

Basic Principles of Intersection Signalization

In Chapter 16, various options for intersection control were presented and discussed. Warrants for implementation of traffic control signals at an intersection, presented in the *Manual on Uniform Traffic Control Devices* [1], provide general and specific criteria for selection of an appropriate form of intersection control. At many intersections, the combination of traffic volumes, potential conflicts, overall safety of operation, efficiency of operation, and driver convenience will lead to a decision to install traffic control signals.

The operation of signalized intersections is often complex, involving competing vehicular and pedestrian movements. Appropriate methodologies for design and timing of signals and for the operational analysis of signalized intersections require that the behavior of drivers and pedestrians at a signalized intersection be modeled in a form that can be easily manipulated and optimized. This chapter discusses some of the fundamental operational characteristics at a signalized intersection and the ways in which they may be effectively modeled.

In Chapter 18, these principles are applied to a signalized intersection design and timing process. In Chapter 21, they are augmented and combined into an

overall model of signalized intersection operations. The particular model presented is that of the *Highway Capacity Manual* [2].

This chapter focuses on four critical aspects of signalized intersection operation:

1. Discharge headways, saturation flow rates, and lost times
2. Allocation of time and the critical lane concept
3. The concept of left-turn equivalency
4. Delay as a measure of service quality

There are other aspects of signalized intersection operation that are also important, and the *Highway Capacity Manual* analysis model addresses many of them. These four, however, are central to understanding traffic behavior at signalized intersections and are highlighted here.

17.1 Terms and Definitions

Traffic signals are complex devices that can operate in a variety of different modes. A number of key terms and

17.1 TERMS AND DEFINITIONS

definitions should be understood before pursuing a more substantive discussion.

17.1.1 Components of a Signal Cycle

The following terms describe portions and subportions of a signal cycle. The most fundamental unit in signal design and timing is the *cycle*, as defined below.

1. *Cycle*. A signal cycle is one complete rotation through all of the indications provided. In general, every legal vehicular movement receives a "green" indication during each cycle, although there are some exceptions to this rule.
2. *Cycle length*. The cycle length is the time (in seconds) that it takes to complete one full cycle of indications. It is given the symbol " C ."
3. *Interval*. The interval is a period of time during which no signal indication changes. It is the smallest unit of time described within a signal cycle. There are several types of intervals within a signal cycle:

- (a) *Change interval*. The change interval is the "yellow" indication for a given movement. It is part of the transition from "green" to "red," in which movements about to lose "green" are given a "yellow" signal, while all other movements have a "red" signal. It is timed to allow a vehicle that cannot safely stop when the "green" is withdrawn to enter the intersection legally. The change interval is given the symbol " y_i " for movement(s) i .
- (b) *Clearance interval*. The clearance interval is also part of the transition from "green" to "red" for a given set of movements. During the clearance interval, all movements have a "red" signal. It is timed to allow a vehicle that legally enters the intersection on "yellow" to safely cross the intersection before conflicting flows are released. The clearance interval is given the symbol " ar_i " (for "all red") for movement(s) i .
- (c) *Green interval*. Each movement has one green interval during the signal cycle. During

a green interval, the movements permitted have a "green" light, while all other movements have a "red" light. The green interval is given the symbol " G_i " for movement(s) i .

- (d) *Red interval*. Each movement has a red interval during the signal cycle. All movements not permitted to move have a "green" light. In general, the red interval overlaps the green intervals for all other movements in the intersection. The red interval is given the symbol " R_i " for movement(s) i .
4. *Phase*. A signal phase consists of a green interval, plus the change and clearance intervals that follow it. It is a set of intervals that allows a designated movement or set of movements to flow and to be safely halted before release of a conflicting set of movements.

17.1.2 Types of Signal Operation

Traffic signals can operate on a pretimed basis or may be partially or fully actuated by arriving vehicles sensed by detectors. In networks, or on arterials, signals may be coordinated through computer control.

1. *Pretimed operation*. In pretimed operation, the cycle length, phase sequence, and timing of each interval are constant. Each cycle of the signal follows the same predetermined plan. "Multi-dial" controllers will allow different pretimed settings to be established. An internal clock is used to activate the appropriate timing. In such cases, it is typical to have at least an AM peak, a PM peak, and an off-peak signal timing.
2. *Semi-actuated operation*. In semi-actuated operation, detectors are placed on the minor approach(es) to the intersection; there are no detectors on the major street. The light is green for the major street at all times except when a "call" or actuation is noted on one of the minor approaches. Then, subject to limitations such as a minimum major-street green, the green is transferred to the minor street. The green returns

to the major street when the maximum minor-street green is reached or when the detector senses that there is no further demand on the minor street. Semi-actuated operation is often used where the primary reason for signalization is "interruption of continuous traffic," as discussed in Chapter 16.

3. **Full actuated operation.** In full actuated operation, every lane of every approach must be monitored by a detector. Green time is allocated in accordance with information from detectors and programmed "rules" established in the controller for capturing and retaining the green. In full actuated operation, the cycle length, sequence of phases, and green time split may vary from cycle to cycle. Chapter 20 presents more detailed descriptions of actuated signal operation, along with a methodology for timing such signals.
4. **Computer control.** Computer control is a system term. No individual signal is "computer controlled," unless the signal controller is considered to be a computer. In a computer-controlled system, the computer acts as a master controller, coordinating the timings of a large number (hundreds) of signals. The computer selects or calculates an optimal coordination plan based on input from detectors placed throughout the system. In general, such selections are made only once in advance of an AM or PM peak period. The nature of a system transition from one timing plan to another is sufficiently disruptive to be avoided during peak-demand periods. Individual signals in a computer-controlled system generally operate in the pretimed mode. For coordination to be effective, all signals in the network must use the same cycle length (or an even multiple thereof), and it is therefore difficult to maintain a progressive pattern where cycle length or phase splits are allowed to vary.

17.1.3 Treatment of Left Turns

The modeling of signalized intersection operation would be straightforward if left turns did not exist. Left

turns at a signalized intersection can be handled in one of three ways:

1. **Permitted left turns.** A "permitted" left turn movement is one that is made across an opposing flow of vehicles. The driver is permitted to cross through the opposing flow, but must select an appropriate gap in the opposing traffic stream through which to turn. This is the most common form of left-turn phasing at signalized intersections, used where left-turn volumes are reasonable and where gaps in the opposing flow are adequate to accommodate left turns safely.
2. **Protected left turns.** A "protected" left turn movement is made without an opposing vehicular flow. The signal plan protects left-turning vehicles by stopping the opposing through movement. This requires that the left turns and the opposing through flow be accommodated in separate signal phases and leads to multiphase (more than two) signalization. In some cases, left turns are "protected" by geometry or regulation. Left turns from the stem of a T-intersection, for example, face no opposing flow, as there is no opposing approach to the intersection. Left turns from a one-way street similarly do not face an opposing flow.
3. **Compound left turns.** More complicated signal timing can be designed in which left turns are protected for a portion of the signal cycle and are permitted in another portion of the cycle. Protected and permitted portions of the cycle can be provided in any order. Such phasing is also referred to as *protected plus permitted* or *permitted plus protected*, depending upon the order of the sequence.

The permitted left turn movement is very complex. It involves the conflict between a left turn and an opposing through movement. The operation is affected by the left-turn flow rate and the opposing flow rate, the number of opposing lanes, whether left turns flow from an exclusive left-turn lane or from a shared lane, and the details of the signal timing. Modeling the interaction

among these elements is a complicated process, one that often involves iterative elements.

The terms *protected* and *permitted* may also be applied to right turns. In this case, however, the conflict is between the right-turn vehicular movement and the pedestrian movement in the conflicting crosswalk. The vast majority of right turns at signalized intersections are handled on a permitted basis. Protected right turns generally occur at locations where there are overpasses or underpasses provided for pedestrians. At these locations, pedestrians are prohibited from making surface crossings; barriers are often required to enforce such a prohibition.

17.2 Discharge Headways, Saturation Flow, Lost Times, and Capacity

The fundamental element of a signalized intersection is the periodic stopping and restarting of the traffic stream. Figure 17.1 illustrates this process. When the light turns GREEN, there is a queue of stored vehicles that were stopped during the preceding RED phase, waiting to be discharged. As the queue of vehicles moves, headway measurements are taken as follows:

- The first headway is the time lapse between the initiation of the GREEN signal and the time that the front wheels of the first vehicle cross the stop line.
- The second headway is the time lapse between the time that the first vehicle's front wheels cross the stop line and the time that the second vehicle's front wheels cross the stop line.
- Subsequent headways are similarly measured.
- Only headways through the last vehicle in queue (at the initiation of the GREEN light) are considered to be operating under "saturated" conditions.

If many queues of vehicles are observed at a given location and the average headway is plotted vs. the queue position of the vehicle, a trend similar to that shown in Figure 17.1 (b) emerges.

The first headway is relatively long. The first driver must go through the full perception-reaction sequence, move his or her foot from the brake to the accelerator, and accelerate through the intersection. The second headway is shorter, because the second driver can overlap the perception-reaction and acceleration process of the first driver. Each successive headway is a little bit smaller. Eventually, the headways tend to level out. This generally occurs when queued vehicles have fully accelerated by the time they cross the stop line. At this point, a stable moving queue has been established.

17.2.1 Saturation Headway and Saturation Flow Rate

As noted, average headways will tend towards a constant value. In general, this occurs from the fourth or fifth headway position. The constant headway achieved is referred to as the *saturation headway*, as it is the average headway that can be achieved by a saturated, stable moving queue of vehicles passing through the signal. It is given the symbol "*h*," in units of seconds/vehicle.

It is convenient to model behavior at a signalized intersection by assuming that every vehicle (in a given lane) consumes an average of "*h*" seconds of green time to enter the intersection. If every vehicle consumes "*h*" seconds of green time and if the signal were *always* green, then "*s*" vehicles per hour could enter the intersection. This is referred to as the *saturation flow rate*:

$$s = \frac{3,600}{h} \quad (17-1)$$

where: *s* = saturation flow rate, vehicles per hour of green per lane (veh/hg/ln)

h = saturation headway, seconds/vehicle (s/veh)

Saturation flow rate can be multiplied by the number of lanes provided for a given set of movements to obtain a saturation flow rate for a lane group or approach.

The saturation flow rate is, in effect, the capacity of the approach lane or lanes if they were available for use all of the time (i.e., if the signal were always GREEN). The signal, of course, is not always GREEN for any given movement. Thus, some mechanism (or

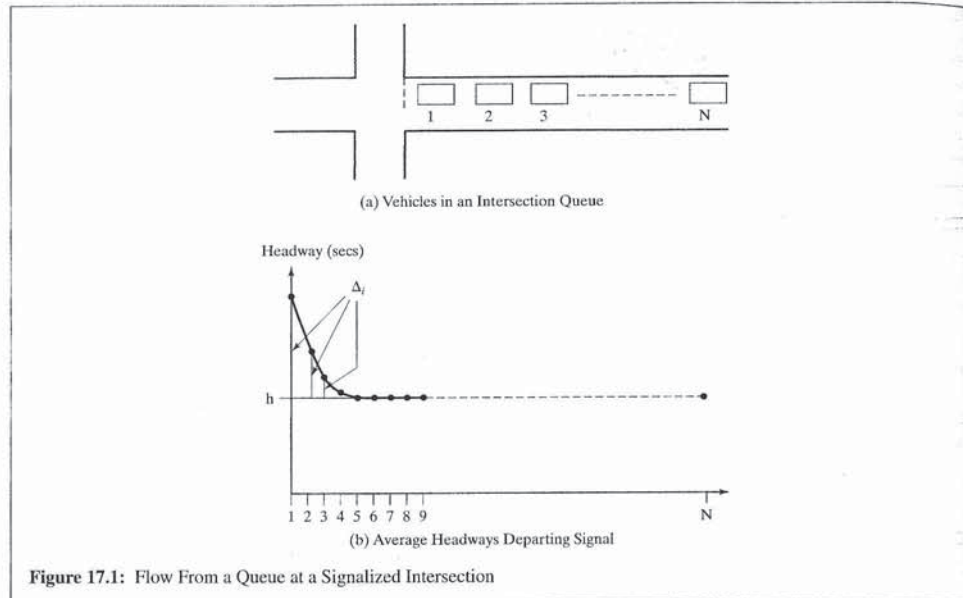


Figure 17.1: Flow From a Queue at a Signalized Intersection

model) for dealing with the cyclic starting and stopping of movements must be developed.

17.2.2 Start-Up Lost Time

The average headway per vehicle is actually greater than " h " seconds. The first several headways are, in fact, larger than " h " seconds, as illustrated in Exhibit 17-1 (b). The first three or four headways involve additional time as drivers react to the GREEN signal and accelerate. The additional time involved in each of these initial headways (above and beyond " h " seconds) is noted by the symbol Δ_i (for headway i). These additional times are added, and are referred to as the *start-up lost time*:

$$\ell_1 = \sum_i \Delta_i \quad (17-2)$$

where: ℓ_1 = start-up lost time, s/phase
 Δ_i = incremental headway (above " h " seconds) for vehicle i , s

Thus, it is possible to model the amount of GREEN time required to discharge a queue of " n " vehicles as:

$$T_n = \ell_1 + nh \quad (17-3)$$

where: T_n = GREEN time required to move queue of " n " vehicles through a signalized intersection, s

ℓ_1 = start-up lost time, s/phase
 n = number of vehicles in queue
 h = saturation headway, s/veh

While this particular model is not of great use, it does illustrate the basic concepts of saturation headway and start-up lost times. The start-up lost time is thought of as a period of time that is "lost" to vehicle use. Remaining

GREEN time, however, may be assumed to be usable at a rate of h s/veh.

17.2.3 Clearance Lost Time

The start-up lost time occurs every time a queue of vehicles starts moving on a GREEN signal. There is also a lost time associated with stopping the queue at the end of the GREEN signal. This time is more difficult to observe in the field, as it requires that the standing queue of vehicles be large enough to consume all of the GREEN time provided. In such a situation, the clearance lost time, ℓ_2 , is defined as the time interval between the last vehicle's front wheels crossing the stop line, and the initiation of the GREEN for the next phase. The clearance lost time occurs each time a flow of vehicles is stopped.

17.2.4 Total Lost Time and the Concept of Effective GREEN Time

If the start-up lost time occurs each time a queue starts to move and the clearance lost time occurs each time the flow of vehicles stops, then for each GREEN phase:

$$t_L = \ell_1 + \ell_2 \quad (17-4)$$

where: t_L = total lost time per phase, s/phase

All other variables are as previously defined.

The concept of lost times leads to the concept of *effective green time*. The actual signal goes through a sequence of intervals for each signal phase:

- Green
- Yellow
- All-red
- Red

The "yellow" and "all-red" intervals are a transition between GREEN and RED. This must be provided because vehicles cannot stop instantaneously when the light changes. The "all-red" is a period of time during which all lights in all directions are red. During the RED interval for one set of movements, another set of movements goes

through the green, yellow, and all-red intervals. These intervals are defined more precisely in Chapter 18.

In terms of modeling, there are really only two time periods of interest: *effective green time* and *effective red time*. For any given set of movements, *effective green time* is the amount of time that vehicles are moving (at a rate of one vehicle every h seconds). The *effective red time* is the amount of time that they are not moving. Effective green time is related to actual green time as follows:

$$g_i = G_i + Y_i - t_{Li} \quad (17-5)$$

where: g_i = effective green time for movement(s) i , s

G_i = actual green time for movement(s) i , s

Y_i = sum of yellow and all red intervals for movement(s) i , s, ($Y_i = y_i + ar_i$)

y_i = yellow interval for movement(s) i , s

ar_i = all-red interval for movement(s) i , s

t_{Li} = total lost time for movement(s) i , s

This model results in an effective green time that may be fully utilized by vehicles at the saturation flow rate (i.e., at an average headway of h s/veh).

17.2.5 Capacity of an Intersection Lane or Lane Group

The saturation flow rate(s) represents the capacity of an intersection lane or lane group assuming that the light is always GREEN. The portion of real time that is effective green is defined by the "green ratio," the ratio of the effective green time to the cycle length of the signal (g/C). The capacity of an intersection lane or lane group may then be computed as:

$$c_i = s_i * \left(\frac{g_i}{C} \right) \quad (17-6)$$

where: c_i = capacity of lane or lane group i , veh/h

s_i = saturation flow rate for lane or lane group i , veh/hg

g_i = effective green time for lane or lane group i , s

C = signal cycle length, s

A Sample Problem

These concepts are best illustrated using a sample problem. Consider a given movement at a signalized intersection with the following known characteristics:

- Cycle length, $C = 60$ s
- Green time, $G = 27$ s
- Yellow plus all-red time, $Y = 3$ s
- Saturation headway, $h = 2.4$ s/veh
- Start-up lost time, $\ell_1 = 2.0$ s
- Clearance lost time, $\ell_2 = 1.0$ s

For these characteristics, what is the capacity (per lane) for this movement?

The problem will be approached in two different ways. In the first, a ledger of time within the hour is created. Once the amount of time per hour used by vehicles at the saturation flow rate is established, capacity can be found by assuming that this time is used at a rate of one vehicle every h seconds. Since the characteristics stated are given on a *per phase* basis, these would have to be converted to a *per hour* basis. This is easily done knowing the number of signal cycles that occur within an hour. For a 60-s cycle, there are $3,600/60 = 60$ cycles within the hour. The subject movements will have one GREEN phase in each of these cycles. Then:

- Time in hour: 3,600 s
- RED time in hour: $(60 - 27 - 3) \times 60 = 1,800$ s
- Lost time in hour: $(2 + 1) \times 60 = 180$ s
- Remaining time in hour: $3,600 - 1,800 - 180 = 1,620$ s

The 1,620 remaining seconds of time in the hour represent the amount of time that can be used at a rate of one vehicle every h seconds, where $h = 2.4$ s/veh in this case. This number was calculated by deducting the periods during which no vehicles (in the subject movements) are effectively moving. These periods include the RED time as well as the start-up and clearance lost times in each signal cycle. The capacity of this movement may then be computed as:

$$c = \frac{1620}{2.4} = 675 \text{ veh/h/ln}$$

A second approach to this problem utilizes Equation 17-6, with the following values:

$$s = \frac{3,600}{h} = \frac{3,600}{2.4} = 1,500 \text{ veh/hg/ln}$$

$$g = G + Y - t_L = 27 + 3 - 3 = 27 \text{ s}$$

$$c = s \left(\frac{g}{C} \right) = 1,500 \left(\frac{27}{60} \right) = 675 \text{ veh/h/ln}$$

The two results are, as expected, the same. Capacity is found by isolating the effective green time available to the subject movements and by assuming that this time is used at the saturation flow rate (or headway).

17.2.6 Notable Studies on Saturation Headways, Flow Rates, and Lost Times

For purposes of illustrating basic concepts, subsequent sections of this chapter will assume that the value of saturation flow rate (or headway) is known. In reality, the saturation flow rate varies widely with a variety of prevailing conditions, including lane widths, heavy-vehicle presence, approach grades, parking conditions near the intersection, transit bus presence, vehicular and pedestrian flow rates, and other conditions.

The first significant studies of saturation flow were conducted by Bruce Greenshields in the 1940s [3]. His studies resulted in an average saturation flow rate of 1,714 veh/hg/ln and a start-up lost time of 3.7 s. The study, however, covered a variety of intersections with varying underlying characteristics. A later study in 1978 [4] reexamined the Greenshields hypothesis; it resulted in the same saturation flow rate (1,714 veh/hg/ln) but a lower start-up lost time of 1.1 s. The latter study had data from 175 intersections, covering a wide range of underlying characteristics.

A comprehensive study of saturation flow rates at intersections in five cities was conducted in 1987–1988 [5] to determine the effect of opposed left turns. It also produced, however, a good deal of data on saturation flow rates in general. Some of the results are summarized in Table 17.1.

These results show generally lower saturation flow rates (and higher saturation headways) than previous studies. The data, however, reflect the impact of opposed

Table 17.1: Saturation Flow Rates from a Nationwide Survey

Item	Single-Lane Approaches	Two-Lane Approaches
Number of Approaches	14	26
Number of 15-Minute Periods	101	156
Saturation Flow Rates		
Average	1,280 veh/hg/ln	1,337 veh/hg/ln
Minimum	636 veh/hg/ln	748 veh/hg/ln
Maximum	1,705 veh/hg/ln	1,969 veh/hg/ln
Saturation Headways		
Average	2.81 s/veh	2.69 s/veh
Minimum	2.11 s/veh	1.83 s/veh
Maximum	5.66 s/veh	4.81 s/veh

left turns, truck presence, and a number of other “non-standard” conditions, all of which have a significant impeding effect. The most remarkable result of this study, however, was the wide variation in measured saturation flow rates, both over time at the same site and from location to location. Even when underlying conditions remained fairly constant, the variation in observed saturation flow rates at a given location was as large as 20%–25%. In a doctoral dissertation using the same data, Prassas demonstrated that saturation headways and flow rates have a significant stochastic component, making calibration of stable values difficult [6].

The study also isolated saturation flow rates for “ideal” conditions, which include all passenger cars, no turns, level grade, and 12-ft lanes. Even under these conditions, saturation flow rates varied from 1,240 pc/hg/ln to 2,092 pc/hg/ln for single-lane approaches, and from 1,668 pc/hg/ln to 2,361 pc/hg/ln for multilane approaches. The difference between observed saturation flow rates at single and multilane approaches is also interesting. Single-lane approaches have a number of unique characteristics that are addressed in the *Highway Capacity Manual* model for analysis of signalized intersections (see Chapters 21 and 22).

Current standards in the *Highway Capacity Manual* [1] use an ideal saturation flow rate of 1,900 pc/hg/ln for both single and multilane approaches. This ideal rate is then adjusted for a variety of prevailing conditions. The manual also provides default values for lost times. The default value for start-up lost time (ℓ_1) is 2.0 s. For the

clearance lost time (ℓ_2), the default value varies with the “yellow” and “all-red” timings of the signal:

$$\ell_2 = y + ar - e \quad (17-7)$$

where: ℓ_2 = clearance lost time, s

y = length of yellow interval, s

ar = length of all-red interval, s

e = encroachment of vehicles into yellow and all-red, s

A default value of 2.0 s is used for e .

17.3 The Critical-Lane and Time-Budget Concepts

In signal analysis and design, the “critical-lane” and “time budget” concepts are closely related. The time budget, in its simplest form, is the allocation of time to various vehicular and pedestrian movements at an intersection through signal control. Time is a constant: there are always 3,600 seconds in an hour, and all of them must be allocated. In any given hour, time is “budgeted” to legal vehicular and pedestrian movements and to lost times.

The “critical-lane” concept involves the identification of specific lane movements that will control the timing of a given signal phase. Consider the situation illustrated in Figure 17.2. A simple two-phase signal

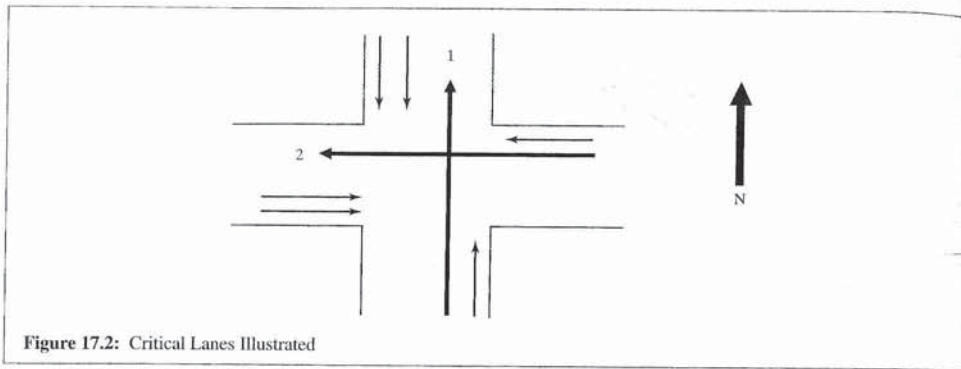


Figure 17.2: Critical Lanes Illustrated

controls the intersection. Thus, all E-W movements are permitted during one phase, and all N-S movements are permitted in another phase. During each of these phases, there are four lanes of traffic (two in each direction) moving simultaneously. Demand is not evenly distributed among them; one of these lanes will have the most intense traffic demand. The signal must be timed to accommodate traffic in this lane—the “critical lane” for the phase.

In the illustration of Figure 17.2, the signal timing and design must accommodate the total demand flows in lanes 1 and 2. As these lanes have the most intense demand, if the signal accommodates them, all other lanes will be accommodated as well. Note that the critical lane is identified as the lane with the most intense traffic demand, not the lane with the highest volume. This is because there are many variables affecting traffic flow. A lane with many left-turning vehicles, for example, may require more time than an adjacent lane with no turning vehicles, but a higher volume. Determining the intensity of traffic demand in a lane involves accounting for prevailing conditions that may affect flow in that particular lane.

In establishing a time budget for the intersection of Figure 17.2, time would have to be allocated to four elements:

- Movement of vehicles in critical lane 1
- Movement of vehicles in critical lane 2

- Start-up and clearance lost times for vehicles in critical lane 1
- Start-up and clearance lost times for vehicles in critical lane 2

This can be thought of in the following way: lost times are not used by any vehicle. When deducted from total time, remaining time is effective green time and is allocated to critical-lane demands—in this case, in lanes 1 and 2. The total amount of effective green time, therefore, must be sufficient to accommodate the total demand in lanes 1 and 2 (the critical lanes). These critical demands must be accommodated one vehicle at a time, as they cannot move simultaneously.

The example of Figure 17.2 is a relatively simple case. In general, the following rules apply to the identification of critical lanes:

- There is a critical lane and a critical-lane flow for each discrete signal phase provided.
- Except for lost times, when no vehicles move, there must be one and only one critical lane moving during every second of effective green time in the signal cycle.
- Where there are overlapping phases, the potential combination of lane flows yielding the highest sum of critical lane flows while preserving the requirement of item (b) identifies critical lanes.

Chapter 20 contains a detailed discussion of how to identify critical lanes for any signal timing and design.

17.3.1 The Maximum Sum of Critical-Lane Volumes: One View of Signalized Intersection Capacity

It is possible to consider the maximum possible sum of critical-lane volumes to be a general measure of the “capacity” of the intersection. This is not the same as the traditional view of capacity presented in the *Highway Capacity Manual*, but it is a useful concept to pursue.

By definition, each signal phase has one and only one critical lane. Except for lost times in the cycle, one critical lane is always moving. Lost times occur for each signal phase and represent time during which no vehicles in any lane are moving. The maximum sum of critical lane volumes may, therefore, be found by determining how much total lost time exists in the hour. The remaining time (total effective green time) may then be divided by the saturation headway.

To simplify this derivation, it is assumed that the total lost time per phase (t_L) is a constant for all phases. Then, the total lost time per signal cycle is:

$$L = N * t_L \quad (17-8)$$

where: L = lost time per cycle, s/cycle

t_L = total lost time per phase (sum of $\ell_1 + \ell_2$), s/phase

N = number of phases in the cycle

The total lost time in an hour depends upon the number of cycles occurring in the hour:

$$L_H = L * (3,600/C) \quad (17-9)$$

where: L_H = lost time per hour, s/hr

L = lost time per cycle, s/cycle

C = cycle length, s

The remaining time within the hour is devoted to effective green time for critical lane movements:

$$T_G = 3,600 - L_H \quad (17-10)$$

where: T_G = total effective green time in the hour, s

This time may be used at a rate of one vehicle every h seconds, where h is the saturation headway:

$$V_c = (T_G/h) \quad (17-11)$$

where: V_c = maximum sum of critical lane volumes, veh/h

h = saturation headway, s/veh

Merging Equations 17-8 through 17-11, the following relationship emerges:

$$V_c = \frac{1}{h} \left[3,600 - N t_L \left(\frac{3,600}{C} \right) \right] \quad (17-12)$$

where all variables are as previously defined.

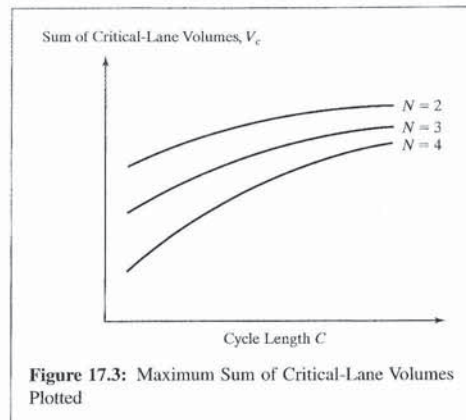
Consider the example of Figure 17.2 again. If the signal at this location has two phases, a cycle length of 60 seconds, total lost times of 4 s/phase, and a saturation headway of 2.5 s/veh, the maximum sum of critical lane flows (the sum of flows in lanes 1 and 2) is:

$$V_c = \frac{1}{2.5} \left[3,600 - 2 * 4 * \left(\frac{3,600}{60} \right) \right] = 1,248 \text{ veh/h}$$

The equation indicates that there are 3,600/60 = 60 cycles in an hour. For each of these, 2 * 4 = 8 s of lost time is experienced, for a total of 8 * 60 = 480 s in the hour. The remaining 3,600 - 480 = 3,120 s may be used at a rate of one vehicle every 2.5 s.

If Equation 17-12 is plotted, an interesting relationship between the maximum sum of critical lane volumes (V_c), cycle length (C), and number of phases (N) may be observed, as illustrated in Figure 17.3.

As the cycle length increases, the “capacity” of the intersection also increases. This is because of lost times, which are constant per cycle. The longer the cycle length, the fewer cycles there are in an hour. This leads to less lost time in the hour, more effective green time in the hour, and a higher sum of critical-lane volumes. Note, however, that the relationship gets flatter as cycle length increases. As a general rule, increasing the cycle length may result in small increases in capacity. On the other hand, capacity can rarely be increased significantly by only increasing the cycle length. Other measures, such as adding lanes, are often also necessary.



Capacity also decreases as the number of phases increases. This is because for each phase, there is one full set of lost times in the cycle. Thus, a two-phase signal has only two sets of lost times in the cycle, while a three-phase signal has three.

These trends provide insight, but also raise an interesting question: Given these trends, it appears that all signals should have two phases and that the maximum practical cycle length should be used in all cases. After all, this combination would, apparently, yield the highest "capacity" for the intersection.

Using the maximum cycle length is not practical unless truly needed. Having a cycle length that is considerably longer than what is desirable causes increases in delay to drivers and passengers. The increase in delay is because there will be times when vehicles on one approach are waiting for the green while there is no demand on conflicting approaches. Shorter cycle lengths yield less delay. Further, there is no incentive to maximize the cycle length. There will always be 3,600 seconds in the hour, and increasing the cycle length to accommodate increasing demand over time is quite simple, requiring only a resetting of the local signal controller. The shortest cycle length consistent with a v/c ratio in the range of 0.80–0.95 is generally used to produce optimal delays. Thus, the view of signal capacity is quite different from that of pavement capacity. When

deciding on the number of lanes on a freeway (or on an intersection approach), it is desirable to build excess capacity (i.e., achieve a low v/c ratio). This is because once built, it is unlikely that engineers will get an opportunity to expand the facility for 20 or more years, and adjacent land development may make such expansion impossible. The 3,600 seconds in an hour, however, are immutable, and retiming the signal to allocate more of them to effective green time is a simple task requiring no field construction.

17.3.2 Finding an Appropriate Cycle Length

If it is assumed that the demands on an intersection are known and that the critical lanes can be identified, then Equation 17-12 could be solved using a known value of V_c to find a minimum acceptable cycle length:

$$C_{\min} = \frac{Nt_L}{1 - \left(\frac{V_c}{3,600/h} \right)} \quad (17-13)$$

Thus, if in the example of Figure 17.2, the actual sum of critical-lane volumes was determined to be 1,000 veh/h, the minimum feasible cycle length would be:

$$C_{\min} = \frac{2 \cdot 4}{1 - \left(\frac{1,000}{3,600/2.5} \right)} = \frac{8}{0.3056} = 26.2 \text{ s}$$

The cycle length could be reduced, in this case, from the given 60 s to 30 s (the effective minimum cycle length used). This computation, however, assumes that the demand (V_c) is uniformly distributed throughout the hour and that every second of effective green time will be used. Neither of these assumptions is very practical. In general, signals would be timed for the flow rates occurring in the peak 15 minutes of the hour. Equation 17-13 could be modified by dividing V_c by a known peak-hour factor (PHF) to estimate the flow rate in the worst 15-minute period of the hour. Similarly, most signals would be timed to have somewhere between 80% and 95% of the available capacity actually used. Due to the

normal stochastic variations in demand on a cycle-by-cycle and daily basis, some excess capacity must be provided to avoid failure of individual cycles or peak periods on a specific day. If demand, V_c , is also divided by the expected utilization of capacity (expressed in decimal form), then this is also accommodated. Introducing these changes transforms Equation 17-13 to:

$$C_{\text{des}} = \frac{Nt_L}{1 - \left[\frac{V_c}{(3,600/h) \cdot \text{PHF} \cdot (v/c)} \right]} \quad (17-14)$$

where: C_{des} = desirable cycle length, s

PHF = peak hour factor

v/c = desired volume to capacity ratio

All other variables are as previously defined.

Returning to the example, if the PHF is 0.95 and it is desired to use no more than 90% of available capacity during the peak 15-minute period of the hour, then:

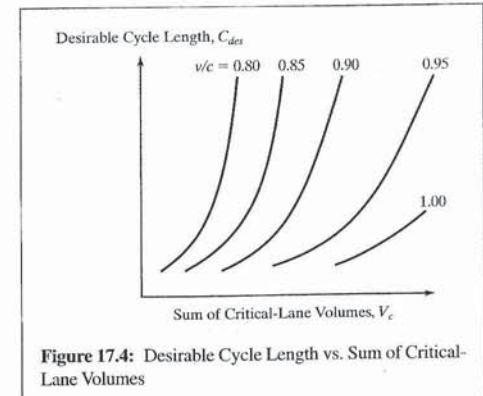
$$C_{\text{des}} = \frac{2 \cdot 4}{1 - \left[\frac{1,000}{(3,600/2.5) \cdot 0.95 \cdot 0.90} \right]} = \frac{8}{0.188} = 42.6 \text{ s}$$

In practical terms, this would lead to the use of a 45-second cycle length.

The relationship between a desirable cycle length, the sum of critical-lane volumes, and the target v/c ratio is quite interesting and is illustrated in Figure 17.4.

Figure 17.4 illustrates a typical relationship for a specified number of phases, saturation headway, lost times, and peak-hour factor. If a vertical is drawn at any specified value of V_c (sum of critical lane volumes), it is clear that the resulting cycle length is very sensitive to the target v/c ratio. As the curves for each v/c ratio are eventually asymptotic to the vertical, it is not always possible to achieve a specified v/c ratio.

Consider the case of a three-phase signal, with $t_L = 4$ s/phase, a saturation headway of 2.2 s/veh, a PHF of 0.90 and $V_c = 1,200$ veh/h. Desirable cycle lengths will be computed for a range of target v/c ratios varying from 1.00 to 0.80.



$$C_{\text{des}} = \frac{3 \cdot 4}{1 - \left[\frac{1,200}{(3,600/2.2) \cdot 0.90 \cdot 1.00} \right]} = \frac{12}{0.1852} = 64.8 \Rightarrow 65 \text{ s}$$

$$C_{\text{des}} = \frac{3 \cdot 4}{1 - \left[\frac{1,200}{(3,600/2.2) \cdot 0.90 \cdot 0.95} \right]} = \frac{12}{0.1423} = 84.3 \Rightarrow 85 \text{ s}$$

$$C_{\text{des}} = \frac{3 \cdot 4}{1 - \left[\frac{1,200}{(3,600/2.2) \cdot 0.90 \cdot 0.90} \right]} = \frac{12}{0.0947} = 126.7 \Rightarrow 130 \text{ s}$$

$$C_{\text{des}} = \frac{3 \cdot 4}{1 - \left[\frac{1,200}{(3,600/2.2) \cdot 0.90 \cdot 0.85} \right]} = \frac{12}{0.0414} = 289.9 \Rightarrow 290 \text{ s}$$

$$C_{des} = \frac{3 \times 4}{1 - \left[\frac{1,200}{(3,600/2.2) \times 0.90 \times 0.80} \right]}$$

$$= \frac{12}{-0.0185} = -648.6 \text{ s}$$

For this case, reasonable cycle lengths can provide target v/c ratios of 1.00 or 0.95. Achieving v/c ratios of 0.90 or 0.85 would require long cycle lengths beyond the practical limit of 120 s for pretimed signals. The 130 s cycle needed to achieve a v/c ratio of 0.90 might be acceptable for an actuated signal location. However, a v/c ratio of 0.80 cannot be achieved under any circumstances. The negative cycle length that results signifies that there is not enough time within the hour to accommodate the demand with the required green time plus the 12 s of lost time per cycle. In effect, more than 3,600 s would have to be available to accomplish this.

A Sample Problem

Consider the intersection shown in Figure 17.5. The critical directional demands for this two-phase signal are shown with other key variables. Using the time-budget and critical-lane concepts, determine the number of lanes required for each of the critical movements and the minimum desirable cycle length that could be used. Note that an initial cycle length is specified, but will be modified as part of the analysis.

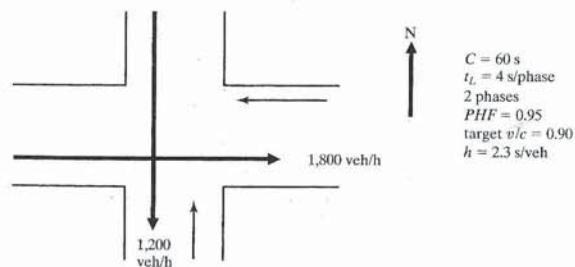


Figure 17.5: Sample Problem Using the Time-Budget and Critical-Lane Concepts

Assuming that the initial specification of a 60-s cycle is correct and given the other specified conditions, the maximum sum of critical lanes that can be accommodated is computed using Equation 17-12:

$$V_c = \frac{1}{2.3} \left[3,600 - 2 \times 4 \times \left(\frac{3,600}{60} \right) \right] = 1,357 \text{ veh/h}$$

The critical SB volume is 1,200 veh/h, and the critical EB volume is 1,800 veh/h. The number of lanes each must be divided into is now to be determined. Whatever combination is used, the sum of the critical-lane volumes for these two approaches must be below 1,357 veh/h. Figure 17.6 shows a number of possible lane combinations and the resulting sum of critical lane volumes. As can be seen from the scenarios of Figure 17.6, in order to have a sum of critical-lane volumes less than 1,357 veh/h, the SB approach must have at least two lanes, and the EB approach must have three lanes. Realizing that these demands probably reverse in the other peak hour (AM or PM), the N-S artery would probably require four lanes, and the E-W artery six lanes.

This is a very basic analysis, and it would have to be modified based on more specific information regarding individual movements, pedestrians, parking needs, and other factors.

If the final scenario is provided, V_c is only 1,200 veh/h. It is possible that the original cycle length of 60 s

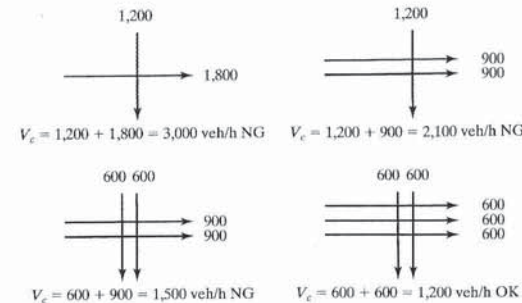


Figure 17.6: Possible Lane Scenarios for Sample Problem

could be reduced. A minimum desirable cycle length may be computed from Equation 17-14:

$$C_{des} = \frac{2 \times 4}{1 - \left[\frac{1,200}{(3,600/2.3) \times 0.95 \times 0.90} \right]}$$

$$= \frac{8}{0.103} = 77.7 \Rightarrow 80 \text{ s}$$

The resulting cycle length is larger than the original 60 s because the equation takes both the PHF and target v/c ratios into account. Equation 17-12 for computing the maximum value of V_c does not and assumes full use of capacity ($v/c = 1.00$) and no peaking within the hour. In essence, the 2×3 lane design proposal should be combined with an 80 s cycle length to achieve the desired results.

This problem illustrates the critical relationship between number of lanes and cycle lengths. Clearly, there are other scenarios that would produce desirable results. Additional lanes could be provided in either direction, which would allow the use of a shorter cycle length. Unfortunately, for many cases, signal timing is considered with a fixed design already in place. Only where right-of-way is available or a new intersection is being constructed can major changes in the number of lanes be considered. Allocation of lanes to various movements is also a consideration. Optimal solutions

are generally found more easily when the physical design and signalization can be treated in tandem.

If, in the problem of Figure 17.5, space limited both the EB and SB approaches to two lanes, the resulting V_c would be 1,500 veh/h. Would it be possible to accommodate this demand by lengthening the cycle length? Again, Equation 17-14 is used:

$$C_{des} = \frac{2 \times 4}{1 - \left[\frac{1,500}{(3,600/2.3) \times 0.95 \times 0.90} \right]}$$

$$= \frac{8}{-0.121} = -66.1 \text{ s NG}$$

The negative result indicates that there is no cycle length that can accommodate a V_c of 1,500 veh/h at this location.

17.4 The Concept of Left-Turn Equivalency

The most difficult process to model at a signalized intersection is the left turn. Left turns are made in several different modes using different design elements. Left turns may be made from a lane shared with through vehicles (shared-lane operation) or from a lane dedicated to left-turning vehicles (exclusive-lane operation). Traffic signals

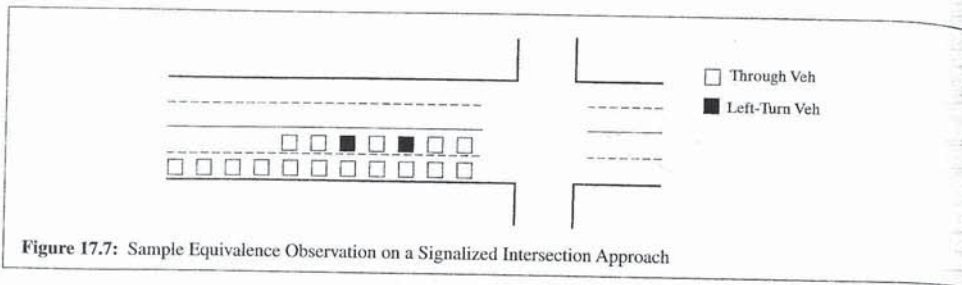


Figure 17.7: Sample Equivalence Observation on a Signalized Intersection Approach

may allow for permitted or protected left turns, or some combination of the two.

Whatever the case, however, a left-turning vehicle will consume more effective green time traversing the intersection than will a similar through vehicle. The most complex case is that of a permitted left turn made across an opposing vehicular flow from a shared lane. A left-turning vehicle in the shared lane must wait for an acceptable gap in the opposing flow. While waiting, the vehicle blocks the shared lane, and other vehicles (including through vehicles) in the lane are delayed behind it. Some vehicles will change lanes to avoid the delay, while others are unable to and must wait until the left-turner successfully completes the turn.

Many models of the signalized intersection account for this in terms of "through vehicle equivalents" (i.e., how many through vehicles would consume the same amount of effective green time traversing the stopline as one left-turning vehicle?). Consider the situation depicted in Figure 17.7. If both the left lane and the right lane were observed, an equivalence similar to the following statement could be determined:

In the same amount of time, the left lane discharges five through vehicles and two left-turning vehicles, while the right lane discharges eleven through vehicles.

In terms of effective green time consumed, this observation means that eleven through vehicles are equivalent to five through vehicles plus two left turning vehicles. If the left-turn equivalent is defined as E_{LT} :

$$11 = 5 + 2E_{LT}$$

$$E_{LT} = \frac{11 - 5}{2} = 3.0$$

It should be noted that this computation holds only for the prevailing characteristics of the approach during the observation period. The left-turn equivalent depends upon a number of factors, including how left turns are made (protected, permitted, compound), the opposing traffic flow, and the number of opposing lanes. Figure 17.8 illustrates the general form of the relationship for through vehicle equivalents of permitted left turns.

The left-turn equivalent, E_{LT} , increases as the opposing flow increases. For any given opposing flow, however, the equivalent decreases as the number of opposing lanes is increased from one to three. This latter relationship is not linear, as the task of selecting a gap through multilane opposing traffic is more difficult than selecting a gap through single-lane opposing traffic. Further, in a multilane traffic stream, vehicles do not pace each other side-by-side, and the gap distribution does not improve as much as the per-lane opposing flow decreases.

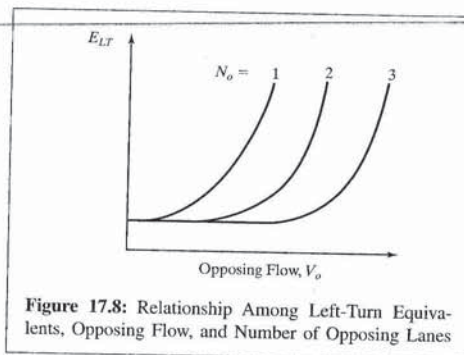


Figure 17.8: Relationship Among Left-Turn Equivalents, Opposing Flow, and Number of Opposing Lanes

To illustrate the use of left-turn equivalents in modeling, consider the following problem:

An approach to a signalized intersection has two lanes, permitted left-turn phasing, 10% left-turning vehicles, and a left-turn equivalent of 5.0. The saturation headway for through vehicles is 2.0 s/veh. Determine the equivalent saturation flow rate and headway for all vehicles on this approach.

The first way to interpret the left-turn equivalent is that each left-turning vehicle consumes 5.0 times the effective green time as a through vehicle. Thus, for the situation described, 10% of the traffic stream has a saturation headway of $2.0 \times 5.0 = 10.0$ s/veh, while the remainder (90%) have a saturation headway of 2.0 s/veh. The average saturation headway for all vehicles is, therefore:

$$h = (0.10 \times 10.0) + (0.90 \times 2.0) = 2.80 \text{ s/veh}$$

This corresponds to a saturation flow rate of:

$$s = \frac{3,600}{2.80} = 1,286 \text{ veh/hg/ln}$$

A number of models, including the *Highway Capacity Manual* approach, calibrate a multiplicative adjustment factor that converts an ideal (or through) saturation flow rate to a saturation flow rate for prevailing conditions:

$$s_{prev} = s_{ideal} * f_{LT}$$

$$f_{LT} = \frac{s_{prev}}{s_{ideal}} = \frac{(3,600/h_{prev})}{(3,600/h_{ideal})}$$

$$= h_{ideal}/h_{prev} \quad (17-15)$$

where: s_{prev} = saturation flow rate under prevailing conditions, veh/hg/ln

s_{ideal} = saturation flow rate under ideal conditions, veh/hg/ln

f_{LT} = left-turn adjustment factor

h_{ideal} = saturation headway under ideal conditions, s/veh

h_{prev} = saturation headway under prevailing conditions, s/veh

In effect, in the first solution, the prevailing headway, h_{prev} , was computed as follows:

$$h_{prev} = (P_{LT} E_{LT} h_{ideal}) + [(1 - P_{LT}) h_{ideal}] \quad (17-16)$$

Combining Equations 17-15 and 17-16:

$$f_{LT} = \frac{h_{ideal}}{(P_{LT} E_{LT} h_{ideal}) + [(1 - P_{LT}) h_{ideal}]}$$

$$f_{LT} = \frac{1}{P_{LT} E_{LT} + (1 - P_{LT})} = \frac{1}{1 + P_{LT}(E_{LT} - 1)} \quad (17-17)$$

The problem posed may now be solved using a left-turn adjustment factor. Note that the saturation headway under ideal conditions is $3,600/2.0 = 1,800$ veh/hg/ln. Then:

$$f_{LT} = \frac{1}{1 + 0.10(5 - 1)} = 0.714$$

$$s_{prev} = 1800 * 0.714 = 1,286 \text{ veh/hg/ln}$$

This, of course, is the same result.

It is important that the concept of left-turn equivalence be understood. Its use in multiplicative adjustment factors often obscures its intent and meaning. The fundamental concept, however, is unchanged—the equivalence is based on the fact that the effective green time consumed by a left-turning vehicle is E_{LT} times the effective green time consumed by a similar through vehicle.

Signalized intersection and other traffic models use other types of equivalents that are similar. Heavy-vehicle, local-bus, and right-turn equivalents have similar meanings and result in similar equations. Some of these have been discussed in previous chapters, and others will be discussed in subsequent chapters.

17.5 Delay as a Measure of Effectiveness

Signalized intersections represent point locations within a surface street network. As point locations, the measures of operational quality or effectiveness used for highway

sections are not relevant. Speed has no meaning at a point, and density requires a section of some length for measurement. A number of measures have been used to characterize the operational quality of a signalized intersection, the most common of which are:

- Delay
- Queuing
- Stops

These are all related. Delay refers to the amount of time consumed in traversing the intersection—the difference between the arrival time and the departure time, where these may be defined in a number of different ways. Queuing refers to the number of vehicles forced to queue behind the stop-line during a RED signal phase; common measures include the average queue length or a percentile queue length. Stops refer to the percentage or number of vehicles that must stop at the signal.

17.5.1 Types of Delay

The most common measure used is delay, with queuing often used as a secondary measure of operational quality. While it is possible to measure delay in the field, it is a difficult process, and different observers may make judgments that could yield different results. For many purposes, it is, therefore, convenient to have a predictive model for the estimate of delay. Delay, however, can be quantified in many different ways. The most frequently used forms of delay are defined as follows:

1. **Stopped-time delay.** Stopped-time delay is defined as the time a vehicle is stopped in queue while waiting to pass through the intersection; average stopped-time delay is the average for all vehicles during a specified time period.
2. **Approach delay.** Approach delay includes stopped-time delay but adds the time loss due to deceleration from the approach speed to a stop and the time loss due to reacceleration back to the desired speed. Average approach delay is the average for all vehicles during a specified time period.
3. **Time-in-queue delay.** Time-in-queue delay is the total time from a vehicle joining an intersection

queue to its discharge across the STOP line on departure. Again, average time-in-queue delay is the average for all vehicles during a specified time period.

4. **Travel time delay.** This is a more conceptual value. It is the difference between the driver's expected travel time through the intersection (or any roadway segment) and the actual time taken. Given the difficulty in establishing a "desired" travel time to traverse an intersection, this value is rarely used, other than as a philosophical concept.
5. **Control delay.** The concept of control delay was developed in the 1994 *Highway Capacity Manual*, and is included in the HCM 2000. It is the delay caused by a control device, either a traffic signal or a STOP-sign. It is approximately equal to time-in-queue delay plus the acceleration-deceleration delay component.

Figure 17.9 illustrates three of these delay types for a single vehicle approaching a RED signal.

Stopped-time delay for this vehicle includes only the time spent stopped at the signal. It begins when the vehicle is fully stopped and ends when the vehicle begins to accelerate. Approach delay includes additional time losses due to deceleration and acceleration. It is

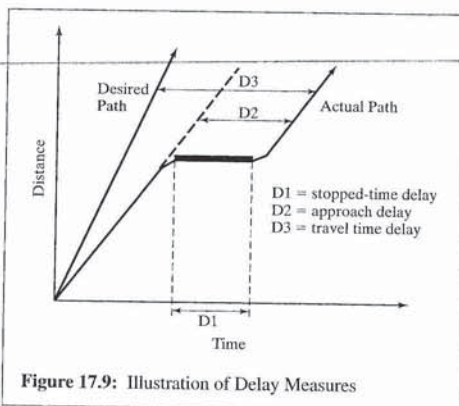


Figure 17.9: Illustration of Delay Measures

found by extending the velocity slope of the approach-vehicle as if no signal existed; the approach delay is the horizontal (time) difference between the hypothetical extension of the approaching velocity slope and the departure slope after full acceleration is achieved. Travel time delay is the difference in time between a hypothetical desired velocity line and the actual vehicle path. Time-in-queue delay cannot be effectively shown using one vehicle, as it involves joining and departing a queue of several vehicles.

Delay measures can be stated for a single vehicle, as an average for all vehicles over a specified time period, or as an aggregate total value for all vehicles over a specified time period. Aggregate delay is measured in total vehicle-seconds, vehicle-minutes, or vehicle-hours for all vehicles in the specified time interval. Average individual delay is generally stated in terms of s/veh for a specified time interval.

17.5.2 Basic Theoretical Models of Delay

Virtually all analytic models of delay begin with a plot of cumulative vehicles arriving and departing vs. time at a given signal location. The time axis is divided into periods of effective green and effective red as illustrated in Figure 17.10.

Vehicles are assumed to arrive at a uniform rate of flow of v vehicles per unit time, seconds in this case. This is shown by the constant slope of the arrival curve. Uniform arrivals assume that the inter-vehicle arrival time between vehicles is a constant. Thus, if the arrival flow rate, v , is 1,800 vehs/h, then one vehicle arrives every $3,600/1,800 = 2.0$ s.

Assuming no preexisting queue, vehicles arriving when the light is GREEN continue through the intersection, (i.e., the departure curve is the same as the arrival curve). When the light turns RED, however, vehicles continue to arrive, but none depart. Thus, the departure curve is parallel to the x-axis during the RED interval. When the next effective GREEN begins, vehicles queued during the RED interval depart from the intersection, now at the saturation flow rate, s , in veh/s. For stable operations, depicted here, the departure curve "catches up" with the arrival curve before the next RED interval begins (i.e., there is no residual or unserved queue left at the end of the effective GREEN).

This simple depiction of arrivals and departures at a signal allows the estimation of three critical parameters:

- The total time that any vehicle i spends waiting in the queue, $W(i)$, is given by the horizontal time-scale difference between the time of arrival and the time of departure.

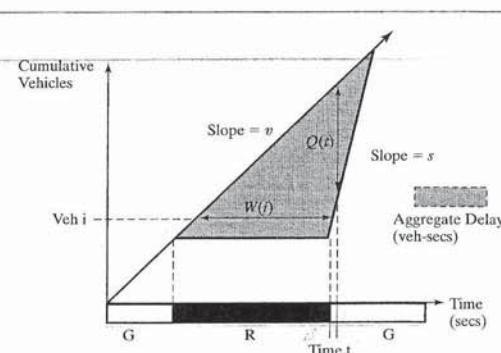


Figure 17.10: Delay, Waiting Time, and Queue Length Illustrated

- The total number of vehicles queued at any time t , $Q(t)$, is the vertical vehicle-scale difference between the number of vehicles that have arrived and the number of vehicles that have departed.
- The aggregate delay for all vehicles passing through the signal is the area between the arrival and departure curves (vehicles $\times$ time).

Note that since the plot illustrates vehicles arriving in queue and departing from queue, this model most closely represents what has been defined as *time-in-queue delay*. There are many simplifications that have been assumed, however, in constructing this simple depiction of delay. It is important to understand the two major simplifications:

- The assumption of a uniform arrival rate is a simplification. Even at a completely isolated location, actual arrivals would be *random* (i.e., would have an average rate over time), but inter-vehicle arrival times would vary around an average rather than being constant. Within coordinated signal systems, however, vehicle arrivals are in platoons.
- It is assumed that the queue is building at a point location (as if vehicles were stacked on top of one another). In reality, as the queue grows, the rate at which vehicles arrive at its end is the arrival rate of vehicles (at a point), plus a component representing the backward growth of the queue in space.

Both of these can have a significant effect on actual results. Modern models account for the former in ways that will be discussed subsequently. The assumption of a "point queue," however is imbedded in many modern applications.

Figure 17.11 expands the range of Figure 17.10 to show a series of GREEN phases and depicts three different types of operation. It also allows for an arrival function, $a(t)$, that varies, while maintaining the departure function, $d(t)$, described previously.

Figure 17.11 (a) shows stable flow throughout the period depicted. No signal cycle "fails" (i.e., ends with some vehicles queued during the preceding RED unserved). During every GREEN phase, the departure

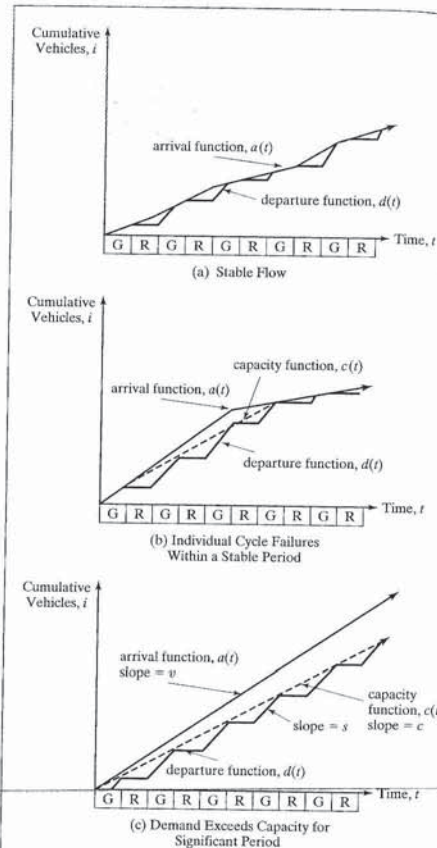


Figure 17.11: Three Delay Scenarios (Adapted with permission of Transportation Research Board, National Research Council, Washington DC, from V.F. Hurdle, "Signalized Intersection Delay Model: A Primer for the Uninitiated," *Transportation Research Record* 971, pgs. 97, 98, 1984.)

function "catches up" with the arrival function. Total aggregate delay during this period is the total of all the triangular areas between the arrival and departure

curves. This type of delay is often referred to as "uniform delay."

In Figure 17.11 (b), some of the signal phases "fail." At the end of the second and third GREEN intervals, some vehicles are not served (i.e., they must wait for a second GREEN interval to depart the intersection). By the time the entire period ends, however, the departure function has "caught up" with the arrival function and there is no residual queue left unserved. This case represents a situation in which the overall period of analysis is stable (i.e., total demand does not exceed total capacity). Individual cycle failures within the period, however, have occurred. For these periods, there is a second component of delay in addition to uniform delay. It consists of the area between the arrival function and the dashed line, which represents the capacity of the intersection to discharge vehicles, and has the slope c . This type of delay is referred to as "overflow delay."

Figure 17.11 (c) shows the worst possible case: Every GREEN interval "fails" for a significant period of time, and the residual, or unserved, queue of vehicles continues to grow throughout the analysis period. In this case, the overflow delay component grows over time, quickly dwarfing the uniform delay component.

The latter case illustrates an important practical operational characteristic. When demand exceeds capacity ($v/c > 1.00$), the delay depends upon the length of time that the condition exists. In Figure 17.11 (b), the condition exists for only two phases. Thus, the queue and the resulting overflow delay is limited. In Figure 17.11 (c), the condition exists for a long time, and the delay continues to grow throughout the oversaturated period.

Components of Delay

In analytic models for predicting delay, there are three distinct components of delay that may be identified:

- *Uniform delay* is the delay based on an assumption of uniform arrivals and stable flow with no individual cycle failures.
- *Random delay* is the additional delay, above and beyond uniform delay, because flow is randomly distributed rather than uniform at isolated intersections.
- *Overflow delay* is the additional delay that occurs when the capacity of an individual phase or

series of phases is less than the demand or arrival flow rate.

In addition, the delay impacts of platoon flow (rather than uniform or random) is treated as an adjustment to uniform delay. Many modern models combine the random and overflow delays into a single function, which is referred to as "overflow delay," even though it contains both components.

The differences between uniform, random, and platooned arrivals are illustrated in Figure 17.12. As noted, the analytic basis for most delay models is the assumption of uniform arrivals, which are depicted in Figure 17.12 (a). Even at isolated intersections, however, arrivals would be random, as shown in Figure 17.12 (b). With random arrivals, the underlying rate of arrivals is a constant, but the inter-arrival times are exponentially distributed around an average. In most urban and suburban cases, where a signalized intersection is likely to be part of a coordinated signal system, arrivals will be in organized platoons that move down the arterial in a cohesive group. The exact time that a platoon arrives at a downstream signal has an enormous potential effect on delay. A platoon of vehicles arriving at the beginning of the RED forces most vehicles to stop for the entire length of the RED phase. The same platoon of vehicles arriving at the beginning of the GREEN phase may flow through the intersection without any vehicles stopping. In both cases, the arrival flow, v , and the capacity of the intersection, c , are the same. The resulting delay, however, would vary significantly. The existence of platoon

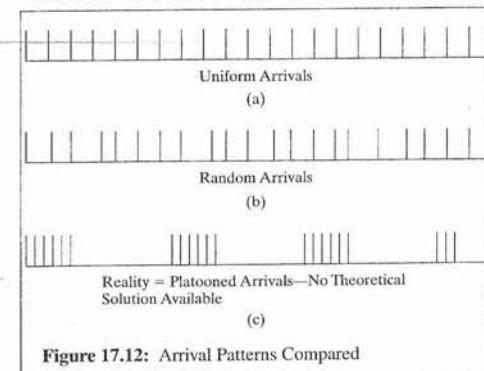


Figure 17.12: Arrival Patterns Compared

arrivals, therefore, necessitates a significant adjustment to models based on theoretically uniform or random flow.

Webster's Uniform Delay Model

Virtually every model of delay starts with Webster's model of uniform delay. Initially published in 1958 [7], this model begins with the simple illustration of delay depicted in Figure 17.10, with its assumptions of stable flow and a simple uniform arrival function. As noted previously, aggregate delay can be estimated as the area between the arrival and departure curves in the figure. Thus, Webster's model for uniform delay is the area of the triangle formed by the arrival and departure functions. For clarity, this triangle is shown again in Figure 17.13.

The area of the aggregate delay triangle is simply one-half the base times the height, or:

$$UD_a = \frac{1}{2}RV$$

where: UD_a = aggregate uniform delay, veh-secs

R = length of the RED phase, s

V = total vehicles in queue, vehs

By convention, traffic models are not developed in terms of RED time. Rather, they focus on GREEN time.

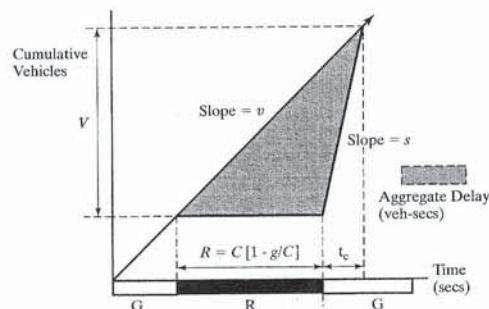


Figure 17.13: Webster's Uniform Delay Model Illustrated

Thus, Webster substitutes the following equivalence for the length of the RED phase:

$$R = C \left[1 - \left(\frac{g}{C} \right) \right]$$

where: C = cycle length, s

g = effective green time, s

In words, the RED time is the portion of the cycle length that is not effectively green.

The height of the triangle, V , is the total number of vehicles in the queue. In effect, it includes vehicles arriving during the RED phase, R , plus those that join the end of the queue while it is moving out of the intersection (i.e., during time t_c in Figure 17.13). Thus, determining the time it takes for the queue to clear, t_c , is an important part of the model. This is done by setting the number of vehicles arriving during the period $R + t_c$ equal to the number of vehicles departing during the period t_c , or:

$$v(R + t_c) = st_c$$

$$R + t_c = \left(\frac{s}{v} \right) t_c$$

$$R = t_c \left(\frac{s}{v} - 1 \right)$$

$$t_c = \frac{R}{\left(\frac{s}{v} - 1 \right)}$$

17.5 DELAY AS A MEASURE OF EFFECTIVENESS

Then, substituting for t_c :

$$V = v(R + t_c) = v \left[R + \frac{R}{\left(\frac{s}{v} - 1 \right)} \right] = R \left[\frac{vs}{s - v} \right]$$

and for R :

$$V = C \left[1 - \left(\frac{g}{C} \right) \right] \left[\frac{vs}{s - v} \right]$$

Then, aggregate delay can be stated as:

$$UD_a = \frac{1}{2}RV = \frac{1}{2}C^2 \left[1 - \frac{g}{C} \right]^2 \left[\frac{vs}{s - v} \right] \quad (17-18)$$

where all variables are as previously defined.

Equation 17-18 estimates aggregate uniform delay in vehicle-seconds for one signal cycle. To get an estimate of average uniform delay per vehicle, the aggregate is divided by the number of vehicles arriving during the cycle, vC . Then:

$$UD = \frac{1}{2}C \left[\frac{1 - g/C}{1 - v/s} \right]^2 \quad (17-19)$$

Another form of the equation uses the capacity, c , rather than the saturation flow rate, s . Noting that $s = c/(g/C)$, the following form emerges:

$$UD = \frac{1}{2}C \frac{[1 - g/C]^2}{[1 - (g/C)(v/c)]} = \frac{0.50C[1 - g/C]^2}{1 - (g/c)X} \quad (17-20)$$

where: UD = average uniform delay per vehicles, s/veh

C = cycle length, s

g = effective green time, s

v = arrival flow rate, veh/h

c = capacity of intersection approach, veh/h

X = v/c ratio, or degree of saturation

This average includes the vehicles that arrive and depart on green, accruing no delay. This is appropriate. One of the objectives in signalizations is to minimize the number or proportion of vehicles that must stop. Any meaningful quality measure would have to include the positive impact of vehicles that are not delayed.

In Equation 17-20, it must be noted that the maximum value of X (the v/c ratio) is 1.00. As the uniform delay model assumes no overflow, the v/c ratio cannot be more than 1.00.

Modeling Random Delay

The uniform delay model assumes that arrivals are uniform and that no signal phases fail (i.e., that arrival flow is less than capacity during every signal cycle of the analysis period).

At isolated intersections, vehicle arrivals are more likely to be random. A number of stochastic models have been developed for this case, including those by Newall [8], Miller [9,10], and Webster [7]. Such models assume that inter-vehicle arrival times are distributed according to the Poisson distribution, with an underlying average arrival rate of v vehicles/unit time. The models account for both the underlying randomness of arrivals and the fact that some individual cycles within a demand period with $v/c < 1.00$ could fail due to this randomness. This additional delay is often referred to as "overflow delay," but it does not address situations in which $v/c > 1.00$ for the entire analysis period. This text refers to additional delay due to randomness as "random delay," RD , to distinguish it from true overflow delay when $v/c > 1.00$. The most frequently used model for random delay is Webster's formulation:

$$RD = \frac{X^2}{2v(1 - X)} \quad (17-21)$$

where: RD = average random delay per vehicle, s/veh

X = v/c ratio

This formulation was found to somewhat overestimate delay, and Webster proposed that total delay (the sum of uniform and random delay) be estimated as:

$$D = 0.90(UD + RD) \quad (17-22)$$

where: D = sum of uniform and random delay

The inconsistency occurs when the v/c ratio (X) is in the vicinity of 1.00. When the v/c ratio is below 1.00, a random delay model is used, as there is no "overflow" delay in this case. Webster's random delay model (Equation 17-22), however, contains the term $(1-X)$ in the denominator. Thus, as X approaches a value of 1.00, random delay increases asymptotically to an infinite value. When the v/c ratio (X) is greater than 1.00, an overflow delay model is applied. The overflow delay model of Equation 17-24, however, has an overflow delay of 0 when $X = 1.00$, and increases uniformly with increasing values of X thereafter.

Neither model is accurate in the immediate vicinity of $v/c = 1.00$. Delay does not become infinite at $v/c = 1.00$. There is no true "overflow" at $v/c = 1.00$, although individual cycle failures due to random arrivals do occur. Similarly, the overflow model, with overflow delay = 0.0 s/veh at $v/c = 1.00$ is also unrealistic. The additional delay of individual cycle failures due to the randomness of arrivals is not reflected in this model.

In practical terms, most studies confirm that the uniform delay model is a sufficiently predictive tool (except for the issue of platooned arrivals) when the v/c ratio is 0.85 or less. In this range, the true value of random delay is minuscule, and there is no overflow delay. Similarly, the simple theoretical overflow delay model (when added to uniform delay) is a reasonable predictor when $v/c \geq 1.15$ or so. The problem is that the most interesting cases fall in the intermediate range ($0.85 < v/c < 1.15$), for which neither model is adequate. Much of the more recent work in delay modeling involves attempts to bridge this gap, creating a model that closely follows the uniform delay model at low v/c ratios, and approaches the theoretical overflow delay model at high v/c ratios (≥ 1.15), producing "reasonable" delay estimates in between. Figure 17.17 illustrates this as the dashed line.

The most commonly used model for bridging this gap was developed by Akcelik for the Australian Road Research Board's signalized intersection analysis procedure [11, 12]:

$$OD = \frac{cT}{4} \left[(X - 1) + \sqrt{(X - 1)^2 + \left(\frac{12(X - X_o)}{cT} \right)} \right]$$

$$X_o = 0.67 + \left(\frac{sg}{600} \right)$$

$$OD = 0.0 \text{ s/veh for } X \leq X_o \quad (17-26)$$

where: T = analysis period, h
 X = v/c ratio
 c = capacity, veh/h
 s = saturation flow rate, veh/s, (vehs per second of green)
 g = effective green time, s

The only relatively recent study resulting in large amounts of delay measurements in the field was conducted by Reilly, *et al.* [13] in the early 1980s to calibrate a model for use in the 1985 edition of the *Highway Capacity Manual*. The study concluded that Equation 17-26 substantially overestimated field-measured values of delay and recommended that a factor of 0.50 be included in the model to adjust for this. The version of the delay equation that was included in the 1985 *Highway Capacity Manual* ultimately did not follow this recommendation, and included other empiric adjustments to the theoretical equation.

17.5.4 Delay Models in the HCM 2000

The delay model incorporated into the HCM 2000 [2] includes the uniform delay model, a version of Akcelik's overflow delay model, and a term covering delay from an existing or residual queue at the beginning of the analysis period. The model is:

$$d = d_1 PF + d_2 + d_3$$

$$d_1 = \left(\frac{C}{2} \right) * \left\{ \frac{(1 - g/c)^2}{1 - [\min(1, X) * g/c]} \right\}$$

$$d_2 = 900 T \left[(X - 1) + \sqrt{(X - 1)^2 + \left(\frac{8kIX}{cT} \right)} \right] \quad (17-27)$$

where: d = control delay, s/veh
 d_1 = uniform delay component, s/veh
 PF = progression adjustment factor

d_2 = overflow delay component, s/veh
 d_3 = delay due to pre-existing queue, s/veh
 T = analysis period, h
 X = v/c ratio
 C = cycle length, s
 k = incremental delay factor for actuated controller settings; 0.50 for all pre-timed controllers
 I = upstream filtering/metering adjustment factor; 1.00 for all individual intersection analyses
 c = capacity, veh/h

17.5.5 Examples in Delay Estimation

Example 17-1:

Consider the following situation: An intersection approach has an approach flow rate of 1,000 veh/h, a saturation flow rate of 2,800 veh/h, a cycle length of 90 s, and a g/C ratio of 0.55. What average delay per vehicle is expected under these conditions?

Solution:

To begin, the capacity and v/c ratio for the intersection approach must be computed. This will determine what model(s) are most appropriate for application in this case:

$$c = s * (g/c) = 2,800 * 0.55 = 1,540 \text{ veh/h}$$

$$v/c = X = \frac{1,000}{1,540} = 0.649$$

Example 17-2:

How would the above result change if the demand flow rate increased to 1,600 veh/h?

Solution:

In this case, the v/c ratio now changes to 1,600/1,540 = 1.039. This is in the difficult range of 0.85–1.15 for which neither the simple random flow model nor the simple overflow delay model are accurate. The Akcelik model of Equation 17-26 will be used. Total delay, however, includes

The progression factor is an empirically calibrated adjustment to uniform delay that accounts for the effect of platooned arrival patterns. This adjustment is detailed in Chapter 21. The delay due to preexisting queues, d_3 , is found using a relatively complex model. (see Chapter 21).

In the final analysis, all delay modeling is based on the determination of the area between an arrival curve and a departure curve on a plot of cumulative vehicles vs. time. As the arrival and departure functions are permitted to become more complex and as rates are permitted to vary for various sub-parts of the signal cycle, the models become more complex as well.

As this is a relatively low value, the uniform delay equation (Equation 17-19) may be applied directly. There is little random delay at such a v/c ratio and no overflow delay to consider. Thus:

$$d = \left(\frac{C}{2} \right) * \left[\frac{(1 - g/c)^2}{1 - g/s} \right]$$

$$= \left(\frac{90}{2} \right) * \left[\frac{(1 - 0.55)^2}{1 - 1,000/2,800} \right] = 14.2 \text{ s/veh}$$

Note that this solution assumes that arrivals at the subject intersection approach are random. Platooning effects are not taken into account.

both uniform delay and overflow delay. The uniform delay component when $v/c > 1.00$ is given by Equation 17-23:

$$UD = 0.50 C \left(1 - g/c \right) = 0.50 * 90 * (1 - 0.55)$$

$$= 20.3 \text{ s/veh}$$

Use of Akcelik's overflow delay model requires that the analysis period be selected or arbitrarily set. If a one-hour