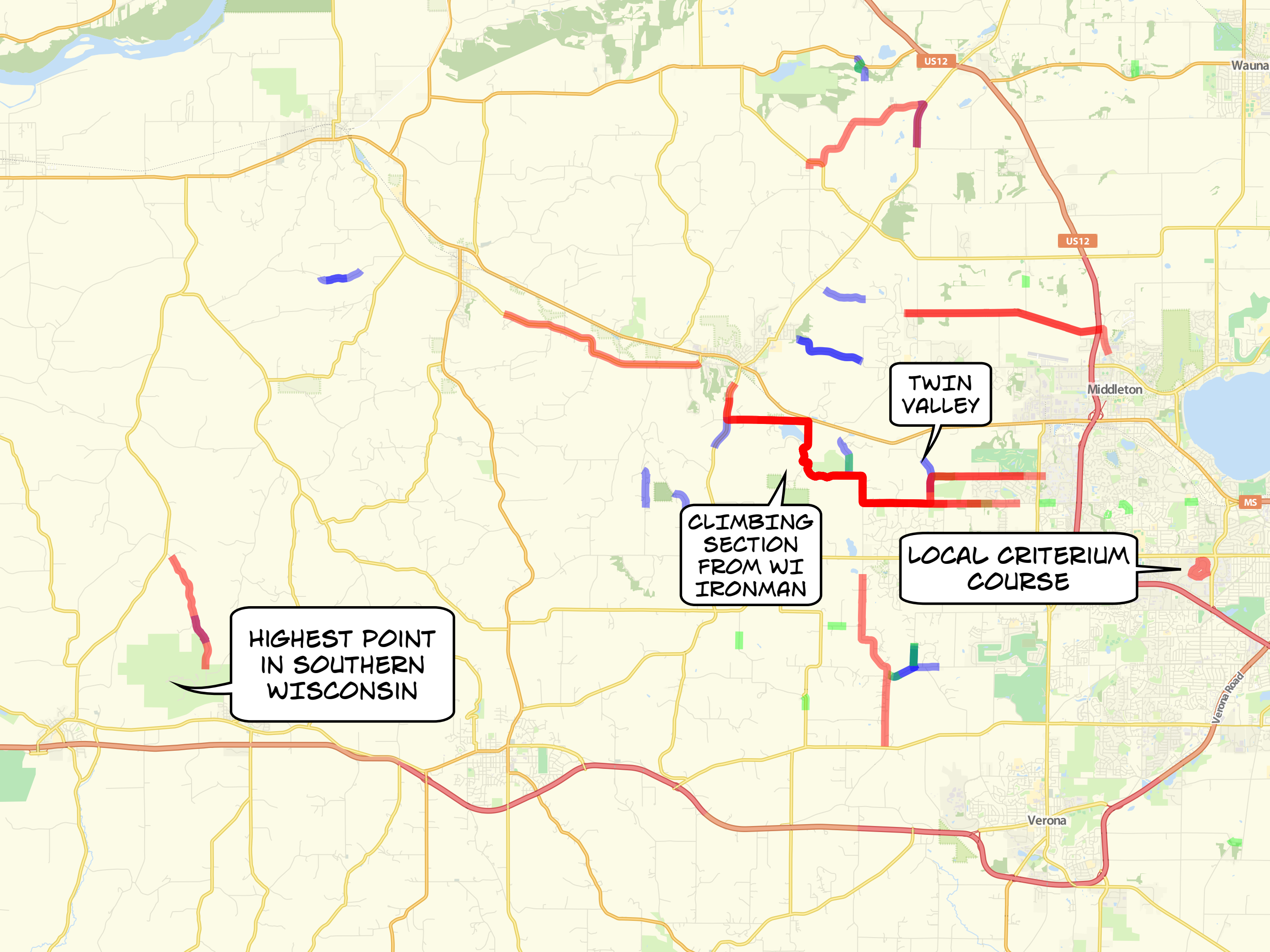
LOCAL CRITERIUM
COURSE



HIGHEST POINT
IN SOUTHERN
WISCONSIN

CLIMBING
SECTION
FROM WI
IRONMAN

LOCAL CRITERIUM
COURSE



HIGHEST POINT
IN SOUTHERN
WISCONSIN

CLIMBING
SECTION
FROM WI
IRONMAN

TWIN
VALLEY

LOCAL CRITERIUM
COURSE

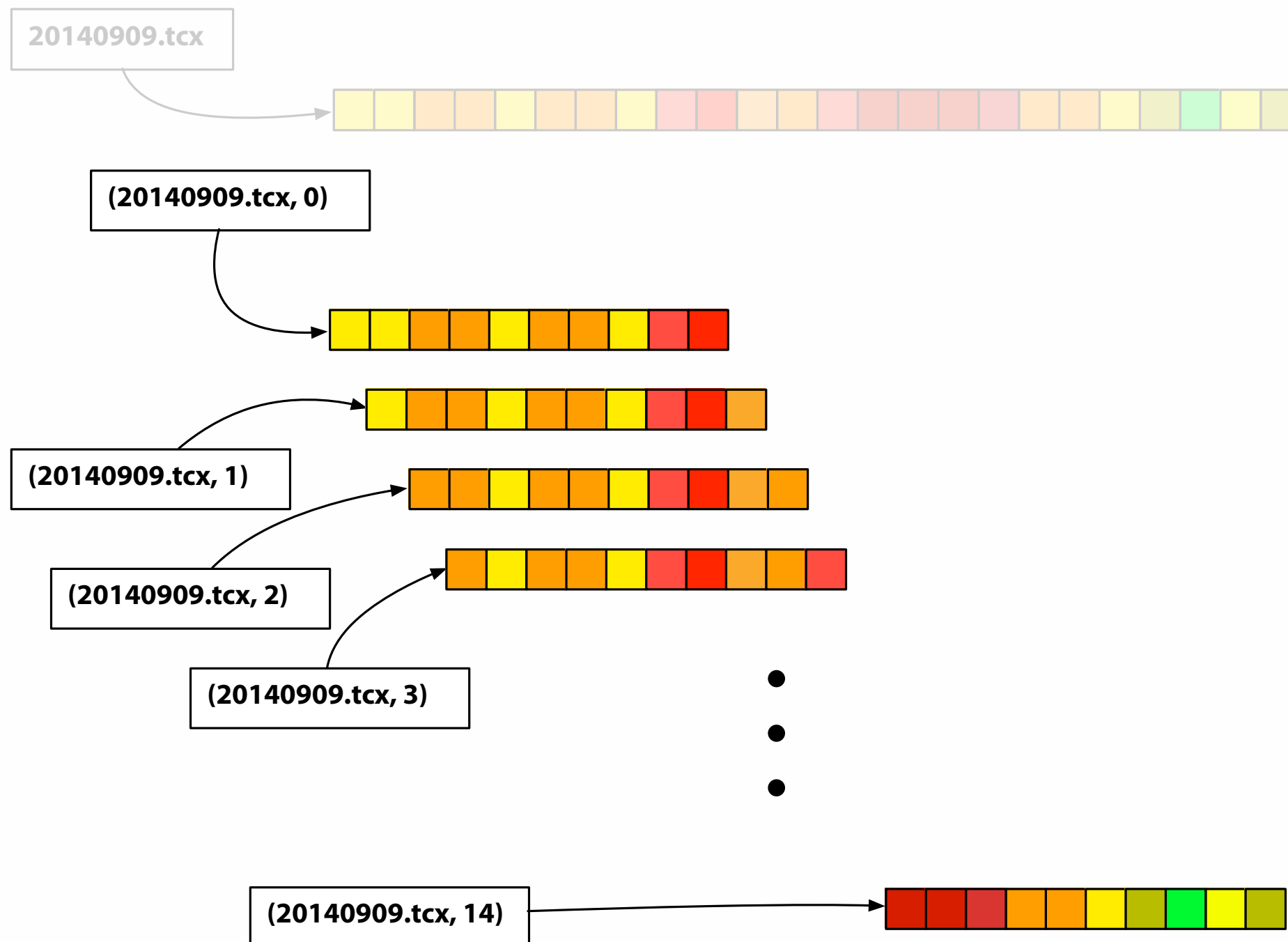
The prototype

Windowed processing

20140909.tcx

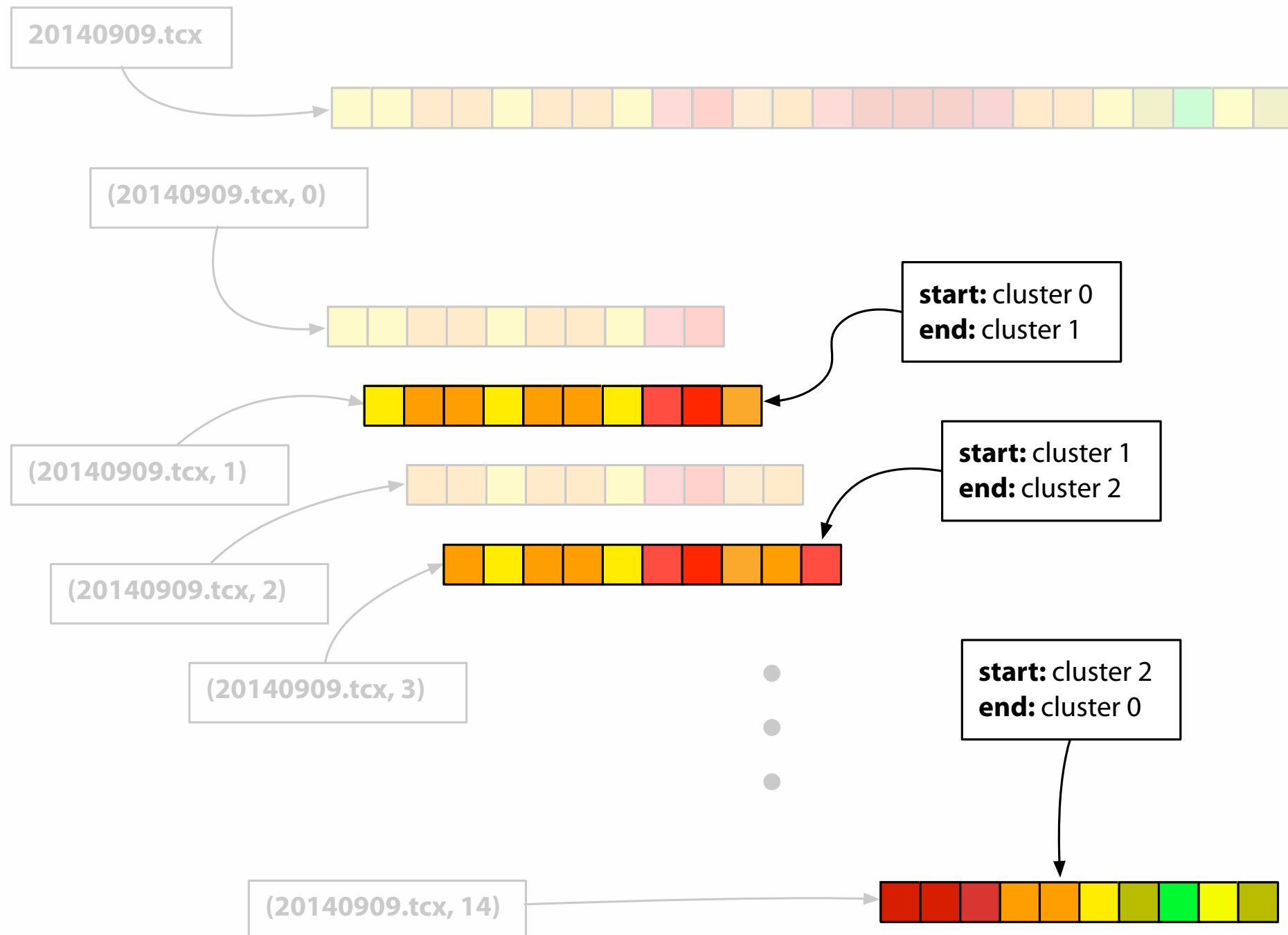


Windowed processing



*KEEP & PLOT THE BEST WINDOWS
FOR EACH SPATIAL CLUSTER PAIR!*

Windowed processing



```
trait ActivitySliding {  
  import org.apache.spark.rdd.RDD  
  import com.freevariable.surplaque.data.{Trackpoint => TP}  
  
  def windowsForActivities[U](data: RDD[TP], period: Int,  
                              xform: (TP => U) = identity _) = {  
    val pairs = data.groupBy((tp:TP) => tp.activity.getOrElse("UNKNOWN"))  
    pairs.flatMap {  
      case (activity:String, stp:Seq[TP]) =>  
        (stp sliding period).zipWithIndex.map {  
          case (s,i) => ((activity, i), s.map(xform))  
        }  
    }  
  }  
  
  def identity(tp: Trackpoint) = tp  
}
```

```
trait ActivitySliding {  
  import org.apache.spark.rdd.RDD  
  import com.freevariable.surplaque.data.{Trackpoint => TP}  
  
  def windowsForActivities[U](data: RDD[TP], period: Int,  
                              xform: (TP => U) = identity _) = {  
    val pairs = data.groupBy((tp:TP) => tp.activity.getOrElse("UNKNOWN"))  
    pairs.flatMap {  
      case (activity:String, stp:Seq[TP]) =>  
        (stp sliding period).zipWithIndex.map {  
          case (s,i) => ((activity, i), s.map(xform))  
        }  
    }  
  }  
  
  def identity(tp: Trackpoint) = tp  
}
```

```
trait ActivitySliding {  
  import org.apache.spark.rdd.RDD  
  import com.freevariable.surplaque.data.{Trackpoint => TP}  
  
  def windowsForActivities[U](data: RDD[TP], period: Int,  
                             xform: (TP => U) = identity _) = {  
    val pairs = data.groupBy((tp:TP) => tp.activity.getOrElse("UNKNOWN"))  
    pairs.flatMap {  
      case (activity:String, stp:Seq[TP]) =>  
        (stp sliding period).zipWithIndex.map {  
          case (s,i) => ((activity, i), s.map(xform))  
        }  
    }  
  }  
  
  def identity(tp: Trackpoint) = tp  
}
```

```
trait ActivitySliding {  
  import org.apache.spark.rdd.RDD  
  import com.freevariable.surplaque.data.{Trackpoint => TP}  
  
  def windowsForActivities[U](data: RDD[TP], period: Int,  
                              xform: (TP => U) = identity _) = {  
    val pairs = data.groupBy((tp:TP) => tp.activity.getOrElse("UNKNOWN"))  
    pairs.flatMap {  
      case (activity:String, stp:Seq[TP]) =>  
        (stp sliding period).zipWithIndex.map {  
          case (s,i) => ((activity, i), s.map(xform))  
        }  
    }  
  }  
  
  def identity(tp: Trackpoint) = tp  
}
```

```
trait ActivitySliding {  
  import org.apache.spark.rdd.RDD  
  import com.freevariable.surplaque.data.{Trackpoint => TP}  
  
  def windowsForActivities[U](data: RDD[TP], period: Int,  
                             xform: (TP => U) = identity _) = {  
    val pairs = data.groupBy((tp:TP) => tp.activity.getOrElse("UNKNOWN"))  
    pairs.flatMap {  
      case (activity:String, stp:Seq[TP]) =>  
        (stp sliding period).zipWithIndex.map {  
          case (s,i) => ((activity, i), s.map(xform))  
        }  
    }  
  }  
  
  def identity(tp: Trackpoint) = tp  
}
```

```
trait ActivitySliding {  
  import org.apache.spark.rdd.RDD  
  import com.freevariable.surplaque.data.{Trackpoint => TP}  
  
  def windowsForActivities[U](data: RDD[TP], period: Int,  
                             xform: (TP => U) = identity _) = {  
    val pairs = data.groupBy((tp:TP) => tp.activity.getOrElse("UNKNOWN"))  
    pairs.flatMap {  
      case (activity:String, stp:Seq[TP]) =>  
        (stp sliding period).zipWithIndex.map {  
          case (s,i) => ((activity, i), s.map(xform))  
        }  
    }  
  }  
  
  def identity(tp: Trackpoint) = tp  
}
```

```
trait ActivitySliding {  
  import org.apache.spark.rdd.RDD  
  import com.freevariable.surplaque.data.{Trackpoint => TP}  
  
  def windowsForActivities[U](data: RDD[TP], period: Int,  
                              xform: (TP => U) = identity _) = {  
    val pairs = data.groupBy((tp:TP) => tp.activity.getOrElse("UNKNOWN"))  
    pairs.flatMap {  
      case (activity:String, stp:Seq[TP]) =>  
        (stp sliding period).zipWithIndex.map {  
          case (s,i) => ((activity, i), s.map(xform))  
        }  
    }  
  }  
  
  def identity(tp: Trackpoint) = tp  
}
```

TRANSFORM AN *RDD OF TRACKPOINTS...*

```
trait ActivitySliding {  
  import org.apache.spark.rdd.RDD  
  import com.freevariable.surlaplaque.data.{Trackpoint => TP}  
  
  def windowsForActivities[U](data: RDD[TP], period: Int,  
                             xform: (TP => U) = identity _) = {  
    val pairs = data.groupBy((tp:TP) => tp.activity.getOrElse("UNKNOWN"))  
    pairs.flatMap {  
      case (activity:String, stp:Seq[TP]) =>  
        (stp sliding period).zipWithIndex.map {  
          case (s,i) => ((activity, i), s.map(xform))  
        }  
    }  
  }  
  
  def identity(tp: Trackpoint) = tp  
}
```

*...TO AN RDD OF WINDOW
IDS AND SAMPLE WINDOWS*



```
def bestsForPeriod(data: RDD[TP], period: Int, app: SLP, model: KMeansModel) = {  
  val windowedSamples = windowsForActivities(data, period, minify _).cache  
  val clusterPairs = windowedSamples.map {  
    case ((act, idx), samples) =>  
      ((act, idx), clusterPairsForWindow(samples, model))  
  }  
  val mmps = windowedSamples.map {  
    case ((act, idx), samples) =>  
      ((act, idx), samples.map(_.watts).reduce(_ + _) / samples.size)  
  }  
  
  // continued...
```

```
def bestsForPeriod(data: RDD[TP], period: Int, app: SLP, model: KMeansModel) = {  
  val windowedSamples = windowsForActivities(data, period, minify _).cache  
  val clusterPairs = windowedSamples.map {  
    case ((act, idx), samples) =>  
      ((act, idx), clusterPairsForWindow(samples, model))  
  }  
  val mmps = windowedSamples.map {  
    case ((act, idx), samples) =>  
      ((act, idx), samples.map(_.watts).reduce(_ + _) / samples.size)  
  }  
  
  // continued...
```

***DIVIDE THE INPUT DATA INTO
OVERLAPPING WINDOWS, KEYED
BY ACTIVITY AND OFFSET (WE'LL
CALL THIS KEY A WINDOW ID)***

```
def bestsForPeriod(data: RDD[TP], period: Int, app: SLP, model: KMeansModel) = {  
  val windowedSamples = windowsForActivities(data, period, minify _).cache  
  val clusterPairs = windowedSamples.map {  
    case ((act, idx), samples) =>  
      ((act, idx), clusterPairsForWindow(samples, model))  
  }  
  val mmps = windowedSamples.map {  
    case ((act, idx), samples) =>  
      ((act, idx), samples.map(_.watts).reduce(_ + _) / samples.size)  
  }  
  
  // continued...
```

**IDENTIFY THE *SPATIAL*
CLUSTERS THAT EACH WINDOW
*STARTS AND ENDS IN***

```
def bestsForPeriod(data: RDD[TP], period: Int, app: SLP, model: KMeansModel) = {  
  val windowedSamples = windowsForActivities(data, period, minify _).cache  
  val clusterPairs = windowedSamples.map {  
    case ((act, idx), samples) =>  
      ((act, idx), clusterPairsForWindow(samples, model))  
  }  
  val mmps = windowedSamples.map {  
    case ((act, idx), samples) =>  
      ((act, idx), samples.map(_.watts).reduce(_ + _) / samples.size)  
  }  
  
  // continued...
```

**IDENTIFY THE MEAN WATTAGE
FOR EACH WINDOW OF SAMPLES**



```
def bestsForPeriod(data: RDD[TP], period: Int, app: SLP, model: KMeansModel) = {  
  val windowedSamples = /* window IDs and raw sample windows */  
  val clusterPairs = /* window IDs and spatial cluster pairs */  
  val mmgs = /* window IDs and mean wattages for each window */  
  
  val top20 = mmgs.join(clusterPairs)  
    .map { case ((act, idx), (watts, (cls1, cls2))) =>  
            ((cls1, cls2), (watts, (act, idx))) }  
    .reduceByKey ((a, b) => if (a._1 > b._1) a else b)  
    .map { case ((cls1, cls2), (watts, (act, idx))) => (watts, (act, idx)) }  
    .sortByKey(false)  
    .take(20)  
  
  app.context.parallelize(top20)  
    .map { case (watts, (act, idx)) => ((act, idx), watts) }  
    .join(windowedSamples)  
    .map { case ((act, idx), (watts, samples)) => (watts, samples) }  
    .collect  
}
```

```
def bestsForPeriod(data: RDD[TP], period: Int, app: SLP, model: KMeansModel) = {  
  val windowedSamples = /* window IDs and raw sample windows */  
  val clusterPairs = /* window IDs and spatial cluster pairs */  
  val mmps = /* window IDs and mean wattages for each window */  
  
  val top20 = mmps.join(clusterPairs)  
    .map { case ((act, idx), (watts, (cls1, cls2))) =>  
      ((cls1, cls2), (watts, (act, idx))) }  
    .reduceByKey ((a, b) => if (a._1 > b._1) a else b)  
    .map { case ((cls1, cls2), (watts, (act, idx))) => (watts, (act, idx)) }  
    .sortByKey(false)  
    .take(20)  
  
  app.context.parallelize(top20)  
    .map { case (watts, (act, idx)) => ((act, idx), watts) }  
    .join(windowedSamples)  
    .map { case ((act, idx), (watts, samples)) => (watts, samples) }  
    .collect  
}
```

**FOR EACH WINDOW ID, JOIN ITS MEAN
WATTAGES WITH ITS SPATIAL CLUSTERS**

```
def bestsForPeriod(data: RDD[TP], period: Int, app: SLP, model: KMeansModel) = {  
  val windowedSamples = /* window IDs and raw sample windows */  
  val clusterPairs = /* window IDs and spatial cluster pairs */  
  val mmps = /* window IDs and mean wattages for each window */  
  
  val top20 = mmps.join(clusterPairs)  
    .map { case ((act, idx), (watts, (cls1, cls2))) =>  
      ((cls1, cls2), (watts, (act, idx))) }  
    .reduceByKey ((a, b) => if (a._1 > b._1) a else b)  
    .map { case ((cls1, cls2), (watts, (act, idx))) => (watts, (act, idx)) }  
    .sortByKey(false)  
    .take(20)  
  
  app.context.parallelize(top20)  
    .map { case (watts, (act, idx)) => ((act, idx), watts) }  
    .join(windowedSamples)  
    .map { case ((act, idx), (watts, samples)) => (watts, samples) }  
    .collect  
}
```

***TRANSPOSE THESE TUPLES SO THEY ARE
KEYED BY SPATIAL CLUSTER PAIRS***

```
def bestsForPeriod(data: RDD[TP], period: Int, app: SLP, model: KMeansModel) = {  
  val windowedSamples = /* window IDs and raw sample windows */  
  val clusterPairs = /* window IDs and spatial cluster pairs */  
  val mmps = /* window IDs and mean wattages for each window */  
  
  val top20 = mmps.join(clusterPairs)  
    .map { case ((act, idx), (watts, (cls1, cls2))) =>  
      ((cls1, cls2), (watts, (act, idx))) }  
    .reduceByKey ((a, b) => if (a._1 > b._1) a else b)  
    .map { case ((cls1, cls2), (watts, (act, idx))) => (watts, (act, idx)) }  
    .sortByKey(false)  
    .take(20)  
  
  app.context.parallelize(top20)  
    .map { case (watts, (act, idx)) => ((act, idx), watts) }  
    .join(windowedSamples)  
    .map { case ((act, idx), (watts, samples)) => (watts, samples) }  
    .collect  
}
```

***KEEP ONLY THE BEST WATTAGE FOR
EACH SPATIAL CLUSTER PAIR***



```
def bestsForPeriod(data: RDD[TP], period: Int, app: SLP, model: KMeansModel) = {  
  val windowedSamples = /* window IDs and raw sample windows */  
  val clusterPairs = /* window IDs and spatial cluster pairs */  
  val mmgs = /* window IDs and mean wattages for each window */  
  
  val top20 = mmgs.join(clusterPairs)  
    .map { case ((act, idx), (watts, (cls1, cls2))) =>  
            ((cls1, cls2), (watts, (act, idx))) }  
    .reduceByKey ((a, b) => if (a._1 > b._1) a else b)  
    .map { case ((cls1, cls2), (watts, (act, idx))) => (watts, (act, idx)) }  
    .sortByKey(false)  
    .take(20)  
  
  app.context.parallelize(top20)  
    .map { case (watts, (act, idx)) => ((act, idx), watts) }  
    .join (windowedSamples)  
    .map { case ((act, idx), (watts, samples)) => (watts, samples) }  
    .collect  
}
```

PROJECT AWAY THE CLUSTER CENTERS

```
def bestsForPeriod(data: RDD[TP], period: Int, app: SLP, model: KMeansModel) = {  
  val windowedSamples = /* window IDs and raw sample windows */  
  val clusterPairs = /* window IDs and spatial cluster pairs */  
  val mmps = /* window IDs and mean wattages for each window */  
  
  val top20 = mmps.join(clusterPairs)  
    .map { case ((act, idx), (watts, (cls1, cls2))) =>  
      ((cls1, cls2), (watts, (act, idx))) }  
    .reduceByKey ((a, b) => if (a._1 > b._1) a else b)  
    .map { case ((cls1, cls2), (watts, (act, idx))) => (watts, (act, idx)) }  
    .sortByKey(false)  
    .take(20)  
  
  app.context.parallelize(top20)  
    .map { case (watts, (act, idx)) => ((act, idx), watts) }  
    .join(windowedSamples)  
    .map { case ((act, idx), (watts, samples)) => (watts, samples) }  
    .collect  
}
```

***SORT BY WATTAGE IN DESCENDING
ORDER; KEEP THE BEST TWENTY***

```
def bestsForPeriod(data: RDD[TP], period: Int, app: SLP, model: KMeansModel) = {  
  val windowedSamples = /* window IDs and raw sample windows */  
  val clusterPairs = /* window IDs and spatial cluster pairs */  
  val mmps = /* window IDs and mean wattages for each window */  
  
  val top20 = mmps.join(clusterPairs)  
    .map { case ((act, idx), (watts, (cls1, cls2))) =>  
      ((cls1, cls2), (watts, (act, idx))) }  
    .reduceByKey ((a, b) => if (a._1 > b._1) a else b)  
    .map { case ((cls1, cls2), (watts, (act, idx))) => (watts, (act, idx)) }  
    .sortByKey(false)  
    .take(20)  
  
  app.context.parallelize(top20)  
    .map { case (watts, (act, idx)) => ((act, idx), watts) }  
    .join(windowedSamples)  
    .map { case ((act, idx), (watts, samples)) => (watts, samples) }  
    .collect  
}
```

***RE-KEY THE BEST EFFORTS
BY WINDOW ID***

```
def bestsForPeriod(data: RDD[TP], period: Int, app: SLP, model: KMeansModel) = {  
  val windowedSamples = /* window IDs and raw sample windows */  
  val clusterPairs = /* window IDs and spatial cluster pairs */  
  val mmps = /* window IDs and mean wattages for each window */  
  
  val top20 = mmps.join(clusterPairs)  
    .map { case ((act, idx), (watts, (cls1, cls2))) =>  
      ((cls1, cls2), (watts, (act, idx))) }  
    .reduceByKey ((a, b) => if (a._1 > b._1) a else b)  
    .map { case ((cls1, cls2), (watts, (act, idx))) => (watts, (act, idx)) }  
    .sortByKey(false)  
    .take(20)  
  
  app.context.parallelize(top20)  
    .map { case (watts, (act, idx)) => ((act, idx), watts) }  
    .join(windowedSamples)  
    .map { case ((act, idx), (watts, samples)) => (watts, samples) }  
    .collect  
}
```

**GET THE ACTUAL SAMPLE WINDOWS
FOR EACH EFFORT; PROJECT AWAY IDS**



```
def bestsForPeriod(data: RDD[TP], period: Int, app: SLP, model: KMeansModel) = {  
  val windowedSamples = /* window IDs and raw sample windows */  
  val clusterPairs = /* window IDs and spatial cluster pairs */  
  val mmgs = /* window IDs and mean wattages for each window */  
  
  val top20 = mmgs.join(clusterPairs)  
    .map { case ((act, idx), (watts, (cls1, cls2))) =>  
            ((cls1, cls2), (watts, (act, idx))) }  
    .reduceByKey ((a, b) => if (a._1 > b._1) a else b)  
    .map { case ((cls1, cls2), (watts, (act, idx))) => (watts, (act, idx)) }  
    .sortByKey(false)  
    .take(20)  
  
  app.context.parallelize(top20)  
    .map { case (watts, (act, idx)) => ((act, idx), watts) }  
    .join (windowedSamples)  
    .map { case ((act, idx), (watts, samples)) => (watts, samples) }  
    .collect  
}
```

**Improving the
prototype**

**Broadcast large
static data**

Broadcast variables

```
// phoneBook maps (given name, surname) -> phone number digits
val phoneBook: Map[(String, String), String] = initPhoneBook()
val names: RDD[(String, String)] = /* ... */

val directory = names.map {
  case name @ (first, last) => (name, phoneBook.getOrElse("555-1212"))
}
```

Broadcast variables

```
// phoneBook maps (given name, surname) -> phone number digits  
val phoneBook: Map[(String, String), String] = initPhoneBook()  
val names: RDD[(String, String)] = /* ... */
```

```
val directory = names.map {  
  case name @ (first, last) => (name, phoneBook.getOrElse("555-1212"))  
}
```

**PHONEBOOK WILL BE COPIED AND
DESERIALIZED FOR EACH TASK!**

Broadcast variables

```
// phoneBook maps (given name, surname) -> phone number digits
val phoneBook: Map[(String, String), String] = initPhoneBook()
val names: RDD[(String, String)] = /* ... */
val pbb = sparkContext.broadcast(phoneBook)

val directory = names.map {
  case name @ (first, last) => (name, pbb.value.getOrElse("555-1212"))
}
```

**BROADCASTING PHONEBOOK MEANS
IT CAN BE DESERIALIZED ONCE AND
CACHED ON EACH NODE!**

```
def bestsForPeriod(data: RDD[Trackpoint], period: Int,  
                  app: SLP, model: KMeansModel) = {  
  val windowedSamples = windowsForActivities(data, period, minify _)  
  val clusterPairs = windowedSamples.map {  
    case ((act, idx), samples) =>  
      ((act, idx), clusterPairsForWindow(samples, model))  
  
    // ...  
  
  }
```

```
def bestsForPeriod(data: RDD[Trackpoint], period: Int,  
                  app: SLP, model: Broadcast[KMeansModel]) = {  
  val windowedSamples = windowsForActivities(data, period, minify _)  
  val clusterPairs = windowedSamples.map {  
    case ((act, idx), samples) =>  
      ((act, idx), clusterPairsForWindow(samples, model.value))  
  
    // rest of function unchanged  
  
  }
```

**Cache only
when necessary**



```
def bestsForPeriod(data: RDD[TP], period: Int, app: SLP, model: KMeansModel) = {  
  val windowedSamples = windowsForActivities(data, period).cache  
  
  // ...  
  
  val top20 = mmps.join(clusterPairs)  
    .map { case ((act, idx), (watts, (cls1, cls2))) =>  
            ((cls1, cls2), (watts, (act, idx))) }  
    .reduceByKey ((a, b) => if (a._1 > b._1) a else b)  
    .map { case ((cls1, cls2), (watts, (act, idx))) => (watts, (act, idx)) }  
    .sortByKey(false)  
    .take(20)  
  
  // ...  
}
```

KEEPING EVERY WINDOW IN MEMORY---

```
def bestsForPeriod(data: RDD[TP], period: Int, app: SLP, model: KMeansModel) = {  
  val windowedSamples = windowsForActivities(data, period).cache  
  
  // ...  
  
  val top20 = mmps.join(clusterPairs)  
    .map { case ((act, idx), (watts, (cls1, cls2))) =>  
      ((cls1, cls2), (watts, (act, idx))) }  
    .reduceByKey ((a, b) => if (a._1 > b._1) a else b)  
    .map { case ((cls1, cls2), (watts, (act, idx))) => (watts, (act, idx)) }  
    .sortByKey(false)  
    .take(20)  
  
  // ...  
}
```

KEEPING EVERY WINDOW IN MEMORY---

```
def bestsForPeriod(data: RDD[TP], period: Int, app: SLP, model: KMeansModel) = {  
  val windowedSamples = windowsForActivities(data, period).cache  
  
  // ...  
  
  val top20 = mmps.join(clusterPairs)  
    .map { case ((act, idx), (watts, (cls1, cls2))) =>  
      ((cls1, cls2), (watts, (act, idx))) }  
    .reduceByKey ((a, b) => if (a._1 > b._1) a else b)  
    .map { case ((cls1, cls2), (watts, (act, idx))) => (watts, (act, idx)) }  
    .sortByKey(false)  
    .take(20)  
  
  // ...  
}
```

---EVEN THOUGH RECOMPUTING WINDOWS
IS INCREDIBLY CHEAP AND YOU'LL NEED
ONLY A TINY FRACTION OF WINDOWS?

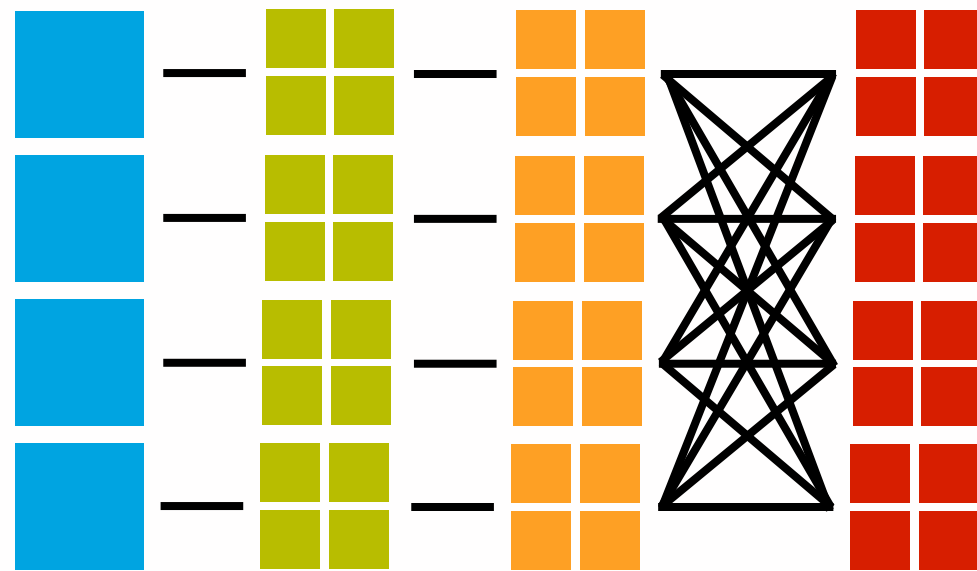


```
def bestsForPeriod(data: RDD[TP], period: Int, app: SLP, model: KMeansModel) = {  
  val windowedSamples = windowsForActivities(data, period).cache  
  
  // ...  
  
  val top20 = mmps.join(clusterPairs)  
    .map { case ((act, idx), (watts, (cls1, cls2))) =>  
            ((cls1, cls2), (watts, (act, idx))) }  
    .reduceByKey ((a, b) => if (a._1 > b._1) a else b)  
    .map { case ((cls1, cls2), (watts, (act, idx))) => (watts, (act, idx)) }  
    .sortByKey(false)  
    .take(20)  
  
  // ...  
}
```

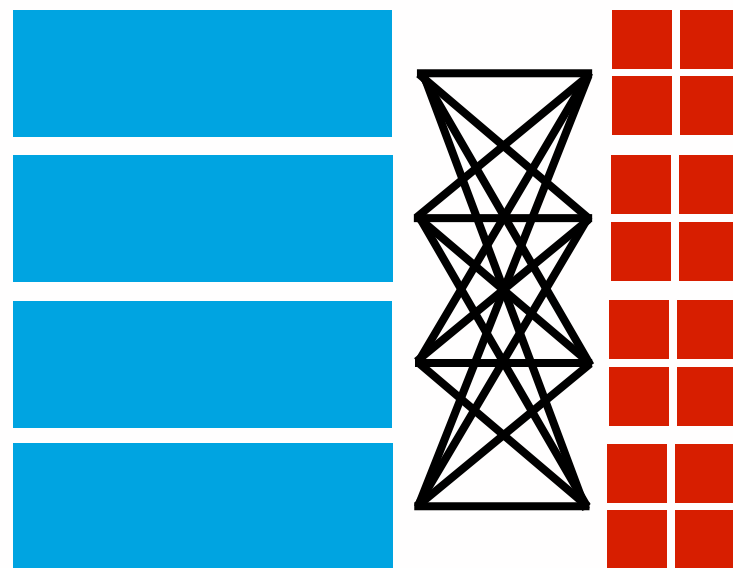
**ELIMINATING UNNECESSARY MEMORY
PRESSURE CAN LEAD TO A SUBSTANTIAL
SPEEDUP!**

**Avoid shuffles
when possible**

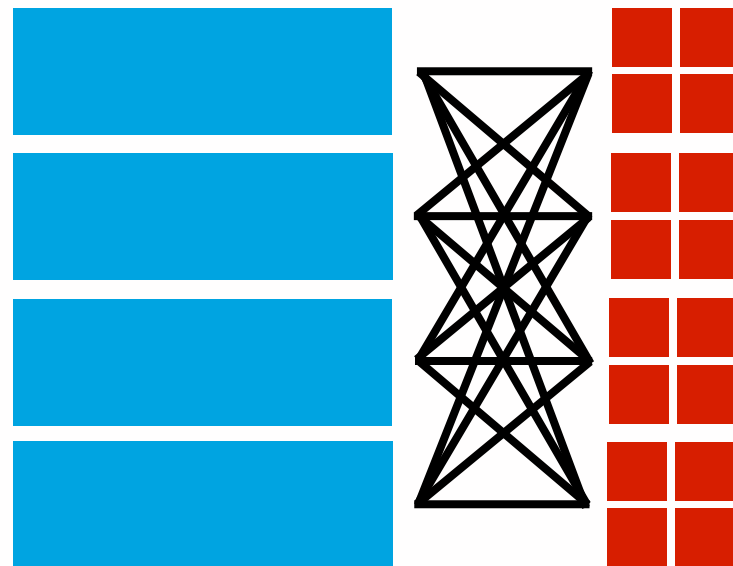
TASKS



STAGES



STAGES



*WE WANT TO AVOID ALL
UNNECESSARY SHUFFLES*



```
def bestsForPeriod(data: RDD[Trackpoint], period: Int,
                  app: SLP, model: Broadcast[KMeansModel]) = {
  val windowedSamples = windowsForActivities(data, period, minify _)

  val bests = windowedSamples.map {
    case ((act, idx), samples) => (
      clusterPairsForWindow(samples, model.value),
      ((act, idx), samples.map(_.watts).reduce(_ + _) / samples.size)
    )
  }.cache

  val top20 = bests.reduceByKey ((a, b) => if (a._2 > b._2) a else b)
    .map { case ((_, _), keep) => keep }
    .takeOrdered(20)(Ordering.by[((String, Int), Double), Double] {
      case ((_, _), watts) => -watts
    })

  app.context.parallelize(top20)
    .join (windowedSamples)
    .map { case ((act, idx), (watts, samples)) => (watts, samples) }
    .collect
}
```



```
def bestsForPeriod(data: RDD[Trackpoint], period: Int,
                  app: SLP, model: Broadcast[KMeansModel]) = {
  val windowedSamples = windowsForActivities(data, period, minify _)

  val bests = windowedSamples.map {
    case ((act, idx), samples) => (
      clusterPairsForWindow(samples, model.value),
      ((act, idx), samples.map(_.watts).reduce(_ + _) / samples.size)
    )
  }.cache

  val top20 = bests.reduceByKey ((a, b) => if (a._2 > b._2) a else b)
    .map { case ((_, _), keep) => keep }
    .takeOrdered(20)(Ordering.by[((String, Int), Double), Double] {
      case ((_, _), watts) => -watts
    })

  app.context.parallelize(top20)
    .join (windowedSamples)
    .map { case ((act, idx), (watts, samples)) => (watts, samples) }
    .collect
}
```

```
def bestsForPeriod(data: RDD[Trackpoint], period: Int,
                  app: SLP, model: Broadcast[KMeansModel]) = {
  val windowedSamples = windowsForActivities(data, period, minify _)
  START AND END CLUSTERS
  val bests = windowedSamples.map {
    case ((act, idx), samples) => (
      clusterPairsForWindow(samples, model.value),
      ((act, idx), samples.map(_.watts).reduce(_ + _) / samples.size)
    )
  }.cache

  val top20 = bests.reduceByKey ((a, b) => if (a._2 > b._2) a else b)
    .map { case ((_, _), keep) => keep }
    .takeOrdered(20)(Ordering.by[((String, Int), Double), Double] {
      case ((_, _), watts) => -watts
    })

  app.context.parallelize(top20)
    .join (windowedSamples)
    .map { case ((act, idx), (watts, samples)) => (watts, samples) }
    .collect
}
```

```
def bestsForPeriod(data: RDD[Trackpoint], period: Int,
                  app: SLP, model: Broadcast[KMeansModel]) = {
  val windowedSamples = windowsForActivities(data, period, minify _)
  START AND END CLUSTERS
  val bests = windowedSamples.map {
    case ((act, idx), samples) => (
      clusterPairsForWindow(samples, model.value),
      ((act, idx), samples.map(_.watts).reduce(_ + _) / samples.size)
    ) WINDOW IDS AND MEAN WATTAGES
  }.cache

  val top20 = bests.reduceByKey ((a, b) => if (a._2 > b._2) a else b)
    .map { case ((_, _), keep) => keep }
    .takeOrdered(20)(Ordering.by[((String, Int), Double), Double] {
      case ((_, _), watts) => -watts
    })

  app.context.parallelize(top20)
    .join (windowedSamples)
    .map { case ((act, idx), (watts, samples)) => (watts, samples) }
    .collect
}
```

```
def bestsForPeriod(data: RDD[Trackpoint], period: Int,
                  app: SLP, model: Broadcast[KMeansModel]) = {
  val windowedSamples = windowsForActivities(data, period, minify _)

  val bests = windowedSamples.map {
    case ((act, idx), samples) => (
      clusterPairsForWindow(samples, model.value),
      ((act, idx), samples.map(_.watts).reduce(_ + _) / samples.size)
    )
  }.cache
```

ELIMINATE A JOIN AND A TRANSPOSE

```
val top20 = bests.reduceByKey ((a, b) => if (a._2 > b._2) a else b)
  .map { case (_, _), keep => keep }
  .takeOrdered(20)(Ordering.by[((String, Int), Double), Double] {
    case (_, _), watts => -watts
  })
```

```
app.context.parallelize(top20)
  .join (windowedSamples)
  .map { case ((act, idx), (watts, samples)) => (watts, samples) }
  .collect
```

```
}
```

```
def bestsForPeriod(data: RDD[Trackpoint], period: Int,
                  app: SLP, model: Broadcast[KMeansModel]) = {
  val windowedSamples = windowsForActivities(data, period, minify _)

  val bests = windowedSamples.map {
    case ((act, idx), samples) => (
      clusterPairsForWindow(samples, model.value),
      ((act, idx), samples.map(_.watts).reduce(_ + _) / samples.size)
    )
  }.cache

  val top20 = bests.reduceByKey ((a, b) => if (a._2 > b._2) a else b)
    .map { case ((_, _), keep) => keep }
    .takeOrdered(20)(Ordering.by[((String, Int), Double), Double] {
      case ((_, _), watts) => -watts
    })
    (USE THE RIGHT API CALLS!)
    app.context.parallelize(top20)
    .join (windowedSamples)
    .map { case ((act, idx), (watts, samples)) => (watts, samples) }
    .collect
}
```



```
def bestsForPeriod(data: RDD[Trackpoint], period: Int,
                  app: SLP, model: Broadcast[KMeansModel]) = {
  val windowedSamples = windowsForActivities(data, period, minify _)

  val bests = windowedSamples.map {
    case ((act, idx), samples) => (
      clusterPairsForWindow(samples, model.value),
      ((act, idx), samples.map(_.watts).reduce(_ + _) / samples.size)
    )
  }.cache

  val top20 = bests.reduceByKey ((a, b) => if (a._2 > b._2) a else b)
    .map { case ((_, _), keep) => keep }
    .takeOrdered(20)(Ordering.by[((String, Int), Double), Double] {
      case ((_, _), watts) => -watts
    })

  app.context.parallelize(top20)
    .join (windowedSamples)
    .map { case ((act, idx), (watts, samples)) => (watts, samples) }
    .collect
}
```

ELIMINATE A TRANSPOSE



```
def bestsForPeriod(data: RDD[Trackpoint], period: Int,
                  app: SLP, model: Broadcast[KMeansModel]) = {
  val windowedSamples = windowsForActivities(data, period, minify _)

  val bests = windowedSamples.map {
    case ((act, idx), samples) => (
      clusterPairsForWindow(samples, model.value),
      ((act, idx), samples.map(_.watts).reduce(_ + _) / samples.size)
    )
  }.cache

  val top20 = bests.reduceByKey ((a, b) => if (a._2 > b._2) a else b)
    .map { case ((_, _), keep) => keep }
    .takeOrdered(20)(Ordering.by[((String, Int), Double), Double] {
      case ((_, _), watts) => -watts
    })

  app.context.parallelize(top20)
    .join (windowedSamples)
    .collect
    .map { case ((act, idx), (watts, samples)) => (watts, samples) }
}
```

ELIMINATE A TRANSPOSE...OR TWO!

**Embrace
laziness**

Embrace laziness

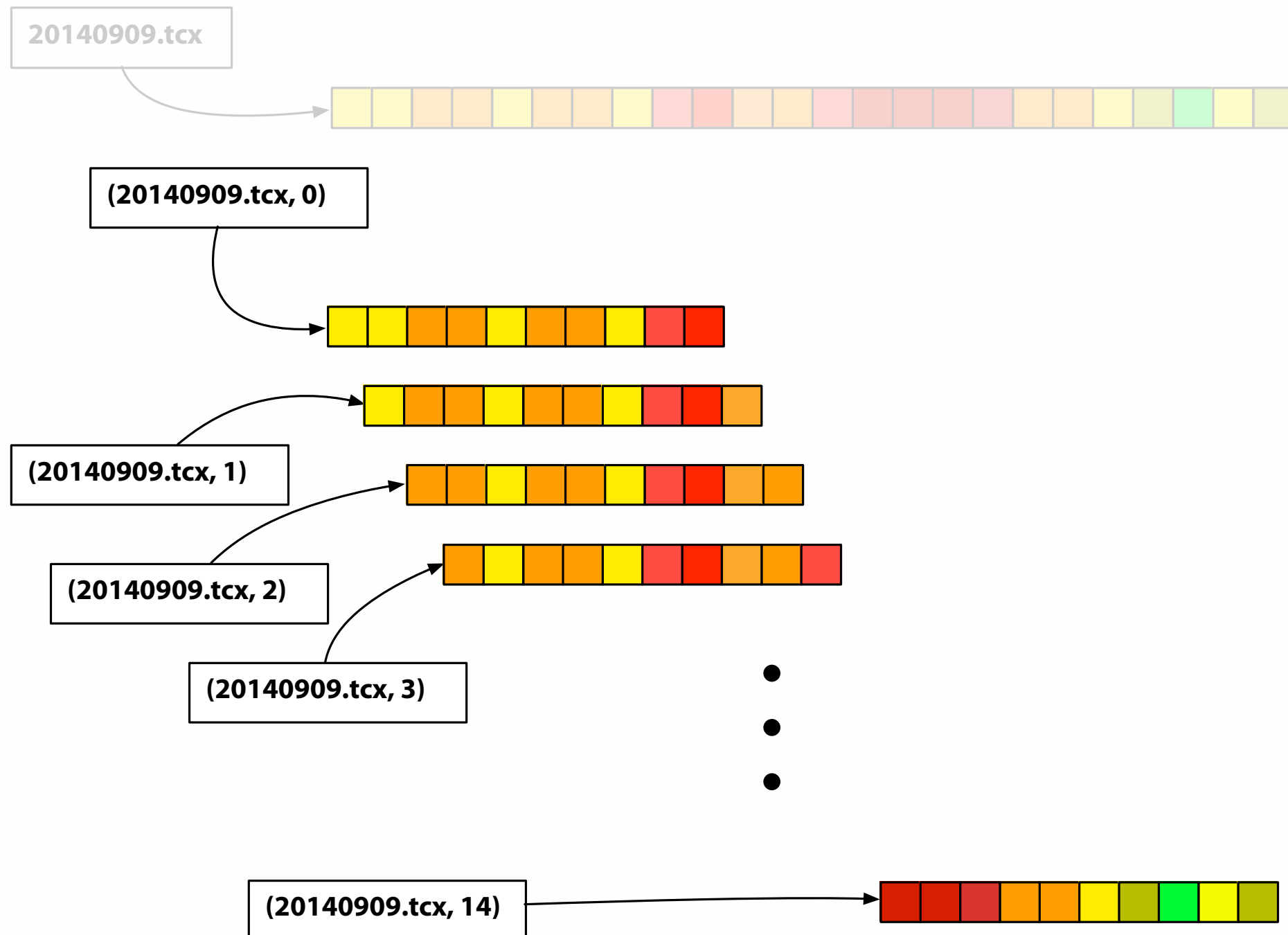
(ONLY PAY FOR WHAT YOU USE)

Windowed processing redux

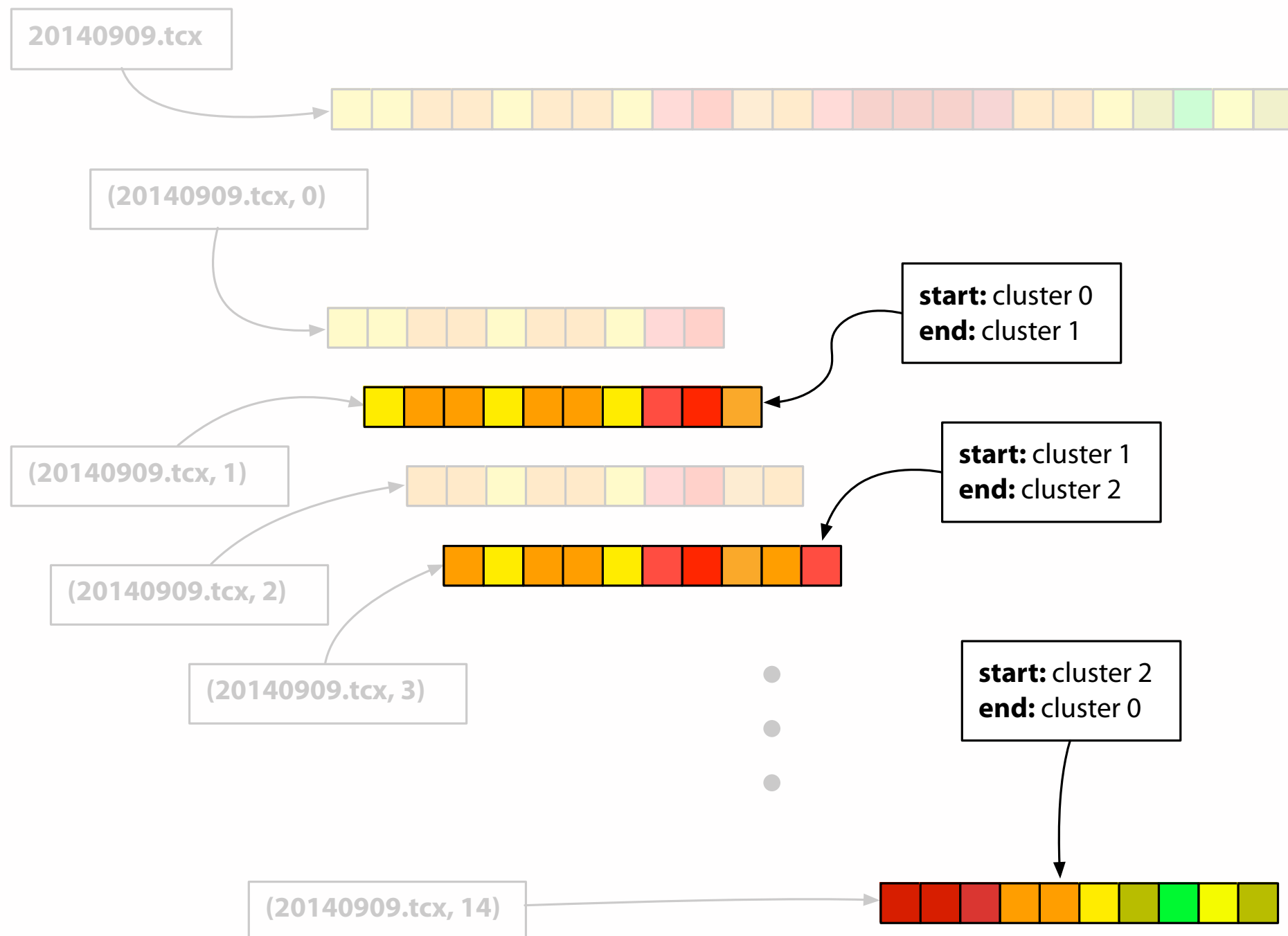
20140909.tcx



Windowed processing redux



Windowed processing redux

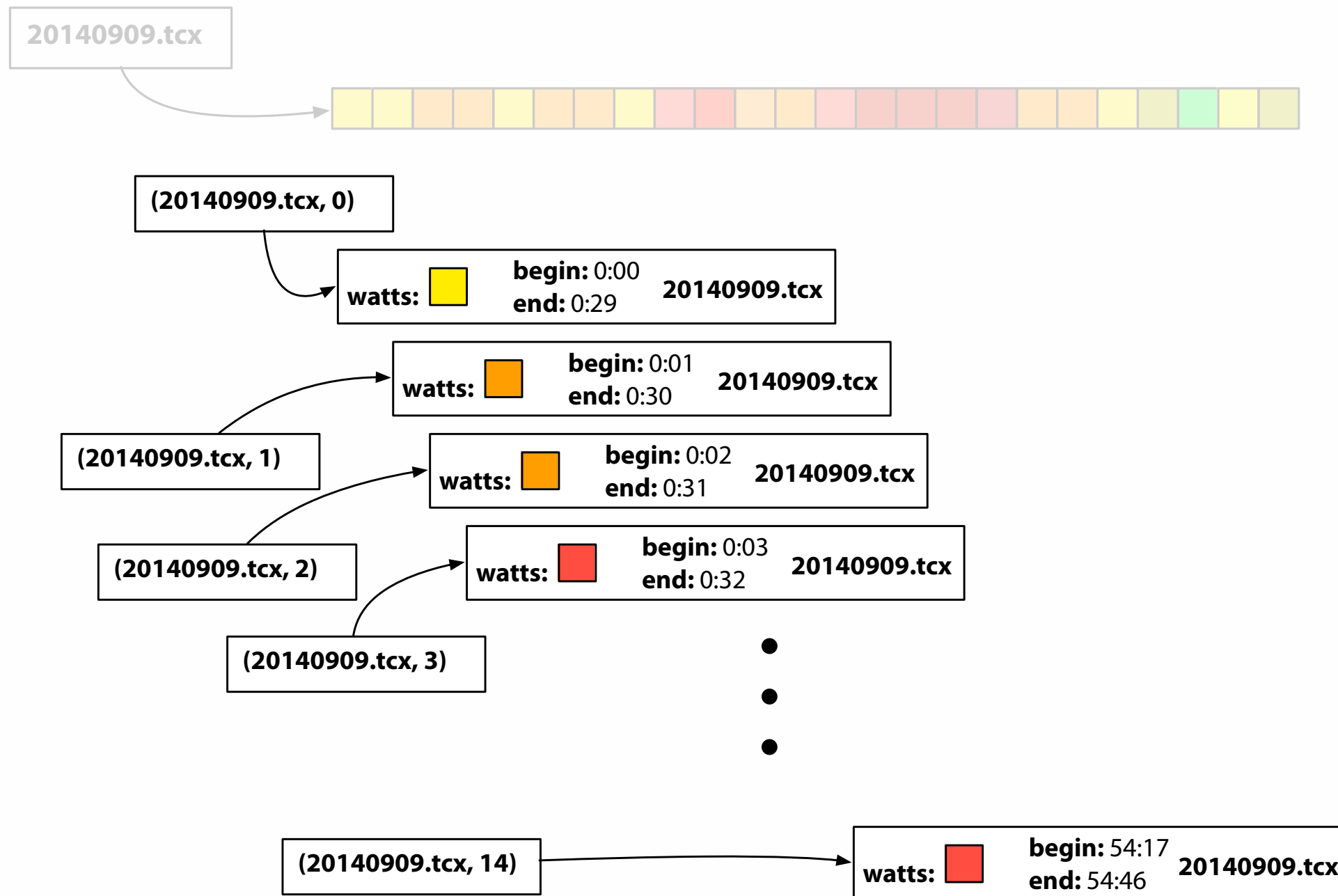


Lazy windowed processing

20140909.tcx

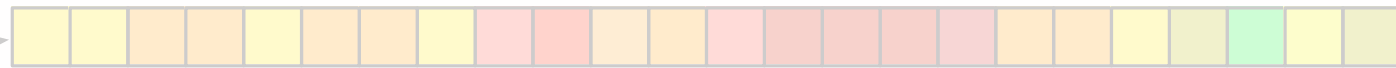


Lazy windowed processing



Lazy windowed processing

20140909.tcx



watts:  begin: 0:00 end: 0:29 20140909.tcx

start: cluster 0
end: cluster 1

watts:  begin: 0:01 end: 0:30 20140909.tcx

start: cluster 0
end: cluster 1

watts:  begin: 0:02 end: 0:31 20140909.tcx

start: cluster 1
end: cluster 2

watts:  begin: 0:03 end: 0:32 20140909.tcx

start: cluster 1
end: cluster 2

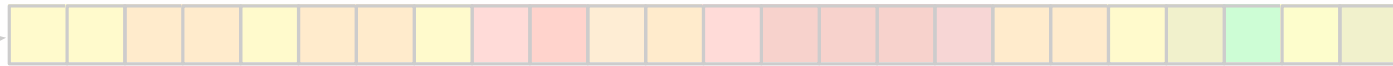
•
•
•

start: cluster 2
end: cluster 0

watts:  begin: 54:17 end: 54:46 20140909.tcx

Lazy windowed processing

20140909.tcx



watts:  **begin: 0:01**
end: 0:30 20140909.tcx

start: cluster 0
end: cluster 1

watts:  **begin: 0:03**
end: 0:32 20140909.tcx

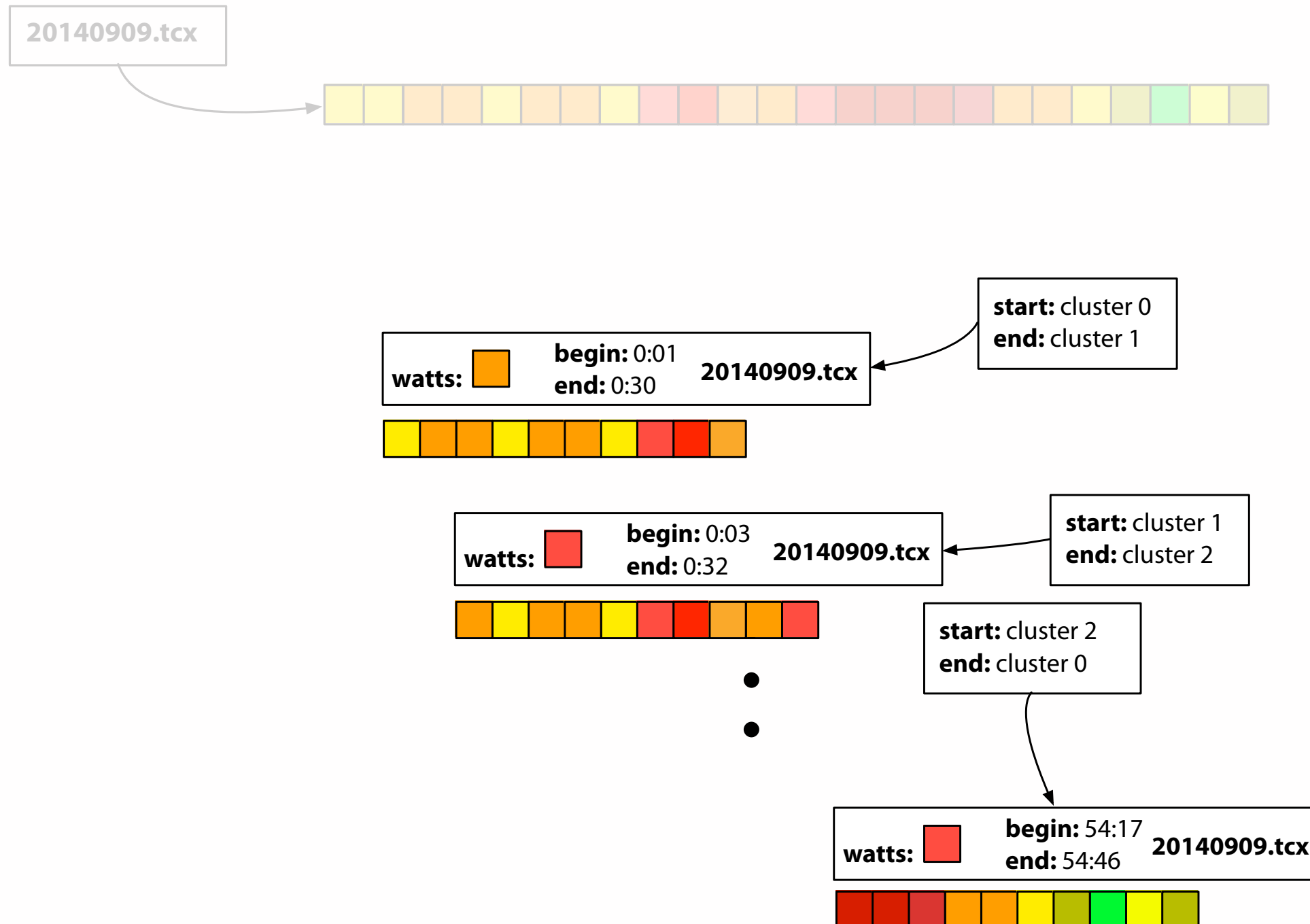
start: cluster 1
end: cluster 2

start: cluster 2
end: cluster 0

•
•
•

watts:  **begin: 54:17**
end: 54:46 20140909.tcx

Lazy windowed processing



```
trait ActivitySliding {  
  import org.apache.spark.rdd.RDD  
  import com.freevariable.surplaque.data.{Trackpoint => TP}  
  
  // ...  
  
  def applyWindowedNoZip[U: ClassTag](data: RDD[TP], period: Int,  
                                       xform: ((String, Seq[TP]) => U)) = {  
    val pairs = data.groupBy((tp:TP) => tp.activity.getOrElse("UNKNOWN"))  
    pairs.flatMap {  
      case (activity: String, stp:Seq[TP]) =>  
        (stp sliding period).map { xform(activity, _) }  
    }  
  }  
}
```

```
trait ActivitySliding {  
  import org.apache.spark.rdd.RDD  
  import com.freevariable.surplaque.data.{Trackpoint => TP}  
  
  // ...  
  
  def applyWindowedNoZip[U: ClassTag](data: RDD[TP], period: Int,  
                                       xform: ((String, Seq[TP]) => U)) = {  
    val pairs = data.groupBy((tp:TP) => tp.activity.getOrElse("UNKNOWN"))  
    pairs.flatMap {  
      case (activity: String, stp:Seq[TP]) =>  
        (stp sliding period).map { xform(activity, _) }  
    }  
  }  
}
```

**PERFORM ARBITRARY TRANSFORMATIONS
ON EACH WINDOW OF SAMPLES**

```
case class Effort(mmp: Double,  
                  activity: String,  
                  startTimestamp: Long,  
                  endTimestamp: Long) {}
```

**MODEL EFFORT SUMMARIES AS
SIMPLE, LIGHTWEIGHT RECORDS**



```
def bestsForPeriod(data: RDD[Trackpoint], period: Int,  
                  app: SLP, model: Broadcast[KMeansModel]) = {  
  val clusteredMMPs = applyWindowedNoZip(data, period, {  
    case (activity:String, samples:Seq[Trackpoint]) =>  
      (  
        clusterPairsForWindow(samples, model.value),  
        Effort(samples.map(_.watts).reduce(_ + _) / samples.size, activity,  
              samples.head.timestamp, samples.last.timestamp)  
      )  
    })  
  })  
  
  // continued  
}
```



```
def bestsForPeriod(data: RDD[Trackpoint], period: Int,
                  app: SLP, model: Broadcast[KMeansModel]) = {
  val clusteredMMPs = applyWindowedNoZip(data, period, {
    case (activity:String, samples:Seq[Trackpoint]) =>
      (
        clusterPairsForWindow(samples, model.value),
        Effort(samples.map(_.watts).reduce(_ + _) / samples.size, activity,
          samples.head.timestamp, samples.last.timestamp)
      )
  })

  // continued
}
```



```
def bestsForPeriod(data: RDD[Trackpoint], period: Int,  
                  app: SLP, model: Broadcast[KMeansModel]) = {  
  val clusteredMMPs = applyWindowedNoZip(data, period, {  
    case (activity:String, samples:Seq[Trackpoint]) =>  
      (  
        clusterPairsForWindow(samples, model.value),  
        Effort(samples.map(_.watts).reduce(_ + _) / samples.size, activity,  
               samples.head.timestamp, samples.last.timestamp)  
      )  
    })  
  
  // continued  
}
```

```
def bestsForPeriod(data: RDD[Trackpoint], period: Int,
                  app: SLP, model: Broadcast[KMeansModel]) = {
  val clusteredMMPs = applyWindowedNoZip(data, period, {
    case (activity:String, samples:Seq[Trackpoint]) =>
      (
        clusterPairsForWindow(samples, model.value),
        Effort(samples.map(_.watts).reduce(_ + _) / samples.size, activity,
              samples.head.timestamp, samples.last.timestamp)
      )
  })

  // continued
}
```



```
def bestsForPeriod(data: RDD[Trackpoint], period: Int,
                  app: SLP, model: Broadcast[KMeansModel]) = {
  val clusteredMMPs = /* map from cluster pairs to effort summary structures */

  clusteredMMPs
    .reduceByKey ((a, b) => if (a.mmp > b.mmp) a else b)
    .takeOrdered(20)(Ordering.by[((Int, Int), Effort), Double] {
      case (_, e: Effort) => -e.mmp
    })
    .map {
      case (_, e: Effort) => (
        e.mmp,
        data.filter {
          case tp: Trackpoint => tp.activity.getOrElse("UNKNOWN") == e.activity &&
            tp.timestamp <= e.endTimeStamp && tp.timestamp >= e.startTimeStamp
        }.collect
      )
    }
}
```



```
def bestsForPeriod(data: RDD[Trackpoint], period: Int,
                  app: SLP, model: Broadcast[KMeansModel]) = {
  val clusteredMMPs = /* map from cluster pairs to effort summary structures */

  clusteredMMPs
    .reduceByKey ((a, b) => if (a.mmp > b.mmp) a else b)
    .takeOrdered(20)(Ordering.by[((Int, Int), Effort), Double] {
      case (_, e: Effort) => -e.mmp
    })
    .map {
      case (_, e: Effort) => (
        e.mmp,
        data.filter {
          case tp: Trackpoint => tp.activity.getOrElse("UNKNOWN") == e.activity &&
            tp.timestamp <= e.endTimeStamp && tp.timestamp >= e.startTimeStamp
        }.collect
      )
    }
}
```



```
def bestsForPeriod(data: RDD[Trackpoint], period: Int,
                  app: SLP, model: Broadcast[KMeansModel]) = {
  val clusteredMMPs = /* map from cluster pairs to effort summary structures */

  clusteredMMPs
    .reduceByKey ((a, b) => if (a.mmp > b.mmp) a else b)
    .takeOrdered(20)(Ordering.by[((Int, Int), Effort), Double] {
      case (_, e: Effort) => -e.mmp
    })
    .map {
      case (_, e: Effort) => (
        e.mmp,
        data.filter {
          case tp: Trackpoint => tp.activity.getOrElse("UNKNOWN") == e.activity &&
            tp.timestamp <= e.endTimeStamp && tp.timestamp >= e.startTimeStamp
        }.collect
      )
    }
}
```



```
def bestsForPeriod(data: RDD[Trackpoint], period: Int,
                  app: SLP, model: Broadcast[KMeansModel]) = {
  val clusteredMMPs = /* map from cluster pairs to effort summary structures */

  clusteredMMPs
    .reduceByKey ((a, b) => if (a.mmp > b.mmp) a else b)
    .takeOrdered(20)(Ordering.by[((Int, Int), Effort), Double] {
      case (_, e: Effort) => -e.mmp
    })
    .map {
      case (_, e: Effort) => (
        e.mmp,
        data.filter {
          case tp: Trackpoint => tp.activity.getOrElse("UNKNOWN") == e.activity &&
            tp.timestamp <= e.endTimeStamp && tp.timestamp >= e.startTimeStamp
        }.collect
      )
    }
}
```



```
def bestsForPeriod(data: RDD[Trackpoint], period: Int,
                  app: SLP, model: Broadcast[KMeansModel]) = {
  val clusteredMMPs = /* map from cluster pairs to effort summary structures */

  clusteredMMPs
    .reduceByKey ((a, b) => if (a.mmp > b.mmp) a else b)
    .takeOrdered(20)(Ordering.by[((Int, Int), Effort), Double] {
      case (_, e: Effort) => -e.mmp
    })
    .map {
      case (_, e: Effort) => (
        e.mmp,
        data.filter {
          case tp: Trackpoint => tp.activity.getOrElse("UNKNOWN") == e.activity &&
            tp.timestamp <= e.endTimeStamp && tp.timestamp >= e.startTimeStamp
        }.collect
      )
    }
}
```

```
def bestsForPeriod(data: RDD[Trackpoint], period: Int,
                  app: SLP, model: Broadcast[KMeansModel]) = {
  val clusteredMMPs = /* map from cluster pairs to effort summary structures */

  clusteredMMPs
    .reduceByKey ((a, b) => if (a.mmp > b.mmp) a else b)
    .takeOrdered(20)(Ordering.by[((Int, Int), Effort), Double] {
      case (_, e: Effort) => -e.mmp
    })
    .map {
      case (_, e: Effort) => (
        e.mmp,
        data.filter {
          case tp: Trackpoint => tp.activity.getOrElse("UNKNOWN") == e.activity &&
            tp.timestamp <= e.endTime && tp.timestamp >= e.startTime
        }.collect
      )
    }
}
```



```
def bestsForPeriod(data: RDD[Trackpoint], period: Int,
                  app: SLP, model: Broadcast[KMeansModel]) = {
  val clusteredMMPs = /* map from cluster pairs to effort summary structures */

  clusteredMMPs
    .reduceByKey ((a, b) => if (a.mmp > b.mmp) a else b)
    .takeOrdered(20)(Ordering.by[((Int, Int), Effort), Double] {
      case (_, e: Effort) => -e.mmp
    })
    .map {
      case (_, e: Effort) => (
        e.mmp,
        data.filter {
          case tp: Trackpoint => tp.activity.getOrElse("UNKNOWN") == e.activity &&
            tp.timestamp <= e.endTimeStamp && tp.timestamp >= e.startTimeStamp
        }.collect
      )
    }
}
```

Conclusions



**Avoid shuffles
when possible**



**Avoid shuffles
when possible**



**Broadcast large
static data**



**Avoid shuffles
when possible**



**Broadcast large
static data**



**Cache only
when necessary**



**Avoid shuffles
when possible**



**Broadcast large
static data**



**Cache only
when necessary**



**Embrace
laziness**

(ONLY PAY FOR WHAT YOU USE)



Avoid shuffles
when possible



Broadcast large
static data

WORK WITH SPARK,
NOT AGAINST IT



Cache only
when necessary



Embrace
laziness

(ONLY PAY FOR WHAT YOU USE)

Thanks!

willb@redhat.com

@willb

<http://chapeau.freevariable.com>