

ÉLECTRONIQUE

Les multivibrateurs

Multivibrateurs : plan

- ▶ **définitions**
- ▶ bistables
- ▶ monostables
- ▶ astables

Définitions

- ▶ **état logique stable**
 - ◆ état qui se maintient indéfiniment tant que l'on n'active pas une entrée particulière appelée "déclenchement" ou "trigger" et destinée à provoquer le changement d'état
- ▶ **état logique métastable ou quasi-stable**
 - ◆ état qui se maintient temporairement
 - ◆ le circuit retourne spontanément dans l'autre état logique
- ▶ **multivibrateurs**
 - ◆ montages logiques à deux états
 - ◆ trois types basés sur le nombre d'états stables en sortie :
 - bistable
 - monostable
 - astable

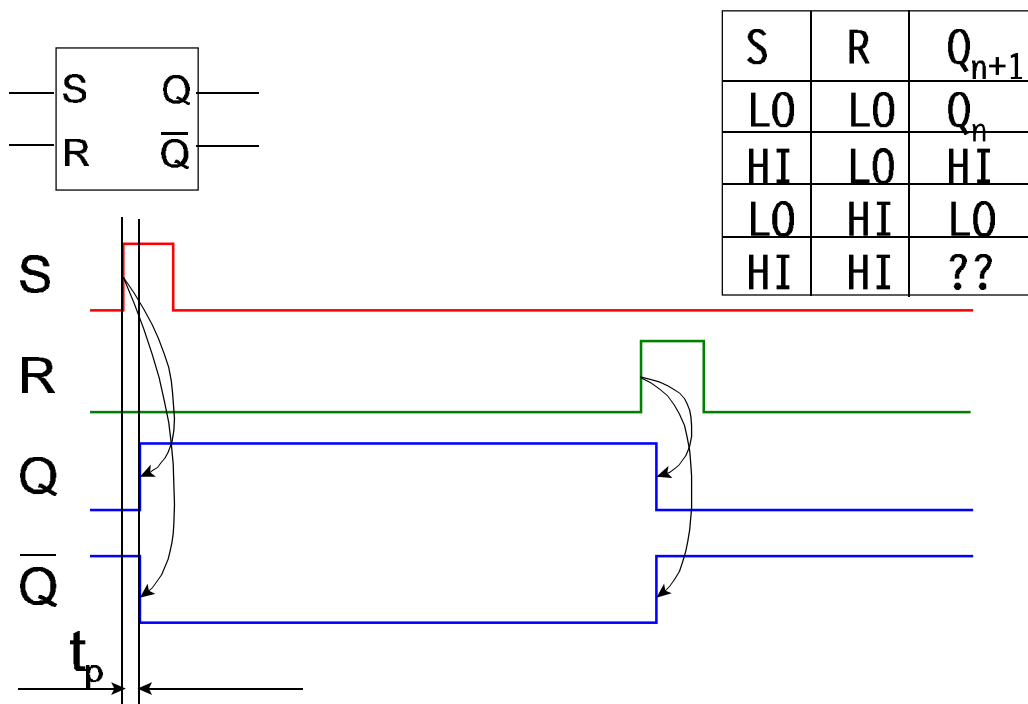
On dit d'un état logique qu'il est stable s'il se maintient indéfiniment, tant que l'on ne coupe pas l'alimentation et que l'on n'active sur une entrée particulière appelée "déclenchement" ou "trigger" destinée à provoquer le changement d'état.

On parle d'état logique métastable ou quasi-stable lorsque le circuit n'y reste que temporairement puis passe spontanément dans l'autre état logique

Les multivibrateurs sont une catégorie particulière de montages logiques, qui se divise en trois types caractérisés par la stabilité de leurs 2 états logiques de sortie :

- les bistables : circuits à 2 états stables
- les monostables : circuits à 1 état stable et 1 état métastable
- les astables : circuits logiques à 2 états métastables entre lesquels la sortie oscille continuellement

Bistable



Le bistable est un circuit logique à 2 états stables.

Il possède donc deux entrées de déclenchement destinées respectivement à l'activation et à la désactivation de la variable de sortie.

La sortie est souvent constituée de 2 signaux opposés, ce qui permet de choisir un signal actif à l'état haut ou à l'état bas pour commander les circuits placés en aval.

Le bistable illustré ici est le plus simple et porte le nom de RS (ou SR), d'après ses deux entrées de déclenchement R(Reset) et S(Set).

En activant l'entrée S on provoque l'activation de la variable de sortie Q, après un temps de propagation t_p . Cet état se maintiendra jusqu'à ce que l'on active l'entrée R, qui remettra Q à l'état inactif.

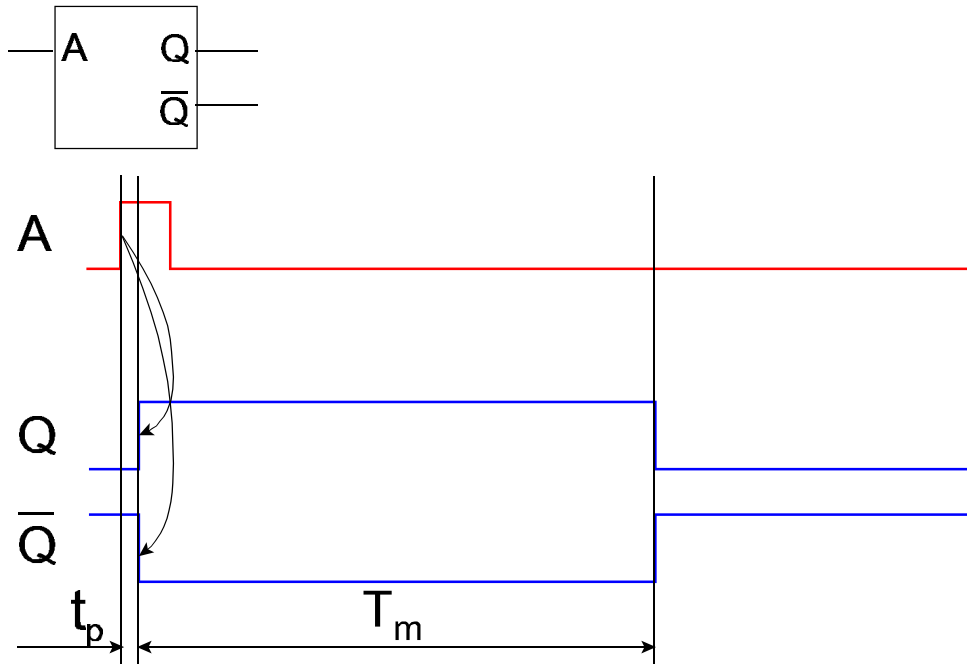
Remarquons que les entrées sont actives à tour de rôle. S est inopérante lorsque la sortie est déjà active et R inopérante lorsque la sortie est déjà inactive.

La table de vérité montre bien le caractère séquentiel du circuit, l'indice de la variable de sortie marquant la succession des états dans le temps. Remarquer dans cette table l'interdiction d'activer simultanément les deux entrées, faute de quoi l'état de sortie est indéterminé. On appliquera donc sur chaque entrée des impulsions qui doivent être remises à leur état inactif avant que l'on active l'autre entrée.

Un bistable constitue le bit élémentaire de mémoire et sert donc de base

- à la logique séquentielle, pour laquelle on a besoin d'une mémorisation de l'état passé du système
- aux circuits de mémoire des ordinateurs

Monostable



Le monostable doit son nom au fait qu'il ne possède qu'un état stable(sortie inactive), l'autre état (sortie active) étant métastable.

Il possède une entrée de déclenchement qui active la sortie pendant la durée déterminée T_m de l'état métastable.

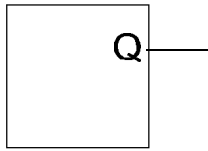
Remarquons que le monostable doit contenir une base de temps pour pouvoir déterminer ce délai.

Un monostable sert le plus souvent à :

- retarder une impulsion d'un délai déterminé: le flanc montant de Q' est retardé de T_m par rapport à celui de A
- mettre en forme une impulsion trop longue ou trop courte; ici, l'impulsion A est transformée en une impulsion plus longue Q .

On en verra plusieurs applications dans la suite du cours.

Astable



Période

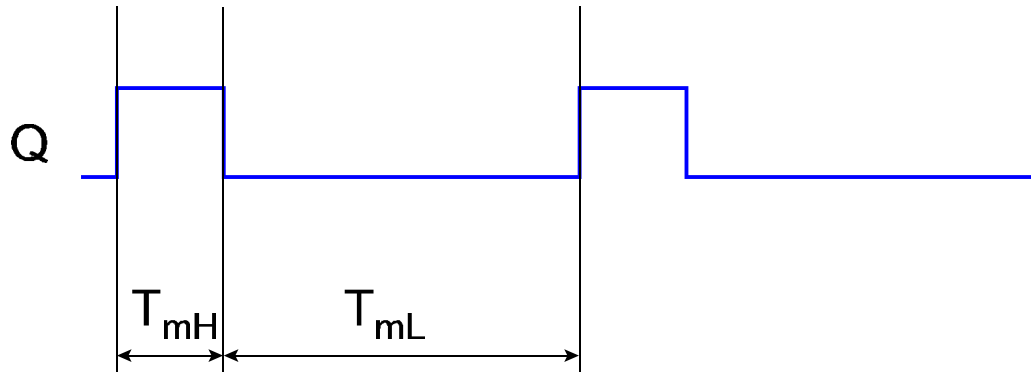
$$T = T_{mH} + T_{mL}$$

Fréquence

$$f = \frac{1}{T}$$

Rapport cyclique

$$\delta = \frac{T_{mH}}{T}$$



L'astable, comme son nom l'indique, ne possède aucun état stable et donc aucune entrée de déclenchement. Après avoir gardé l'état actif pendant un temps T_{mH} , il repasse spontanément à l'état inactif pendant un temps T_{mL} , et ainsi de suite, ...

Il s'agit donc d'un oscillateur non-sinusoïdal, puisqu'il fournit un signal logique rectangulaire.

Il est caractérisé par :

- sa période $T = T_{mH} + T_{mL}$

- sa fréquence $f = 1/T$

- son rapport cyclique ou "duty-cycle" $\delta = T_{mH} / T$. Si $\delta = 50\%$ ($T_{mH} = T_{mL}$), la sortie est une onde carrée.

L'astable sert donc de circuit d'horloge pour les logiques synchrones.

L'astable doit contenir une base de temps pour pouvoir déterminer ses deux délais. Si une grande stabilité et/ou précision est requise, la période de l'astable peut être fixée par un résonateur à quartz.

Multivibrateurs : plan

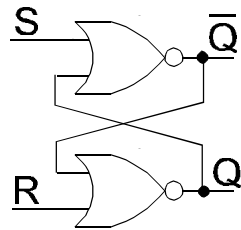
- ▶ définitions
- ▶ **bistables**
 - ◆ **déclenchés par niveaux ou "latches"**
 - ◆ déclenchés par flancs ou "flip-flops"
- ▶ monostables
- ▶ astables

Parmi les bistables, on définit 2 grandes catégories basées sur le type de déclenchement :

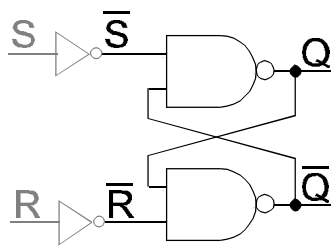
- les **"latches"** (=verrou) pour lesquels les changements d'états et le verrouillage dans un des deux états de sortie dépendent des niveaux logiques aux entrées (on parle de déclenchement par niveau ou "level triggering")
- les **"flip-flops"** pour lesquels les changements d'états à la sortie sont gouvernés par le flanc montant ou descendant d'une entrée spécifique appelée horloge (on parle de déclenchement par flanc ou "edge triggering")

Bistable S-R (latch)

réalisation à l'aide de portes NOR et NAND



S	R	Q_{n+1}
LO	LO	Q_n
HI	LO	HI
LO	HI	LO
HI	HI	??

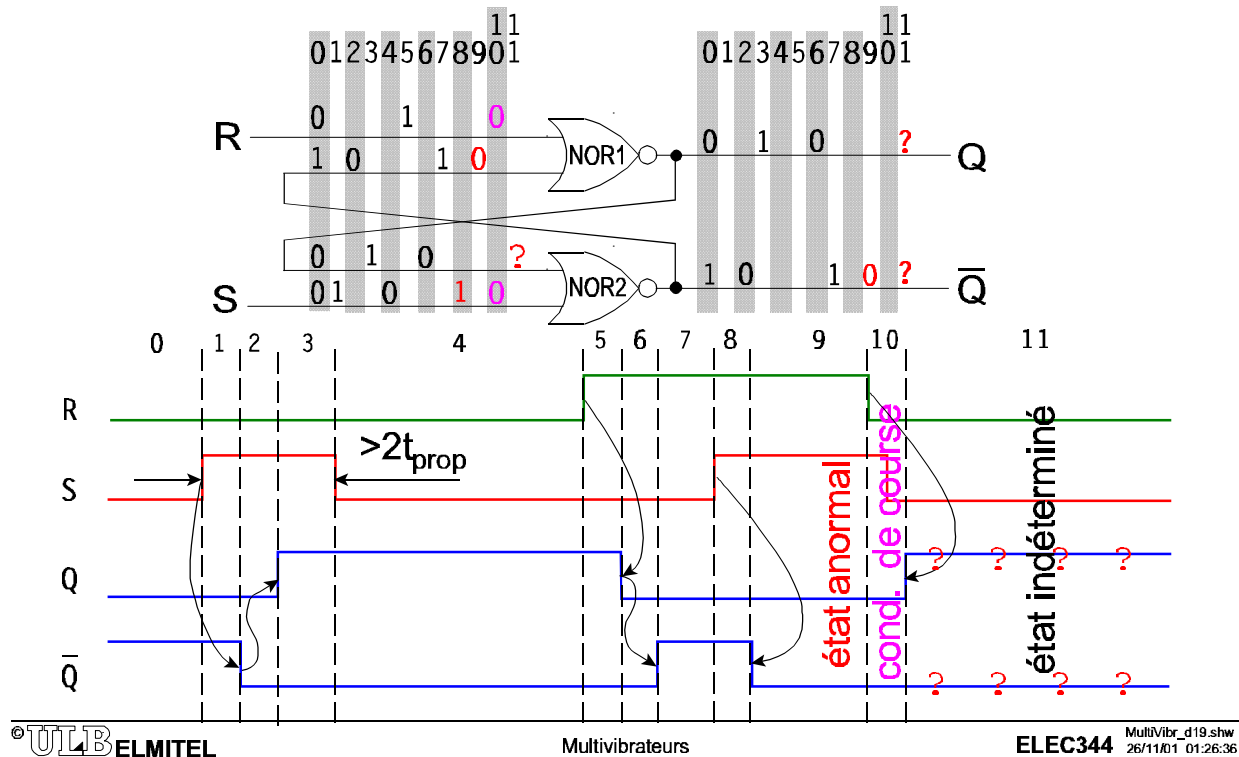


$\bar{S}$	$\bar{R}$	Q_{n+1}
HI	HI	Q_n
LO	HI	HI
HI	LO	LO
LO	LO	??

En pratique, le bistable RS se réalise à l'aide de portes NOR ou NAND. La table de vérité se vérifie aisément.

Si l'on utilise des portes NAND, les entrées sont actives à l'état LO, à moins que l'on ne rajoute deux inverseurs en entrée.

Bistable R-S : chronogramme



17

Supposons que l'état initial du bistable soit $Q=0$ et que la succession des états soit la suivante :

intervalle 0 : $S=0$ $R=0 \Rightarrow$ état stable $Q=0$ $Q'=1$

intervalle 1 : $S=1$ $Q=0 \Rightarrow Q'$ passe à 0 après 1 temps de propagation

intervalle 2 : $Q'=0$ $R=0 \Rightarrow Q$ passe à 1 après 1 temps de propagation

intervalle 3 : S peut repasser à 0 car la porte NOR2 est verrouillée par $Q=1$

intervalle 4 : $S=0$ $R=0 \Rightarrow$ état stable $Q=1$ $Q'=0$

intervalle 5 : $R=1$ $Q'=0 \Rightarrow Q$ passe à 0 après 1 temps de propagation

intervalle 6 : $S=0$ $Q=0 \Rightarrow Q'$ passe à 1 après 1 temps de propagation

intervalle 7 : état stable $Q=0$ $Q'=1$

intervalle 8 : $S=1$ $R=1 \Rightarrow Q=0$ condition anormale; Q' passe à 0 après 1 temps de propagation

intervalle 9 : état anormal $Q=Q'=0$

Supposons qu'à partir de cet état on ramène $S=0$ et $R=0$ simultanément. En réalité "simultanément" n'a pas de sens, il y a toujours un des deux signaux qui prendra un peu d'avance sur l'autre. La figure illustre le cas où R retourne à 0 le premier. Dans ce cas on a

intervalle 10 : $R=0$ $Q'=0 \Rightarrow Q$ passe à 1 après 1 temps de propagation

intervalle 11 : $S=0$ $R'=0 \Rightarrow$ état stable $Q=1$ $Q'=0$

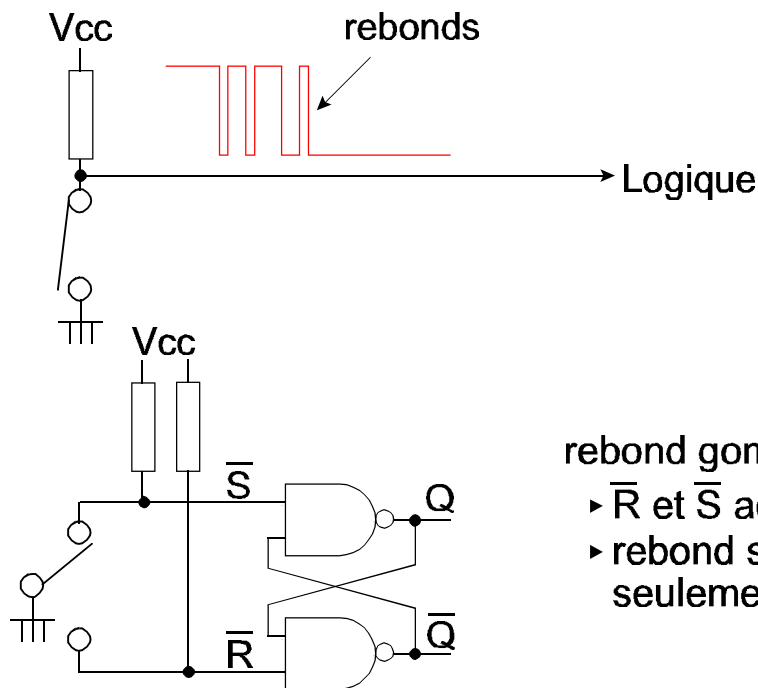
REM 1 : la durée minimum d'une impulsion de déclenchement est $2 \times t_{prop}$, car l'impulsion ne peut retourner à l'état inactif qu'après que le bistable a basculé (cf intervalles 1,2 et 3)

REM 2 : l'état $R=1$ $S=1$ est en fait indéterminé. L'exemple de la figure est une pure conjecture. Si S retourne à 0 avant R , ou si le temps de propagation de NOR1 est plus long que NOR2 on obtient l'autre état stable à l'intervalle 11.

Ce phénomène où un état logique dépend d'une faible avance ou retard d'un signal sur un autre porte le nom de "condition de course".

L'état anormal $Q=Q'$ des intervalles 9 et 10 et l'indétermination de l'état final 11 justifient l'interdiction de la condition d'entrée $R=S=1$.

Bistable R-S : anti-rebond



rebond gommé

- $\bar{R}$ et $\bar{S}$ actifs tour à tour
- rebond sur 1 des contacts seulement

L'interrupteur, ou le bouton poussoir, reste l'un des interfaces homme-machine les plus simples, les plus utilisés et les moins chers. Le "micro switch" est également un des principaux moyens d'indiquer une "fin de course" dans un déplacement mécanique (porte motorisée, chariot de machine-outil, niveau d'une cuve,).

Le défaut le plus gênant est le rebond des contacts mécaniques pendant plusieurs dizaines de ms.

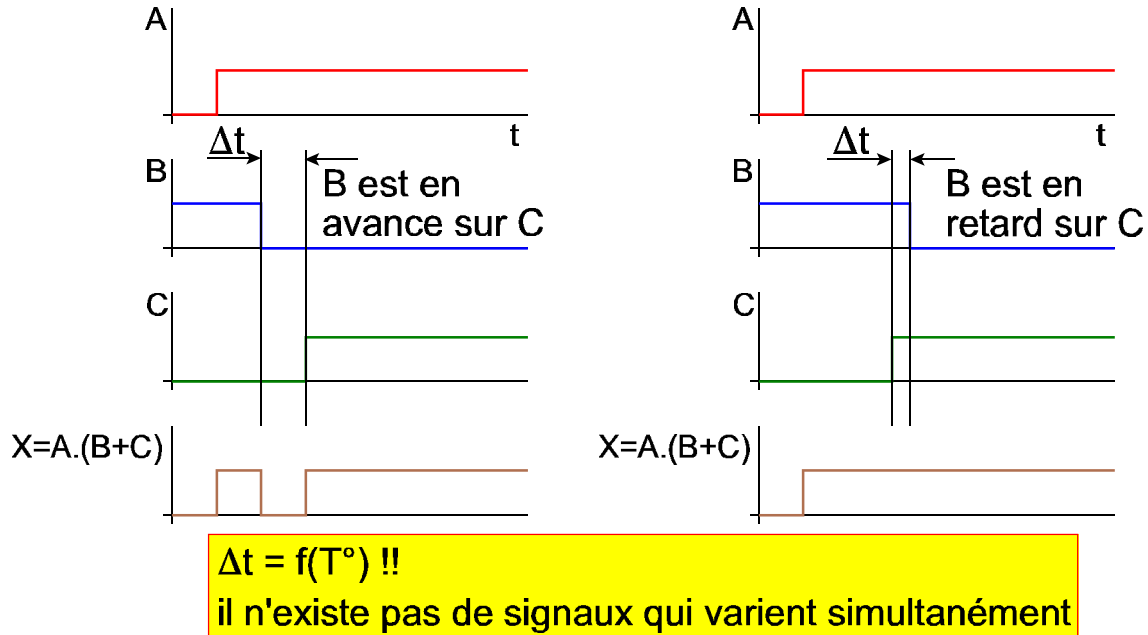
Le montage le plus simple consiste en une résistance "pull-up" qui fixe le niveau logique par défaut en HI, l'interrupteur jouant le rôle de "pull-down" (comme dans une sortie à collecteur ouvert). Les rebonds de l'interrupteur se reflètent alors directement sur le signal logique de sortie. Ces oscillations logiques sont parfois intolérables; citons par exemple:

- les circuits logiques séquentiels dans lesquels deux états LO successifs ont des effets différents sur le système (ex: poussoir d'ouverture-fermeture de porte motorisée)
- les circuits où le nombre de manoeuvres de l'interrupteur fait l'objet d'un comptage.

Dans ce cas, le bistable R-S constitue un excellent "anti-rebond". L'interrupteur simple à 2 pôles doit être remplacé par un inverseur dont la borne commune est à la masse et les deux sorties sont reliées aux entrées R' et S' placées à leur état inactif par des "pull-up". L'action anti-rebond résulte alors de deux propriétés :

- les entrées du bistable ne sont actives qu'à tour de rôle; le premier passage à LO de l'entrée S' fait basculer Q à HI, les rebonds suivants sont ignorés. Il en va de même pour l'entrée R'
- l'amplitude mécanique du rebond est inférieure à la distance qui sépare les deux contacts

Condition de course



En regardant le chronogramme du bistable RS, on voit que les sorties Q et Q' ne basculent pas en même temps. Ceci peut être gênant si l'on exploite les deux sorties dans la logique placée en aval du bistable et si l'on compte sur cette simultanéité. De manière tout à fait générale, la simultanéité n'existe dans aucun circuit logique, même pour des sorties de circuits semblables dont les entrées sont reliées ensemble. Il existe toujours

- les temps de propagation et de transition des portes
- de légères dissymétries des transistors au sein des circuits intégrés
- des différences de temps de propagation du signal au sein des pistes des circuits imprimés

La température joue également son rôle puisqu'elle modifie les propriétés des transistors.

Concevoir des circuits basés sur la simultanéité de certains signaux, ou sur les retards introduits par les portes logiques, constitue dès lors une pratique à déconseiller, sauf cas spéciaux dont nous verrons un exemple.

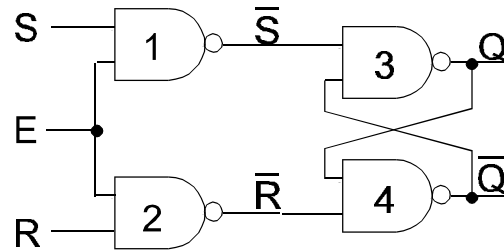
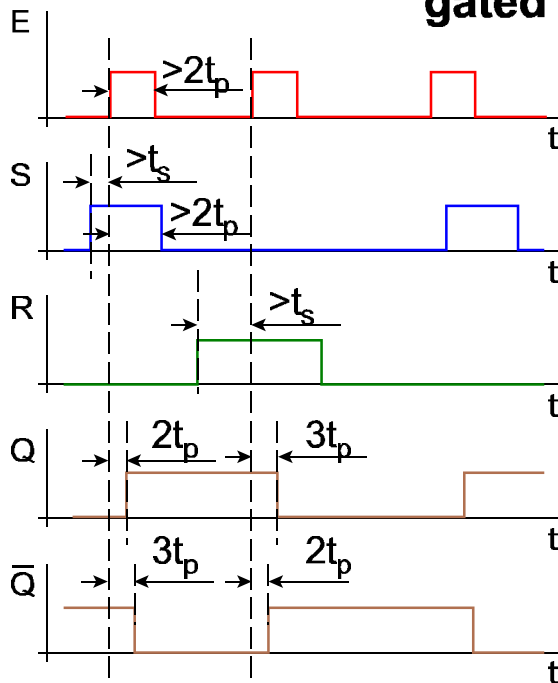
La figure illustre ici la différence entre le signal de sortie X d'une fonction logique simple, suivant que le signal B passe en LO avant ou après le signal C.

Si B et C sont totalement indépendants, le comportement de ce circuit est normal. Par contre, si le délai entre B et C est de l'ordre de quelques ns et dépend de facteurs involontaires (comme la température ou le temps de montée d'un signal externe), il s'agit d'une "condition de course": le résultat X dépend de la course que se livrent les variables B et C, ce qui n'est pas tolérable.

Pour éliminer ces conditions de course on a souvent recours à de la logique synchrone.

Bistable R-S avec entrée d'activation

"gated S-R latch"



E	S	R	Q_{n+1}
LO	X	X	Q_n
HI	LO	LO	Q_n
HI	HI	LO	HI
HI	LO	HI	LO
HI	HI	HI	??

On peut ajouter au bistable RS deux portes pour bloquer/débloquer les entrées par un signal appelé E ou EN (pour "ENable") qui est donc un signal d'inhibition/activation. Ce bistable reste dans la catégorie des latches : le signal E est actif sur son niveau HI, comme le confirme la table de vérité. L'intérêt principal d'une telle entrée est de permettre un fonctionnement synchrone.

En logique synchrone, on n'autorise les changements d'état des entrées logiques qu'à des instants précis, déterminés par des impulsions, généralement périodiques, appelées impulsions d'horloge ou "clock pulses".

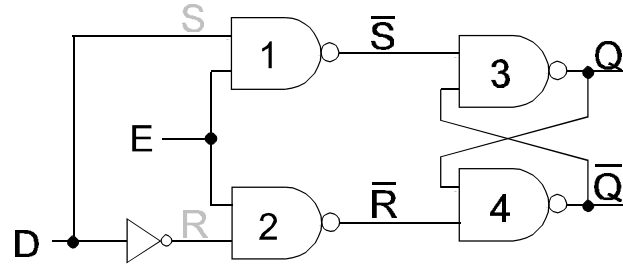
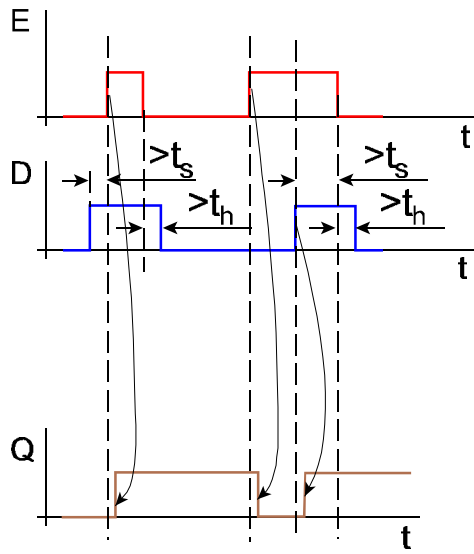
La première partie du chronogramme illustre ce fonctionnement avec l'horloge connectée à l'entrée E. Lorsque le signal E est inactif (LO), les entrées R et S ne sont pas prises en compte. L'activation de l'entrée S n'a d'effet que lorsque l'impulsion d'horloge passe à HI, la porte NAND d'entrée change alors d'état après un temps de propagation t_p , la sortie Q après $2.t_p$ et Q' après $3.t_p$, ce qui verrouille le bistable. On en déduit que plusieurs conditions temporelles doivent être respectées pour un fonctionnement correct :

- le signal S(ou R) doit être actif un peu avant le flanc montant de l'horloge, pour que l'on puisse garantir que le changement d'état du bistable soit bien synchronisé sur l'horloge (temps d'application ou "setup-time" t_s)
- les signaux S (ou R) et E doivent rester simultanément actifs pendant au moins $2.t_p$ pour que le signal interne S' (ou R') n'ait pas disparu avant que Q' (ou Q) ne verrouille le bistable

La seconde partie du chronogramme montre qu'il s'agit bien d'un latch (commande par niveau et non par flanc). Si le signal S (ou R) devient actif pendant le niveau haut de E, il est pris en compte et le bistable bascule sans être synchronisé sur E.

Remarquons enfin que l'état d'entrée $R=S=HI$ reste interdit.

Bistable D ou D latch



E	D	Q_{n+1}
LO	X	Q_n
HI	LO	LO
HI	HI	HI

Un moyen simple d'éviter la condition $R=S=HI$ est de faire en sorte que les deux entrées soient toujours des compléments logiques. On obtient alors un bistable à une seule entrée de commande appelée D (pour DATA). L'entrée d'activation porte le nom de LE ("Latch Enable")

La table de vérité est simple :

- tant que l'entrée LE est active, la sortie Q recopie D (au temps de propagation près) : le bistable est dit "transparent"
- lorsque l'on désactive LE est inactive, l'entrée D devient inopérante; la sortie est verrouillée ("latched") dans le dernier état pris par D avant la désactivation de LE.

"hold time"

Il faut éviter une condition de course entre D et E au moment où l'on désactive LE; on impose donc un temps de maintien ou "hold time" t_h qui assure la stabilité de D pendant la désactivation de E.

"setup time"

Deux types de fonctionnement sont possibles :

1°) On désire que les changements de la sortie soient synchronisés sur le flanc montant de LE (1ère impulsion sur LE dans la figure)

Dans ce cas, le temps d'activation ou "setup time" se réfère au flanc montant de LE

2°) On désire que le latch soit transparent, sauf lorsque l'on désactive LE (2ème impulsion sur LE dans la figure)

Dans ce cas, le temps d'activation ou "setup time" se réfère au flanc descendant de LE.

Nous verrons un exemple d'une telle application du D-Latch dans le démultiplexage de bus des microprocesseurs.

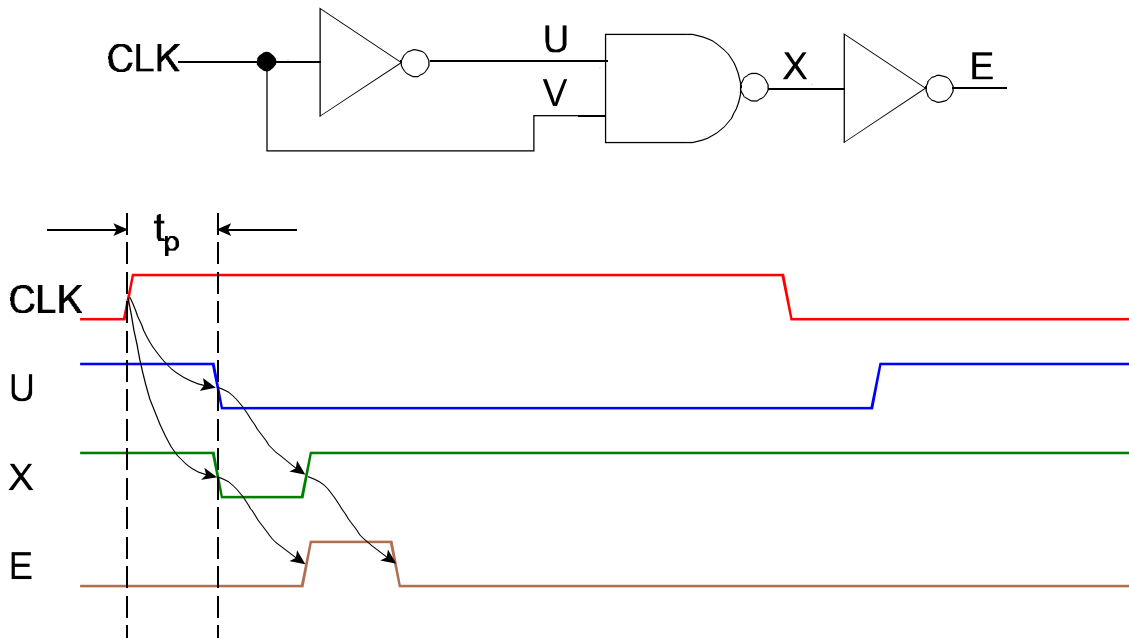
Multivibrateurs : plan

- ▶ définitions
- ▶ **bistables**
 - ◆ déclenchés par niveaux ou "latches"
 - ◆ **déclenchés par flancs ou "flip-flops"**
- ▶ monostables
- ▶ astables

Dans les bistables déclenchés par flanc, on va créer une très courte impulsion lors d'un des deux flancs d'horloge et se servir de cette impulsion pour activer les entrées du bistable. Il faut donc d'abord disposer d'un détecteur de flanc.

Déclenchement par flanc

exemple de détection de flanc

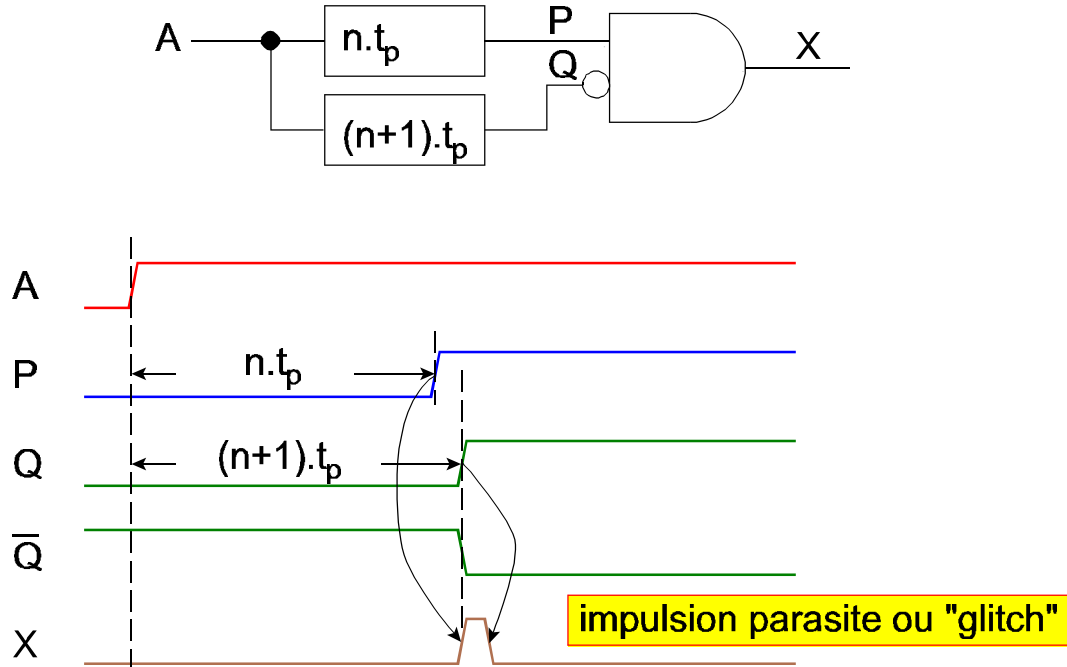


Le circuit présenté ici fournit une impulsion très courte à chaque flanc montant du signal d'horloge.

Il est basé sur l'exploitation volontaire du temps de propagation de l'inverseur pour créer un décalage entre CLK=V et U à l'entrée du NAND. On obtient ainsi une impulsion X dont la largeur est de l'ordre du temps de propagation et que l'on peut éventuellement réinverser pour obtenir un signal E actif à l'état HI.

Il s'agit ici d'une idée de principe. En pratique cette détection de flancs est réalisée par des transistors à l'intérieur des circuits intégrés et doit être conçue pour fonctionner sur toute la gamme de température, de tension d'entrée, et de temps de montée prévus pour la famille logique concernée.

Condition de course (bis)



Insistons sur le fait que le circuit précédent était une exploitation volontaire du temps de propagation.

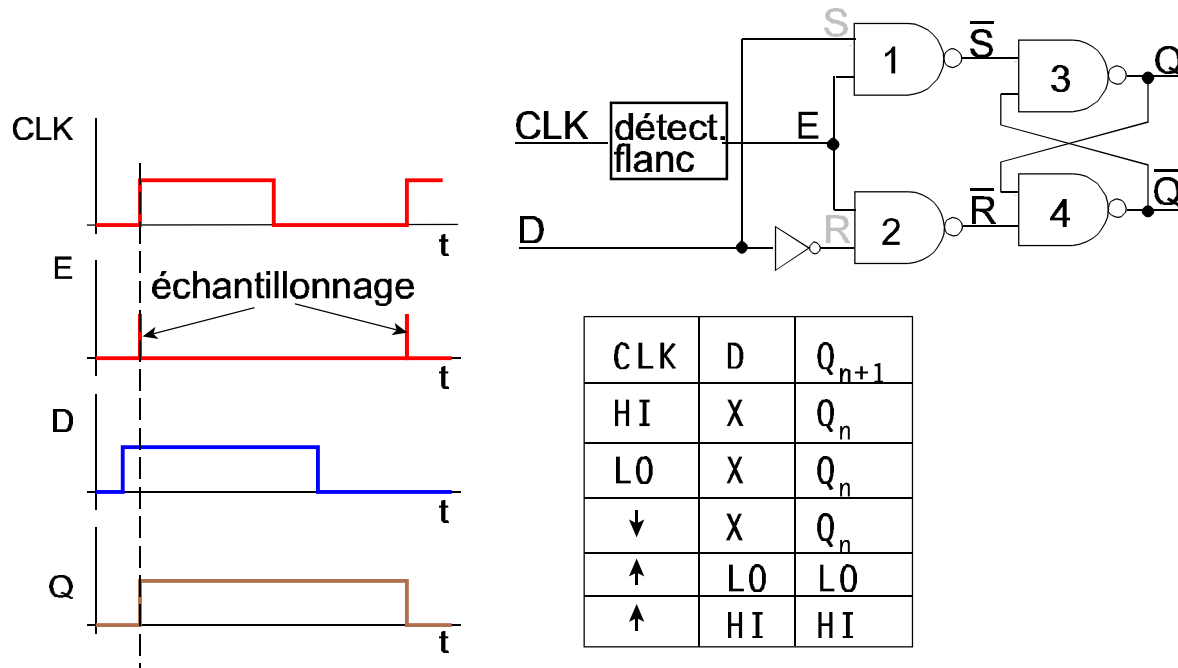
Une même situation peut se produire de manière involontaire lorsque l'on recombine dans une porte logique des signaux dont la transition est gouvernée par le même signal d'entrée, mais qui sont passés par une suite de traitements logiques dont les temps de propagation diffèrent légèrement.

Dans ce cas, on obtient à la sortie X une impulsion logique parasite ou "glitch" qui est potentiellement dangereuse, en particulier parce qu'elle peut être "attrapée" par un bistable et le faire basculer indûment.

La logique asynchrone doit donc être conçue avec soin pour éviter de telles situations. Il faut envisager tous les délais, avec leur plage de variation minimale et maximale (due à la dispersion de fabrication et aux variations de température), afin de trouver le cas le plus défavorable. C'est une opération fastidieuse, qui, dans les cas complexes comme la conception de circuits intégrés, doit être assistée par des programmes de simulation.

Nous verrons que la logique synchrone permet une plus grande sûreté de fonctionnement.

Bistable D sur flanc ou D flip-flop



En reprenant le D latch et en reliant l'entrée E à la sortie d'un détecteur de flanc actionné par le signal d'horloge CLK, on obtient le D flip-flop .

Son fonctionnement idéalisé est le suivant:

- à chaque flanc montant de l'horloge, le détecteur de flanc ouvre brièvement les portes d'entrée du bistable pour prendre en considération l'entrée D, qui est transférée sur Q
- juste après le flanc montant de CLK, la sortie est verrouillée et n'est susceptible de changer qu'au flanc montant suivant de l'horloge CLK

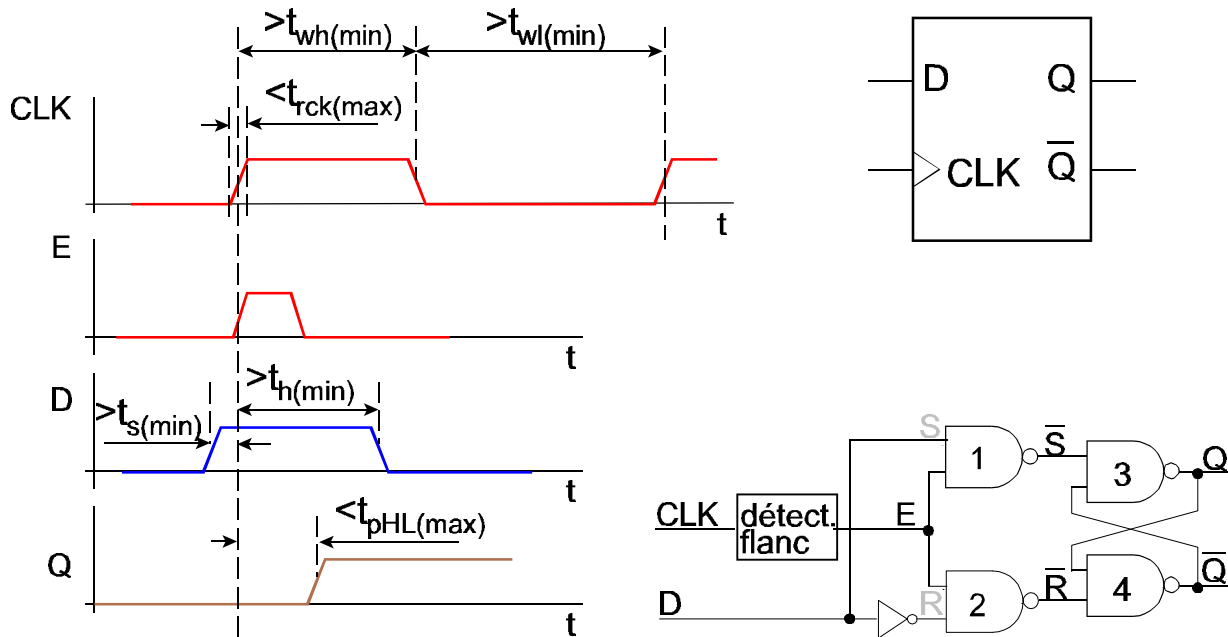
Le flanc montant de l'horloge réalise donc un échantillonnage de l'entrée Q.

La table de vérité montre explicitement que le flanc montant de CLK est l'événement de déclenchement du bistable.

Le fonctionnement correct du D flip-flop est assorti d'un ensemble de conditions temporelles dont nous allons maintenant évoquer les raisons.

Bistable D sur flanc ou D flip-flop

contraintes temporelles



©ULB ELMITEL

Multivibrateurs

ELEC344 MultiVibr_d19.shw 26/11/01 01:26:36

35

Les premières contraintes sont relatives au signal d'horloge CLK lui-même, et sont destinées à assurer un fonctionnement correct du détecteur de flanc :

- temps de montée inférieur à la limite $t_{r(max)}$
- largeur d'impulsion supérieure à la limite $t_{wh(min)}$

Peuvent encore s'y ajouter :

- la fréquence maximum d'horloge compatible avec la technologie utilisée; cette fréquence est soit mentionnée explicitement soit indirectement par la largeur d'impulsion inactive de l'horloge $t_{wl(min)}$, qui ajoutée à la largeur d'impulsion active minimum $t_{wh(min)}$ donnera la période minimum d'horloge
- le rapport cyclique (t_{wh}/T) qui doit parfois rester proche de 50%, notamment dans les circuits qui exploitent en interne les deux flancs d'horloge.

Ensuite, les contraintes déjà évoquées lors de l'étude du RS latch et du D latch permettront un fonctionnement correct du bistable RS interne. On rappellera :

- une largeur courte, mais suffisante, de l'impulsion E, réalisée par le détecteur de flanc
- un temps d'application minimum "setup time" sur l'entrée D, pour que ce soit bien le flanc d'horloge qui provoque le basculement éventuel du bistable
- un temps de maintien ("hold time") de l'entrée D pour éviter une course entre le changement d'état de D et la désactivation de E

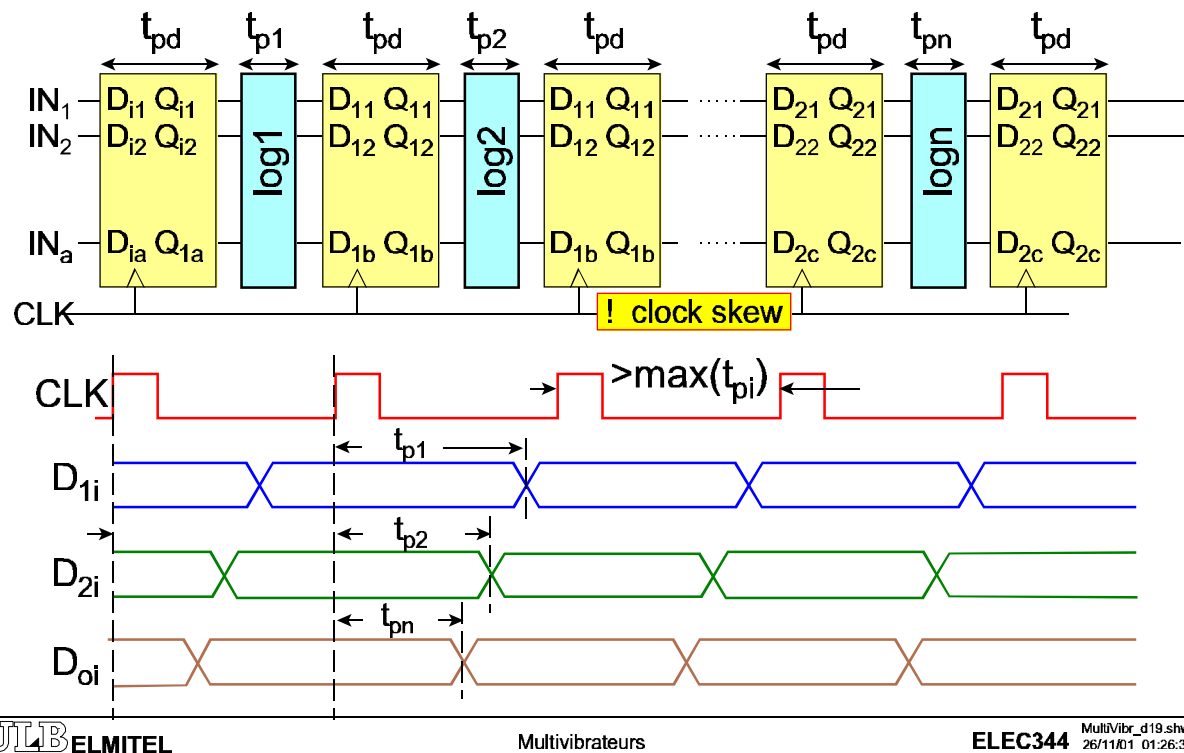
Enfin, si toutes ces conditions sont respectées, le temps de propagation jusqu'à la transition de la sortie sera garanti inférieur aux valeurs $t_{pHL(max)}$ et $t_{pLH(max)}$

Le point de référence de mesure des délais est le flanc actif de l'horloge.

Applications des bistables D

- ▶ systèmes à micro-processeurs
 - ◆ mémoires
 - ◆ registres internes
 - ◆ registres d'entrée-sortie
- ▶ systèmes logiques
 - ◆ réalisation de logiques synchrones
 - ◆ compteurs
 - ◆ registres à décalage

Logique synchrone



©ULB ELMITEL

Multivibrateurs

ELEC344 MultiVibr_d19.shw 26/11/01 01:26:36

39

Soit un circuit logique possédant a entrées et p sorties et effectuant un traitement nécessitant n blocs logiques en cascade $\log_1 \dots \log_n$.

Ces blocs logiques sont combinatoires et asynchrones et présentent éventuellement à leur sortie des impulsions parasites gênantes.

Soient $t_{p1} \dots t_{pn}$ les temps de propagation respectifs de ces blocs, définis comme le délai pour obtenir des sorties stables à partir du moment où les entrées sont stables. Les blocs étant combinatoires, leur sortie ne dépend que de l'état des entrées.

Interposons un jeu de bistables D pour encadrer chacun des blocs logiques et soit t_{pd} le temps de propagation de ces bistables. Tous les bistables sont pilotés par la même entrée d'horloge.

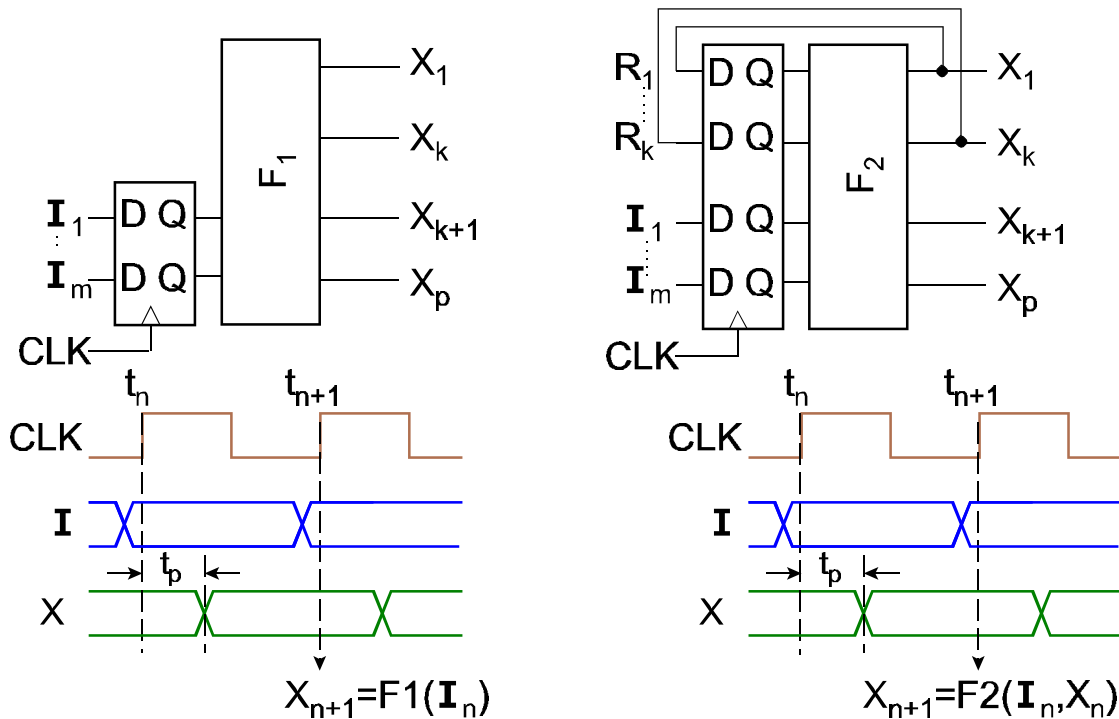
A condition que la période d'horloge soit supérieure au temps de propagation du bloc le plus lent (auquel il faut ajouter t_{pd}), les signaux de sortie de tous les blocs sont stables au moment du flanc actif de l'horloge et il n'y a donc pas de condition de course entre signaux d'entrée des différents blocs. Le prix à payer pour la synchronisation et le fonctionnement sans aléas est un certain ralentissement, dû à la marge de sécurité entre la période d'horloge et le temps de propagation le plus long.

Le signal logique progresse donc d'étage en étage à chaque coup d'horloge; on donne généralement aux sorties logiques un indice correspondant au numéro du coup d'horloge.

REM :

- la vitesse de propagation d'un flanc d'horloge est de l'ordre de 20cm/ns dans les pistes de circuit imprimé; dans les logiques rapides (temps de propagation de 1 ou 2 ns) les différences de trajet dans la distribution de l'horloge sur plusieurs dizaines de cm introduisent donc un décalage d'horloge appelé CLOCK SKEW. Il importe donc d'égaliser les trajets.
- les réflexions sont particulièrement néfastes sur les pistes d'horloge

Logique séquentielle



©ULB ELMITEL

Multivibrateurs

ELEC344 MultiVibr_d19.shw
26/11/01 01:26:36

41

Prenons un circuit combinatoire réalisant la fonction F_1 avec un temps de propagation t_p . Munissons le d'un jeu de bistables D en entrée pour échantillonner les signaux d'entrée $I_1 \dots I_m$. Plaçons-nous au nième coup d'horloge (instant t_n). Un délai t_p plus tard, le vecteur de sortie X sera stable. La période de l'horloge étant supérieure à t_p , la valeur de X au prochain flanc d'horloge sera donc stable et pourra être échantillonnée proprement par un éventuel étage logique en aval. Cette valeur est

$$X_{n+1} = F_1(I_n)$$

Le passage à la logique séquentielle se fait par une rétroaction partielle des sorties $X_1 \dots X_k$ vers les entrées supplémentaires $R_1 \dots R_k$.

Le fonctionnement est identique au cas précédent et l'on a

$$X_{n+1} = F_2(I_n, R_n)$$

or les rétroactions R_n ne sont rien d'autre que les sorties X_n au même moment t_n , et par conséquent

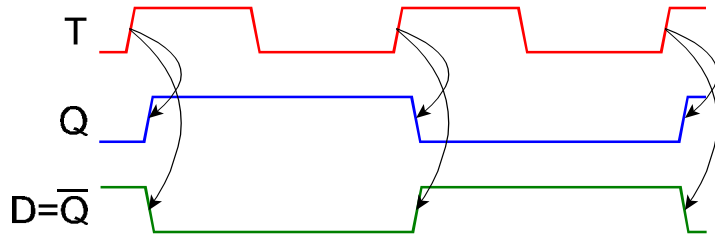
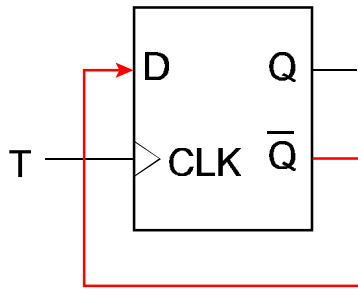
$$X_{n+1} = F_2(I_n, X_n)$$

La sortie X du circuit dépend donc à la fois de l'entrée I et de la sortie X au coup d'horloge précédent, ce qui est par définition de la logique séquentielle. Les bistables D fournissent ici l'élément de mémoire indispensable qui retient l'état précédent du système.

Si l'on met plusieurs étages de bistables en cascade, on peut mémoriser des états antérieurs des entrées et des sorties aux instants t_{n-1} , t_{n-2} ,et les faire intervenir dans la fonction logique.

42

Bistable T ou "Toggle flip-flop"



"toggle" = changement d'état à chaque coup d'horloge

le bistable T divise la fréquence d'entrée par 2

Une rétroaction très simple consiste à connecter Q' sur D.

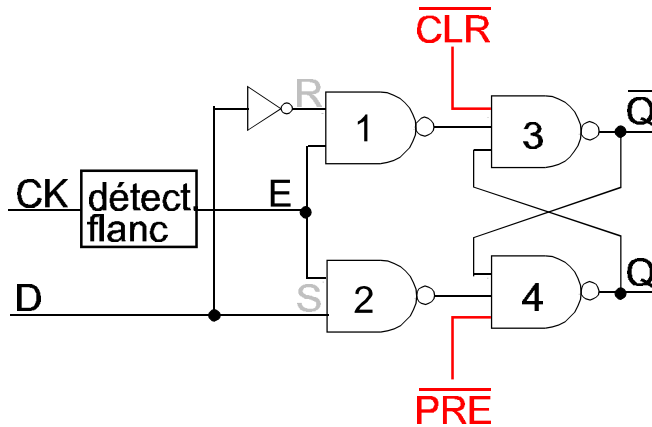
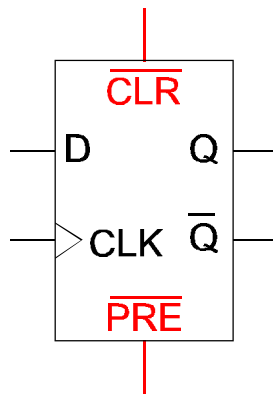
Supposons l'état initial $Q=LO$, $Q'=HI$ avant le premier flanc d'horloge actif.

- au premier flanc d'horloge, on échantillonne $D=Q'=HI$ donc $Q=HI$ et $Q'=LO$ après le temps de propagation
- au flanc suivant, on échantillonne $D=Q'=LO$ donc $Q=LO$ et $Q'=HI$

et ainsi de suite....

Le bistable T (pour "toggle") change donc d'état à chaque flanc actif d'horloge (ici au flanc montant), ce qui en fait un diviseur par 2 : la fréquence de la sortie Q vaut la moitié de celle du signal d'entrée T.

Bistable D : entrées supplémentaires



PRE (preset) ou S_D (set direct)
CLR (clear) ou R_D (reset direct)

	PRE	CLR	CLK	D	Q_{n+1}
preset	LO	HI	X	X	HI
clear	HI	LO	X	X	LO
mem	HI	HI	LO	X	Q_n
	HI	HI	↑	LO	LO
	HI	HI	↑	HI	HI

Il peut être utile d'imposer l'état de la sortie du bistable indépendamment de l'état de l'entrée D.

Cela permet, par exemple

- d'initialiser un bistable T dans un état ou dans l'autre
- de bloquer le fonctionnement sans devoir ajouter une porte pour masquer l'horloge.

Pour ce faire, on ajoute deux bornes supplémentaires :

- PRE (PREset), appelée aussi Set Direct (SD), qui met inconditionnellement $Q \rightarrow HI$ et $Q' \rightarrow LO$
- CLR (CLear), appelée aussi Reset Direct (RD), qui met inconditionnellement $Q \rightarrow LO$ et $Q' \rightarrow HI$

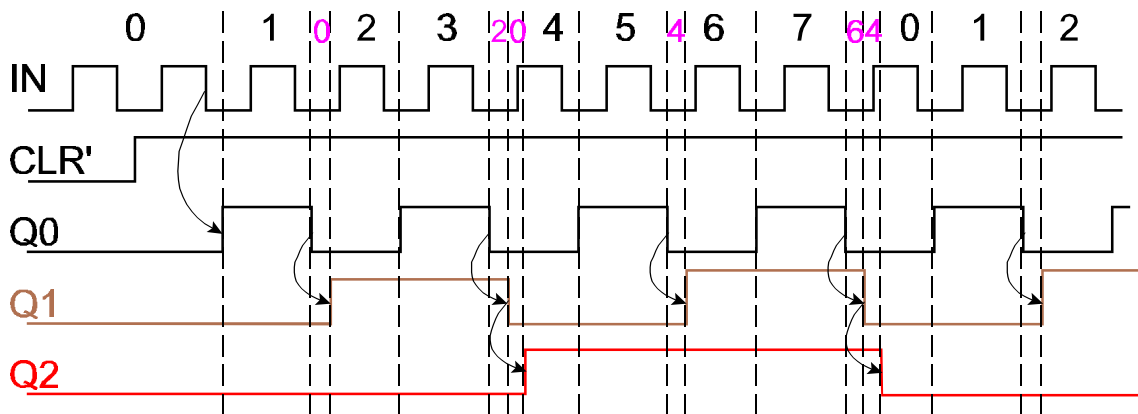
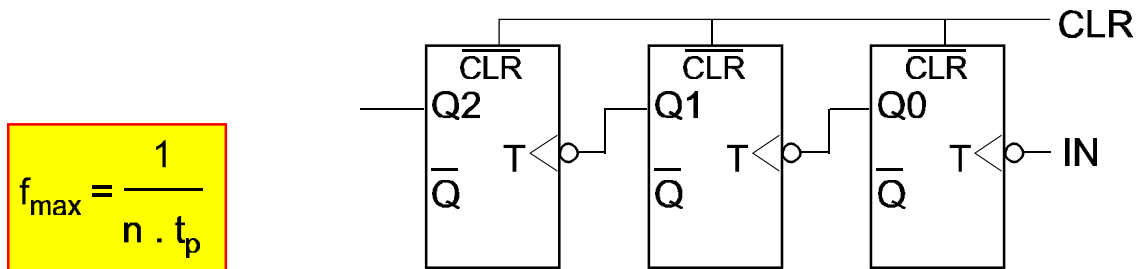
On a donc un mélange dans le même boîtier d'un RS et d'un D flip-flop.

Ces entrées peuvent être :

- asynchrones, c'est-à-dire agir indépendamment de l'horloge (cas de la figure)
- synchrones avec l'horloge, auquel cas elles doivent agir sur l'étage d'entrée pour être prises en compte au moment où l'horloge est active.

Compteur asynchrone binaire

$$f_{\max} = \frac{1}{n \cdot t_p}$$



©ULB ELMITEL

Multivibrateurs

ELEC344 MultiVibr_d19.shw
26/11/01 01:26:36

47

Le compteur le plus simple que l'on puisse réaliser est le compteur binaire asynchrone. Pour ce faire, on cascade des bistables T en connectant le signal d'entrée à la borne T du premier bistable, puis chaque sortie Q à l'entrée T du bistable suivant. On a choisi ici des bistables actifs sur le flanc descendant de l'horloge.

L'utilité de l'entrée CLR' apparaît : on l'active pour mettre toutes les sorties à LO au départ.

Dès que l'on désactive CLR'

- le premier flanc descendant du signal d'entrée provoque le basculement de Q0 → HI
- le flanc suivant fait passer Q0 → LO, ce qui donne un coup d'horloge au bistable suivant et Q1 → HI
- et ainsi de suite ...

On parle de compteur binaire, puisque chaque étage divise par deux. Le nombre {Q2,Q1,Q0} code en binaire le nombre de flancs actifs du signal d'entrée depuis la désactivation du CLR.

Il s'agit d'un compteur MODULO 2^N , où N est le nombre d'étages du compteur.

La dénomination asynchrone vient du fait que les sorties ne changent pas simultanément, à cause du temps de propagation de chaque étage:

- Q0 est en retard de t_p sur l'horloge
- Q1 est en retard de t_p sur Q0 et donc de $2t_p$ sur l'horloge
- ...
- Qn est en retard de $(n+1)t_p$ sur l'horloge

Les compteurs asynchrones présentent donc deux défauts liés à leur principe :

- la fréquence maximum d'entrée décroît avec le nombre d'étages; en effet, un fonctionnement correct impose que :

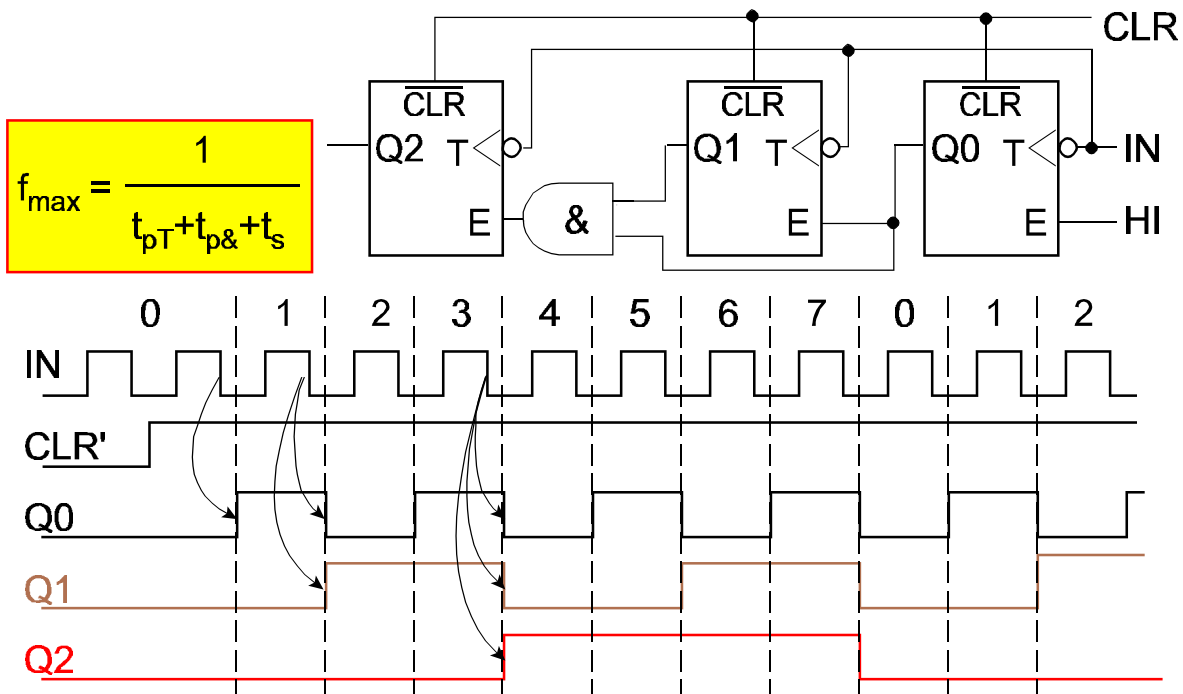
$$(n+1) t_p < T \Rightarrow f_{\max} = 1/(n \cdot t_p)$$

- les états corrects de sortie sont entrecoupés de transitoires, ce que la figure montre bien par la traduction en décimal de la sortie binaire

Le câblage de tels compteurs est par contre remarquablement simple.

48

Compteur synchrone binaire



©ULB ELMITEL

Multivibrateurs

ELEC344 MultiVibr_d19.shw
26/11/01 01:26:36

49

Pour obtenir un fonctionnement synchrone, il faut compliquer le montage et disposer de bistables J-K (voir plus loin) ou de bistables T munis d'une entrée d'activation E supplémentaire. Le synchronisme est assuré par la distribution de l'horloge à toutes les entrées T. On doit alors ajouter des portes supplémentaires pour inhiber le comptage pour les bistables qui ne doivent pas s'incrémenter. Si l'on examine la séquence binaire normale à 3 bits

000
001
010
011
100
101
110
111

on constate:

- que le bit de poids faible change d'état à chaque étape
- qu'un bit de poids fort change lorsque tous les bits de poids plus faible sont à 1

On en déduit le schéma présenté ci-dessus. L'entrée E doit être active (HI) pour que le bistable T change d'état.

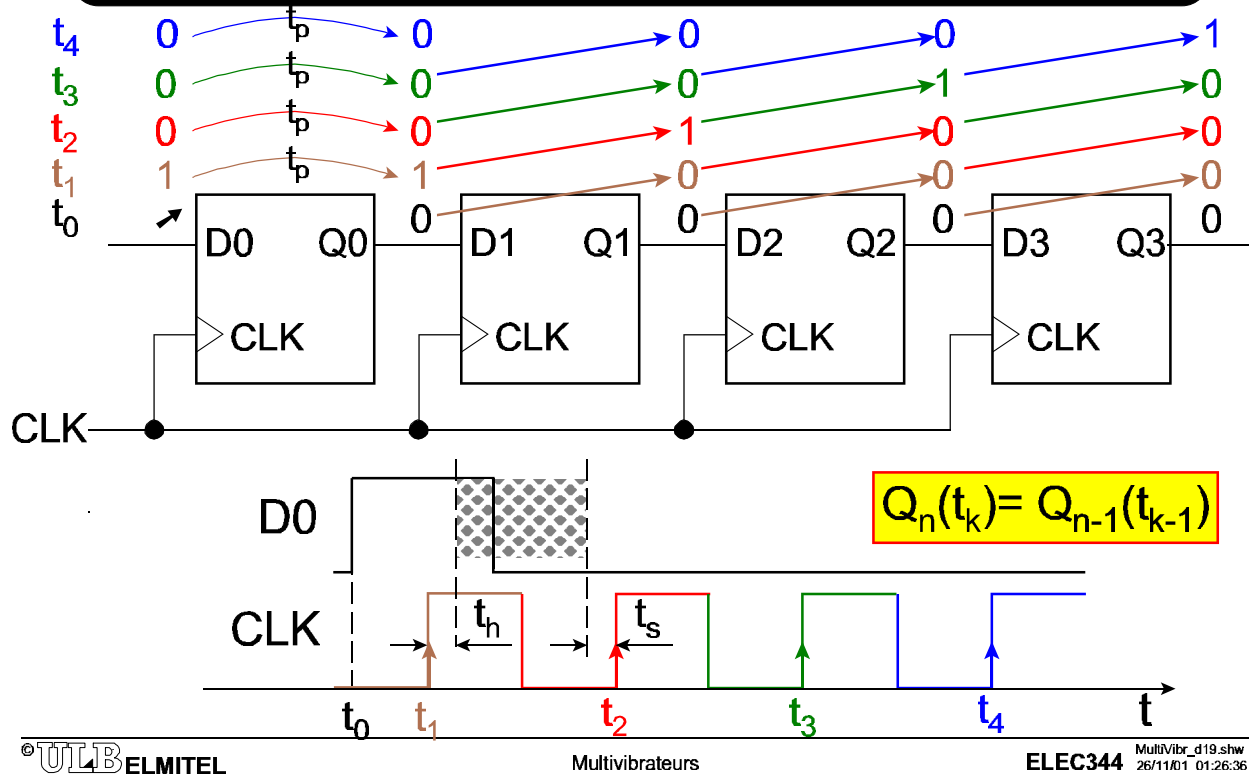
Au prix d'un câblage plus complexe, le compteur synchrone est plus avantageux que l'asynchrone:

- les états parasites ont disparu, puisque toutes les sorties changent d'état "simultanément" (c'est-à-dire à la dispersion près sur le temps de propagation)
- la fréquence maximum du compteur ne dépend plus du nombre d'étages; la contrainte est que le nouveau flanc d'horloge ne peut pas intervenir avant que les entrées E ne soient stabilisées, soit une période maximum valant

$$T < t_{pT} + t_{p\&} + t_s$$

où t_{pT} , $t_{p\&}$ et t_s sont respectivement le temps de propagation du T, de la porte AND et le "setup time" du bistable

Registre à décalage : principe



51

Connectons une série de bistables D en cascade avec une horloge commune. Il s'agit donc d'un circuit synchrone :

- à l'instant initial t_0 , appliquons une impulsion positive sur l'entrée D_0 du 1er bistable, et supposons toutes les sorties à 0 (par exemple par une action antérieure sur les entrées CLR non représentées pour ne pas alourdir le schéma)
- en t_1 se produit le premier flanc d'horloge actif; un temps de propagation t_p plus tard, la sortie Q_0 recopie D_0 et passe à 1; les autres sorties $Q_{1..3}$ recopient les entrées $D_{1..3}$ c'est-à-dire aussi l'état des sorties $Q_{0..2}$ échantillonnées au flanc montant de l'horloge; $Q_{1..3}$ restent donc à 0
- l'impulsion sur D_0 retombe à 0
- en t_2 se produit le 2ème flanc d'horloge actif; un temps de propagation t_p plus tard, la sortie Q_0 recopie D_0 et retombe à 0; Q_1 recopie D_1 et passe à 1, et $Q_{2..3}$ restent à 0

Au fur et à mesure des coups d'horloge, le "1" initialement présent se décale vers la droite et apparaît à la sortie suivante, d'où le nom de registre (=ensemble de bits mémorisés dans un groupe de bistables) à décalage.

Plus généralement, quelle que soit l'information présente dans le registre, elle se décale à droite à chaque coup d'horloge et

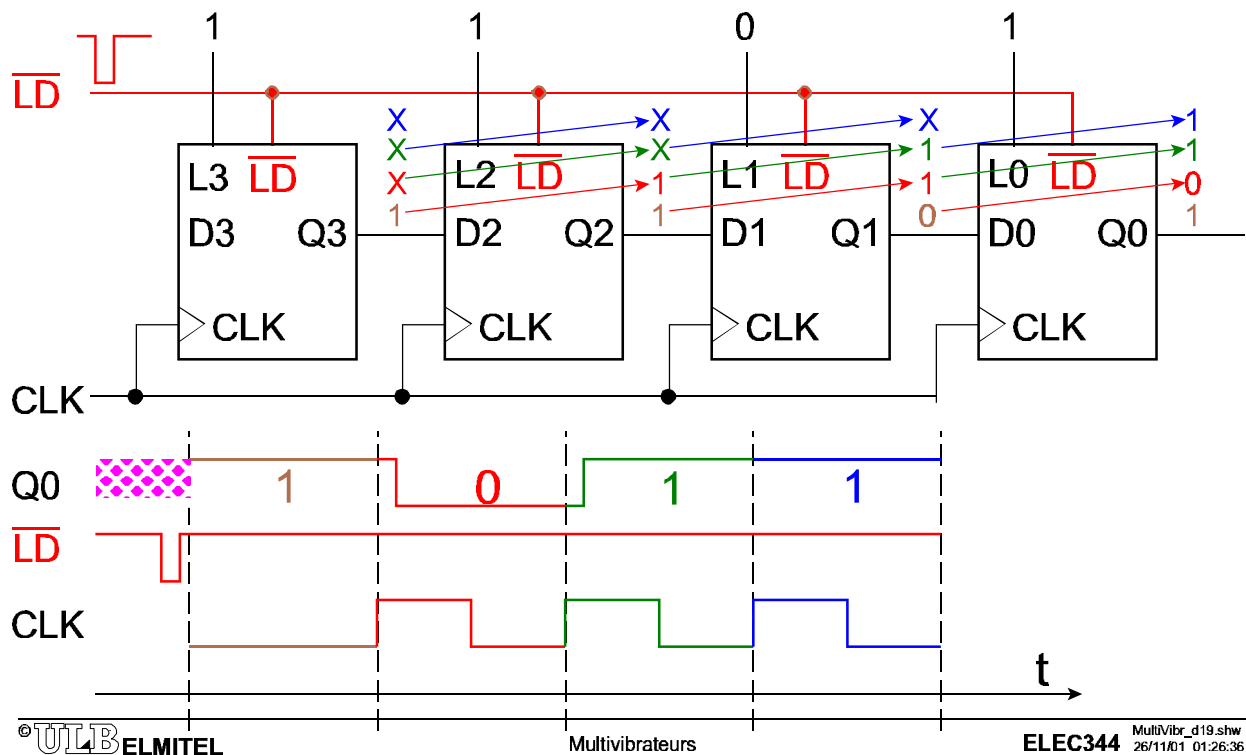
$$Q_n(t_k) = Q_{n-1}(t_{k-1}) \quad \text{où } n \text{ est le rang du bistable et } k \text{ le numéro du coup d'horloge}$$

REM1: les changements d'état du signal D_0 doivent respecter les temps d'activation "setup" et de maintien "hold" par rapport au flanc d'horloge

REM2: si l'on boucle la sortie Q_3 sur l'entrée D_0 et que l'on introduit un "1" dans le registre via les entrées PREset et Clear des bistables, on obtient un anneau avec une circulation du "1" appelé compteur de Johnson. Un tel circuit est utile par exemple pour sélectionner tour à tour les différents digits d'un afficheur et les rafraîchir périodiquement un par un (affichages multiplexés)

52

Registre à décalage : conversion //→ série



53

D'une manière analogue aux entrées PRESET et CLEAR, le bistable D peut être muni d'une entrée de chargement "Load" permettant d'imposer l'état des sorties en fonction du bit présent sur l'entrée D. L'entrée LD peut être totalement asynchrone, ou le chargement peut être synchronisé sur l'horloge, suivant le type de circuit.

La figure représente ici un registre à décalage de 4 bits obtenu par mise en cascade de 4 bistables D. On peut donc y précharger un mot de 4 bits (par exemple 1011) via les entrées D et une impulsion simultanée sur les 4 entrées LD.

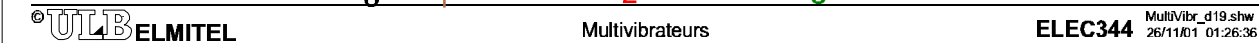
On applique ensuite 4 périodes d'horloge sur l'entrée CLK. On retrouve à la sortie Q₀ une séquence de bits qui est la lecture du mot préchargé, en commençant par le bit le moins significatif.

On a réalisé une conversion parallèle → série.

Le registre à décalage est donc à la base de toutes les transmissions en série. Un processeur qui veut transmettre en caractère à distance placera un mot de 8 ou 16 bits codant ce caractère sur les entrées de chargement parallèle, puis le registre à décalage se chargera la conversion en une séquence de bits, qui seront mis en forme et envoyés sur le canal de communication (ligne téléphonique, transmission radio,)

54

Registre à décalage : conversion série //



Reprenons la séquence de bits créée à la dia précédente et appliquons-la à l'entrée d'un registre à décalage.

L'état du registre au départ est indifférent. A chaque coup d'horloge, le bistable d'entrée prend la valeur du bit suivant dans le flux de données entrant et le bit précédent est décalé dans le bistable de rang inférieur.

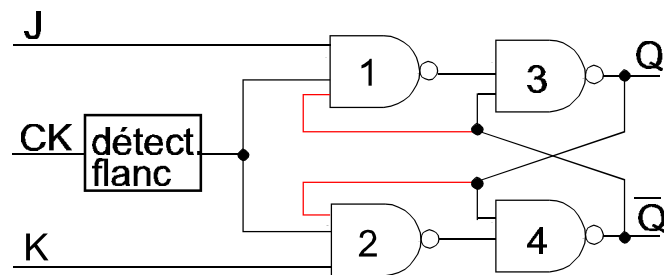
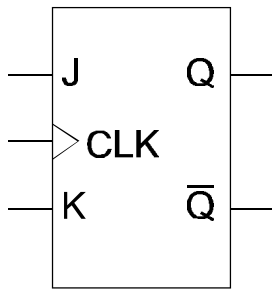
Pour un registre à n bits, on retrouve après n coups d'horloge un mot que l'on peut lire sur les sorties $Q_{n-1} \dots Q_0$. On a donc réalisé la conversion de données de **série vers parallèle**. L'application la plus courante est donc la réception de données sur une ligne de transmission série.

Cette conversion suppose évidemment que l'on dispose d'une horloge CLK à la même fréquence que les données et dont la phase est telle que les temps de "setup" et de "hold" sont respectés (idéalement, les flancs actifs doivent tomber en plein milieu des bits, comme indiqué sur la figure).

Deux cas sont possibles :

- transmission synchrone : l'horloge est fournie en même temps que la donnée par l'émetteur
- transmission asynchrone : l'horloge est fabriquée localement par le récepteur à une fréquence connue d'avance et le calage de phase s'opère par un bit supplémentaire toujours inactif entre chaque mot (appelé "stop bit") suivi d'un bit spécial toujours actif appelé "start bit". On obtient ainsi un flanc de référence temporelle et on décale les flancs d'horloge d'un demi-temps de bit pour échantillonner les données proprement (voir chapitre sur les transmissions en série).

Bistable J-K



J	K	CLK	Q_{n+1}	
LO	LO	X	Q_n	mem
HI	LO	$\uparrow$	HI	set
LO	HI	$\uparrow$	<u>LO</u>	reset
HI	HI	$\uparrow$	$\overline{Q_n}$	toggle

Un autre dérivé du bistable S-R est le bistable J-K. La différence est l'ajout de deux rétroactions :

- de Q vers la porte NAND2 où agit l'entrée K
- de Q' vers la porte NAND1 où agit l'entrée J

On en déduit, si l'on part des l'état inactif du bistable ($Q=LO, Q'=HI$), que :

- la porte NAND2 est bloquée et que toute action sur K ou sur l'horloge est inopérante
- la porte NAND1 fonctionne comme dans le flip-flop R-S : si $J=HI$ et que l'on a un flanc actif de l'horloge, la sortie du NAND1 passe à LO et active le bistable de sortie. J joue donc le rôle de l'entrée S.
- dès que $Q=HI$ et $Q'=LO$, la porte NAND1 est verrouillée et J est inopérant, alors que l'activation de K remettra le bistable de sortie dans l'état inactif. K joue donc le rôle de l'entrée R.

La table de vérité du J-K est donc identique à celle du R-S pour les 3 combinaisons licites de celui-ci. Par contre, l'activation simultanée de J et K est autorisée, alors qu'elle est interdite dans le R-S.

Si $J=K=HI$, le mécanisme de rétroaction que nous venons d'expliquer fait en sorte que seule l'entrée utile pour faire changer le bistable d'état est prise en compte, alors que l'autre est ignorée:

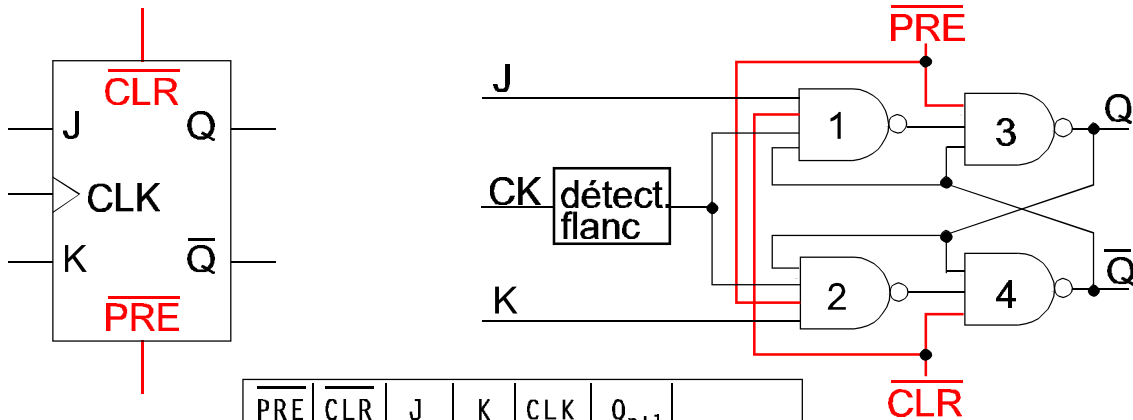
- si $Q=LO$, $J=K=HI$ est ramené à $J=HI$, K est inhibé
- si $Q=HI$, $J=K=HI$ est ramené à $K=HI$, J est inhibé

Le J-K change alors d'état à chaque flanc d'horloge.

On en déduit qu'un J-K dont on connecte les deux entrées ensemble devient un bistable T, dont $J=K$ est l'entrée d'activation. Le J-K est donc la base des échelles de comptage synchrones.

En se basant sur les temps de propagation, on démontrera à titre d'exercice que le fonctionnement correct en "toggle" devient anormal si l'on connecte directement l'horloge aux portes NAND1 et NAND2 sans interposer de détecteur de flanc. La structure de J-K présentée ici est donc exclusivement "edge-triggered".

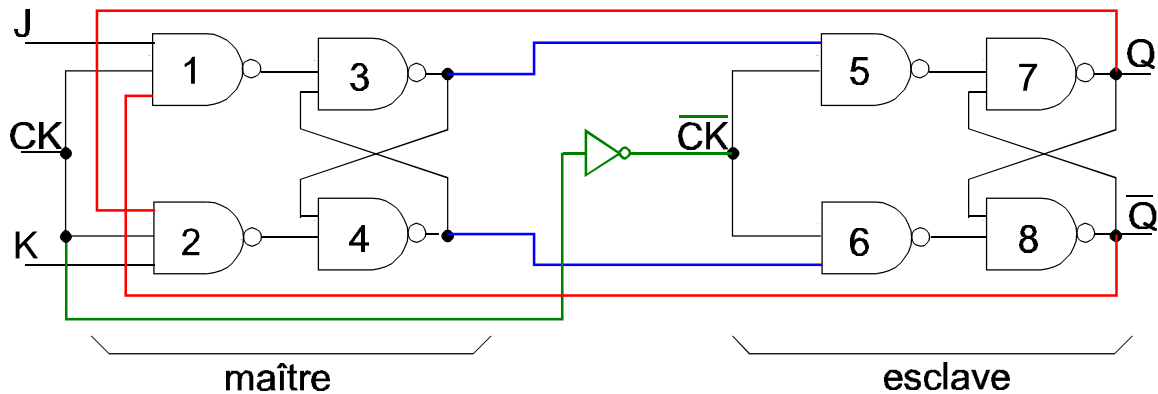
Bistable J-K avec PRE et CLR



$\overline{\text{PRE}}$	$\overline{\text{CLR}}$	J	K	CLK	Q_{n+1}	
LO	HI	X	X	X	HI	preset
HI	LO	X	X	X	LO	clear
HI	HI	LO	LO	X	Q_n	mem
HI	HI	HI	LO	$\uparrow$	HI	set
HI	HI	LO	HI	$\uparrow$	$\overline{\text{LO}}$	reset
HI	HI	HI	HI	$\uparrow$	$\overline{Q_n}$	toggle

Le J-K existe également avec les entrées PREset et CLear permettant d'imposer l'état de sortie de manière asynchrone. Ces entrées agissent à la fois sur le bistable de sortie et sur l'étage d'entrée, où elle inhibent J,K et l'horloge.

Bistable J-K maître-esclave



J	K	CLK	Q_{n+1}	
LO	LO		Q_n	mem
HI	LO		HI	set
LO	HI		LO	reset
HI	HI		$\overline{Q_n}$	toggle

Pour faire un J déclenché par une impulsion et non par un flanc, nous avons vu que la structure précédente ne fonctionne pas correctement.

Une autre structure autorise une commande synchronisée par des impulsions: le J-K maître-esclave ou "master-slave J-K".

On part de deux "latches" R-S, que l'on met en cascade; on ajoute un inverseur sur l'horloge d'entrée pour commander l'esclave. Enfin, on place les rétroactions des sorties Q et Q' de l'esclave vers les entrées du maître, afin de pouvoir créer la fonction "toggle" si les deux entrées sont actives, comme l'exige la table de vérité du J-K.

Le fonctionnement est alors le suivant:

- au flanc montant de l'impulsion d'horloge les portes NAND1 et NAND2 s'ouvrent et libèrent le fonctionnement du maître, alors que NAND5 et NAND6 se ferment et bloquent celui de l'esclave. Le bistable maître bascule dans l'état imposé par J et K, en tenant compte des rétroactions venant des sorties.
- au flanc descendant de l'impulsion d'horloge, c'est le maître qui est inhibé et l'esclave peut recopier et mémoriser l'état du maître.

Pour un tel bistable, les contraintes de "setup time" et de "hold time" sont impératives: les signaux J et K doivent être stables pendant la durée de l'impulsion d'horloge. On s'efforcera donc de réaliser des impulsions d'horloge courtes. On démontrera, à titre d'exercice, que l'arrivée d'une impulsion parasite sur une entrée pendant que l'horloge CK est à l'état HI peut :

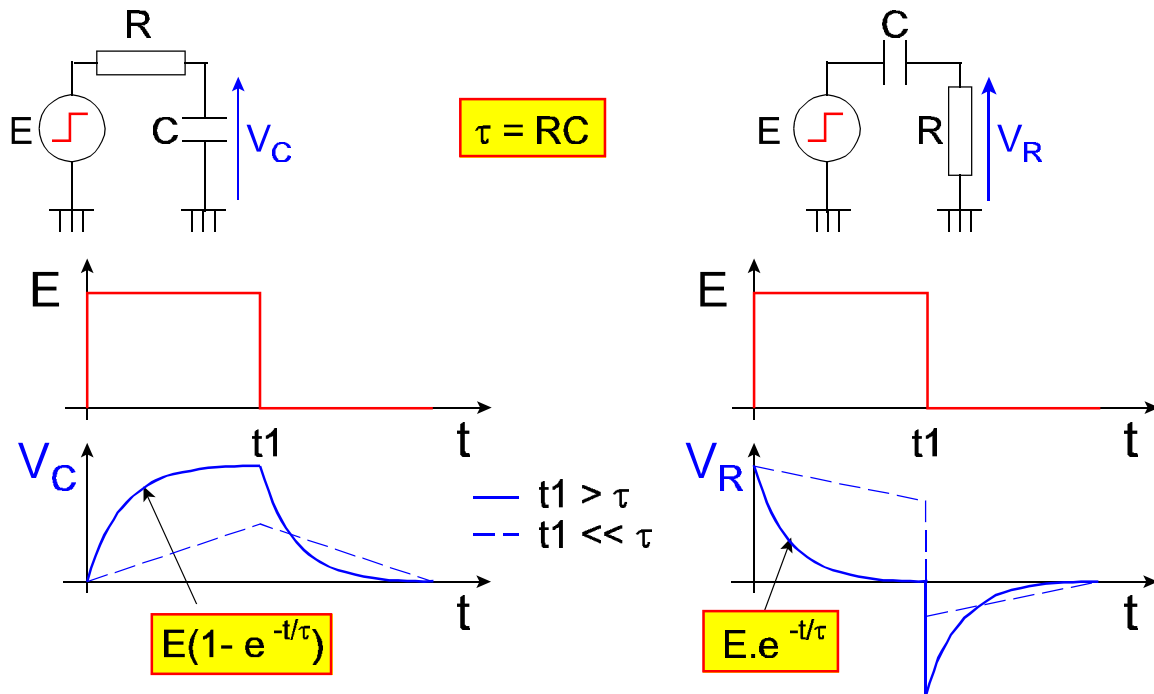
- prendre la table de vérité en défaut
- empêcher un basculement du bistable au coup d'horloge suivant; (ce phénomène est connu sous le nom de "1's catching")

Le J-K maître-esclave est aujourd'hui devenu nettement plus rare que le J-K "edge triggered",

Multivibrateurs : plan

- ▶ définitions
- ▶ bistables
- ▶ **monostables**
 - ◆ **rappels sur le RC**
 - ◆ **monostables réalisés à l'aide de portes**
 - ◆ **monostables intégrés**
- ▶ astables

Rappels sur le RC



©ULB ELMITEL

Multivibrateurs

ELEC344 MultiVibr_d19.shw
26/11/01 01:26:36

65

Puisque le monostable détermine lui-même la durée de son état métastable, il doit contenir sa propre base de temps. Celle-ci est le plus souvent constituée par un circuit RC commandé par les impulsions à la sortie des portes logiques.

La figure rappelle la réponse indicielle des deux types de circuits RC. La grandeur caractéristique est la "constante de temps" $\tau = RC$

A gauche le circuit dit "intégrateur" parce que, pour des temps $t \ll \tau$, la réponse indicielle croît linéairement (développement au premier ordre de l'exponentielle).

A droite le circuit dit "différentiateur" parce que, pour autant que

- le temps de transition (montée/descente) des impulsions soit faible devant τ
 - la durée t_1 des impulsions soit grande devant τ
- la sortie est une courte impulsion dont le signe est celui de la dérivée du signal d'entrée.

Lorsque τ devient grand par rapport à la durée t_1 de l'impulsion, on peut confondre l'exponentielle avec sa tangente à l'origine et la réponse est approximée par des segments de droite.

Par définition le circuit RC est en régime lorsque le condensateur a terminé sa (dé)charge, il n'y a donc plus de courant dans la résistance et la tension V_R est donc nulle.

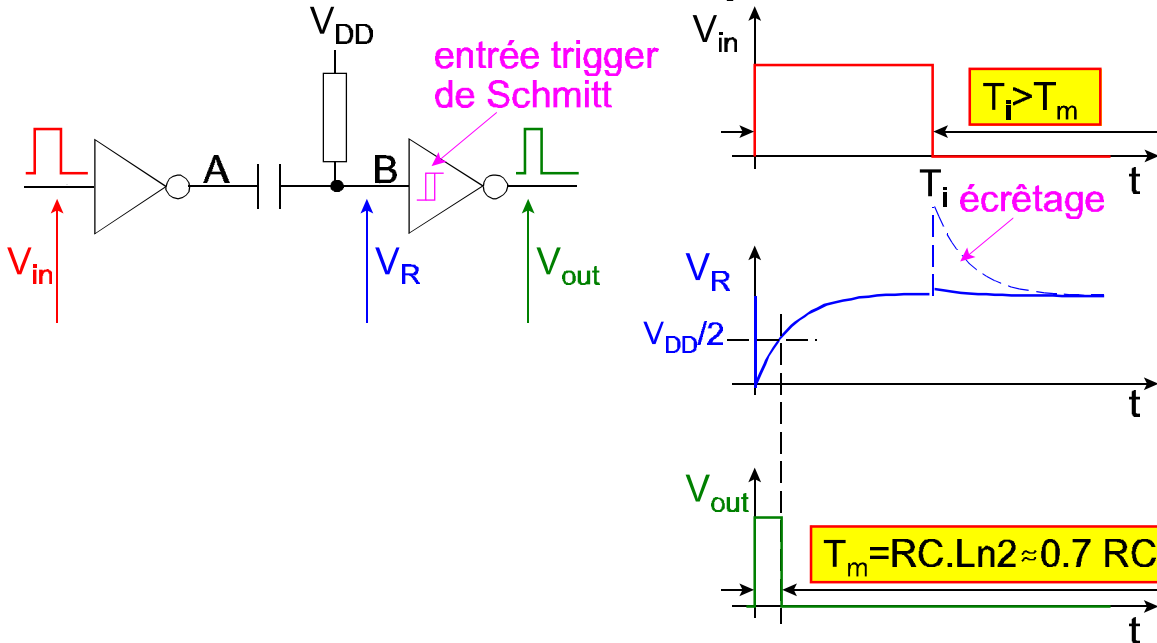
Pour rappel des ordres de grandeur de variation des réponses exponentielles :

- après τ , on est à 63% ($2/3$) de l'asymptote
- après 3τ , on est à 95% de l'asymptote
- après 5τ , on est à 99% de l'asymptote

66

Monostable simple

raccourcir une impulsion



Un monostable simpliste peut être constitué de deux inverseurs et d'un RC dérivateur.

L'état de repos vaut :

- $V_{in} = LO \Rightarrow V_A = HI$
- $V_B = HI$ à cause de R montée en "pull-up" $\Rightarrow V_{out} = LO$
- le condensateur est déchargé

En activant l'entrée V_{in} par une impulsion à l'état HI, on provoque un flanc descendant sur V_A , qui est transmis par le condensateur en V_B , donc V_{out} passe en HI.

Cet état est métastable, car le condensateur va se charger à travers R et la tension V_B remonte exponentiellement à V_{DD} .

En CMOS, la tension de transition est $V_{DD}/2$; lorsque V_B atteint cette valeur, l'inverseur de sortie rebascule dans l'état de repos.

La durée T_m de l'état métastable est donc le délai pour que

$$1 - \exp(-t/RC) = 1/2 \quad \text{soit} \quad T_m = RC \cdot \ln 2 \approx 0,7 RC$$

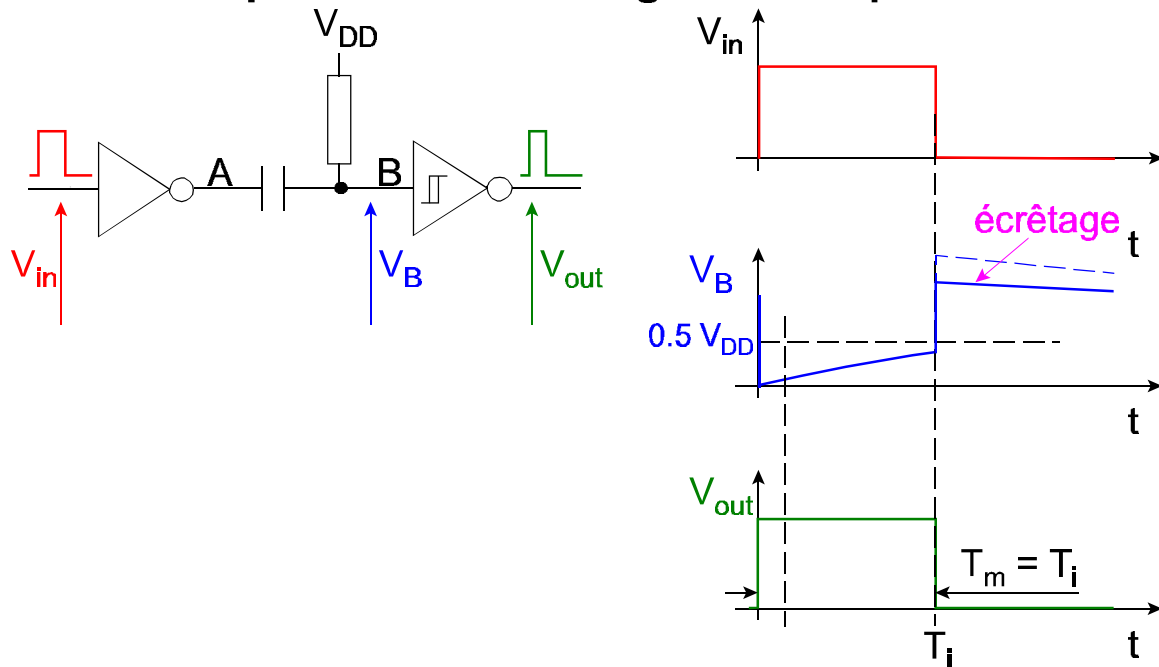
REM1: Pour fonctionner normalement, ce circuit doit être commandé par une impulsion dont la durée T_i est plus élevée que la constante de temps RC (voir dia suivante)

REM2: à la fin de l'impulsion d'entrée, on a un flanc montant sur V_A , donc sur V_B , qui tend vers $2 V_{DD}$. En CMOS, la diode de protection d'entrée va écrêter V_B à la valeur $(V_{DD} + V^*)$. On rappellera que le courant dans cette diode doit être limité. On consultera la notice relative à la famille logique utilisée et l'on ajoutera au besoin une résistance en série avec l'entrée de l'inverseur.

REM3: le temps de montée du signal V_B est en général inférieur au minimum requis par les circuits logiques. V_B passe trop longtemps au voisinage du seuil $V_{DD}/2$ où le gain est élevé, avec des risques d'oscillations. Il faut employer des portes avec hystérèse (Trigger de SCHMITT)

Monostable simple (2)

impossibilité d'allonger une impulsion



Si l'on veut augmenter la durée de l'état métastable, on accroît la valeur de R et/ou de C. Lorsque RC n'est plus suffisamment faible devant la durée de l'impulsion d'entrée T_i , ou, plus précisément, lorsque la durée de l'état métastable devient égale à T_i

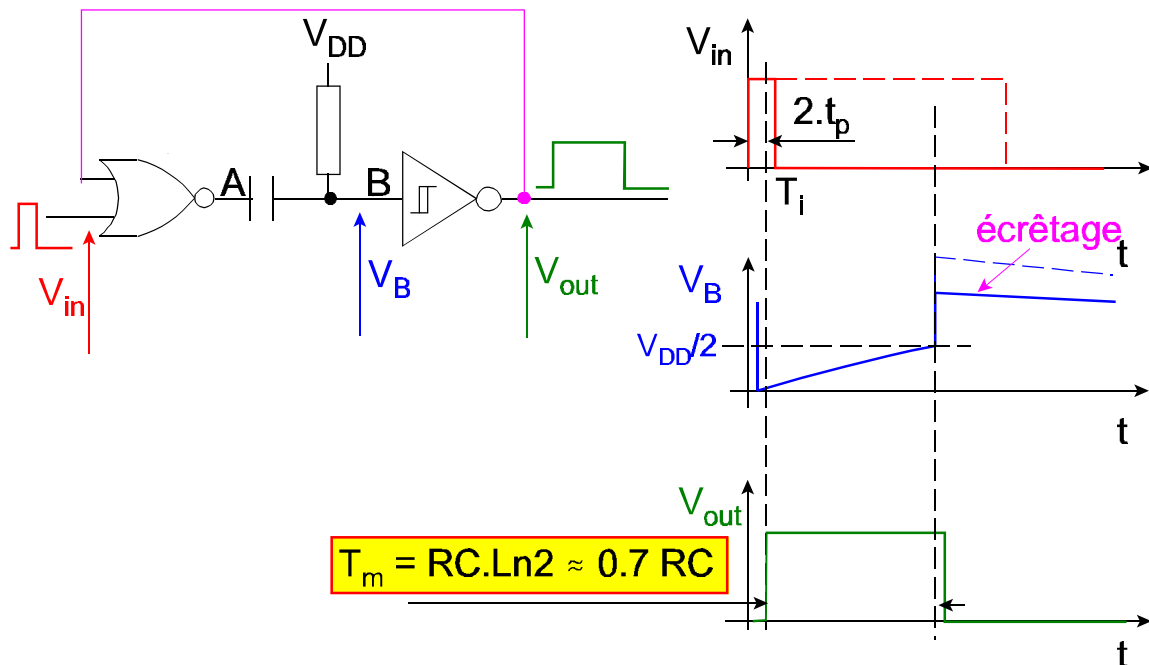
$$T_m = RC \ln 2 = T_i$$

alors le monostable cesse de fonctionner, car c'est l'impulsion d'entrée elle-même qui détermine la durée de l'impulsion de sortie, le RC ne joue plus aucun rôle.

Ce monostable simple ne peut donc que raccourcir l'impulsion d'entrée

$$T_m \leq T_i$$

ajoute d'une rétroaction



Si l'on veut utiliser un monostable pour prolonger une impulsion trop courte, il suffit de mettre une porte NOR à la place de l'inverseur d'entrée et de faire la rétroaction de l'état de sortie sur la deuxième entrée.

A l'état de repos, la sortie est en LO et la porte NOR laisse donc passer le signal d'entrée V_{in} . Deux temps de propagation après la montée de V_{in} , V_{out} passe en HI et vient donc verrouiller la porte NOR. Dès cet instant le signal V_{in} peut repasser à 0 sans affecter le fonctionnement du monostable.

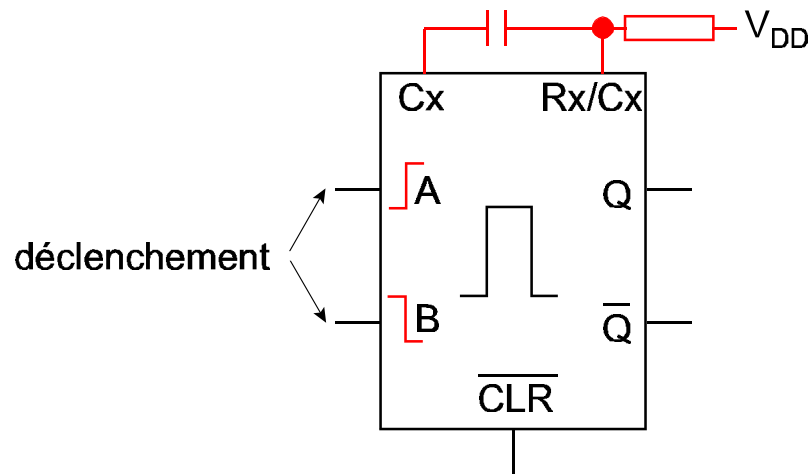
On peut ainsi accepter des impulsions d'entrée très courtes (de l'ordre de $2.t_p$). Ce montage est donc très sensible aux parasites et aux "glitches".

REM1 : si de nouvelles impulsions de déclenchement se produisent pendant l'état métastable, elles seront ignorées à cause de la rétroaction. Ce monostable est dit NON-REDECLANCHABLE ("non-retriggerable one-shot")

REM2 : si la durée de V_{in} est supérieure à l'état métastable, la rétroaction ne sert à rien, mais le monostable fonctionne parfaitement, de manière analogue au circuit plus simple vu précédemment. La durée de l'état métastable est donc devenue indépendante de l'impulsion de commande, qui peut prendre n'importe quelle durée.

REM3: La précision sur T_m est tributaire de la précision sur R et C et de leur coefficient de température. De plus T_m résulte de l'intersection entre une exponentielle et une droite horizontale. Si la constante de temps est élevée, ces courbes se coupent sous un angle faible ce qui donne une mauvaise précision. Il est difficile en pratique de faire mieux que 10%.

Monostables intégrés

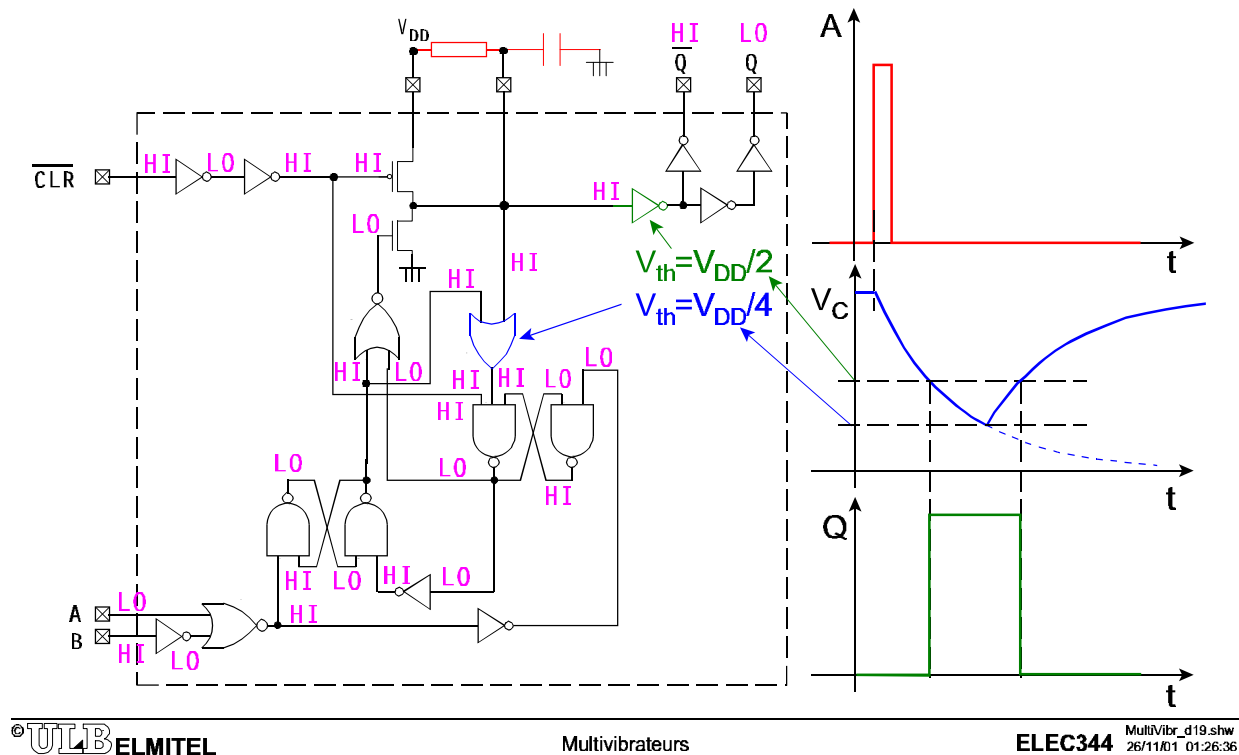


On trouve dans les grandes familles logiques des monostables tout faits sous forme de circuit intégré.

Certains d'entre-eux possèdent un circuit RC interne de faible valeur qui fixe une durée métastable T_m de l'ordre de quelques dizaines de ns. Pour des durées plus grandes, on connecte un RC extérieur. La plupart du temps, on trouve deux bornes pour le condensateur, dont une est commune avec la résistance, l'autre côté de la résistance étant connecté à l'alimentation ou à la masse. Des abaques donnent T_m en fonction de R et de C.

Le déclenchement peut se faire sur flanc montant ou descendant, via deux bornes séparées. Deux sorties complémentaires sont fournies, ainsi qu'une borne de CLR permettant de forcer l'état stable.

Exemple : le 4528



75

Le 4528 est un circuit intégré de la famille CMOS 4000 qui contient 2 monostables identiques indépendants. La figure illustre 1/2 circuit dans son état stable. Il contient plusieurs portes et bistables. Deux portes particulières servent de comparateur pour la tension sur le condensateur externe. L'un des comparateurs a un seuil normal de $V_{DD}/2$, tandis que l'autre est à $V_{DD}/4$.

Les deux entrées de déclenchement sont A et B.
CLR' force l'état de repos.

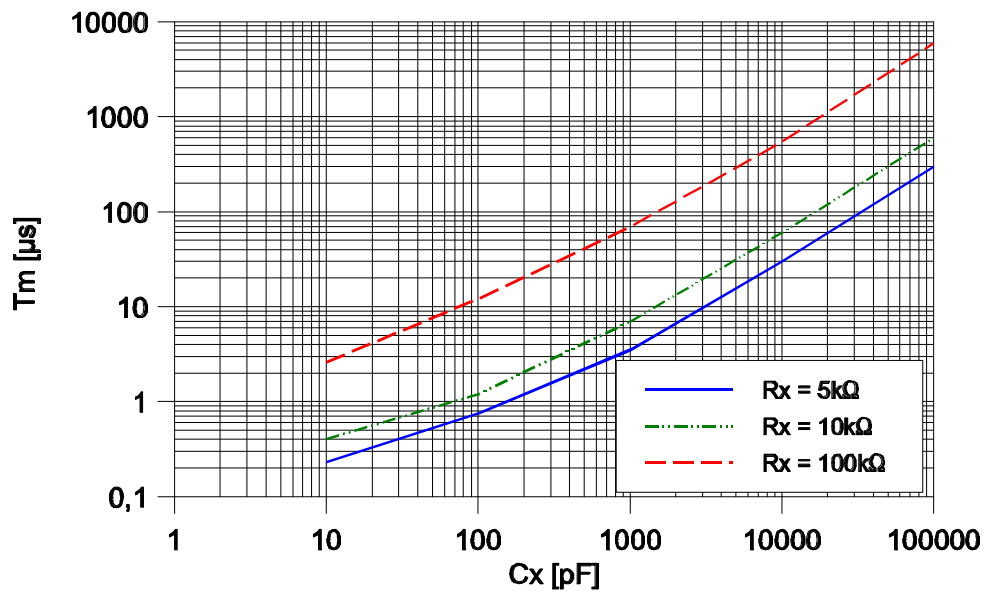
A titre d'exercice:

- vérifier que l'état présenté est l'état stable
- donner l'état actif des entrées A et B
- montrer l'action du CLR'
- montrer que le déclenchement engendre bien les formes de signaux présentées à droite

76

Exemple : 4528

abaque de choix de Rx et Cx



L'abaque fourni par le constructeur nous donne une idée des valeurs courantes des composants externes

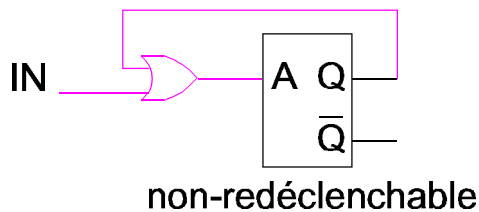
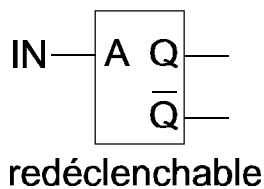
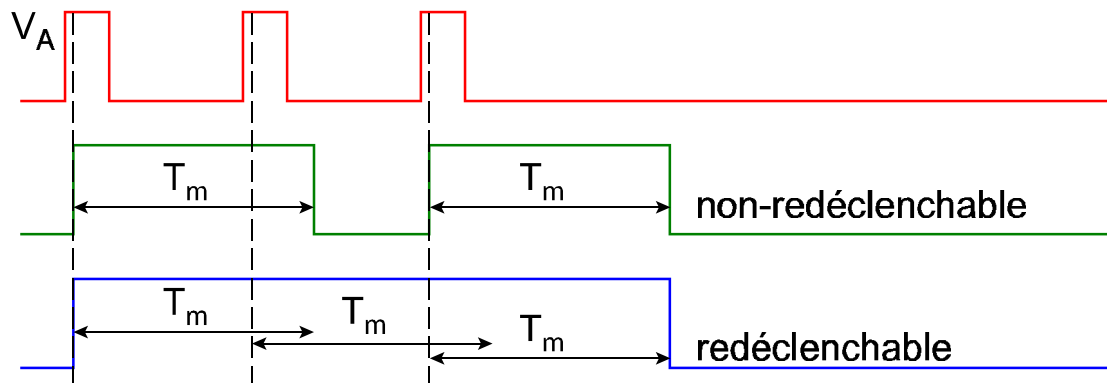
- R_x de quelques $k\Omega$ à $100k\Omega$
- C_x de $10pF$ à $100nF$

pour des durées d'état métastable entre 100 ns et 6 ms .

Nous verrons dans la suite du chapitre qu'il existe des circuits plus appropriés pour produire des délais plus longs.

Monostable redéclenchable

"retriggerable monostable"



Dans les monostables déclenchés par flanc, on distingue :

Les monostables non-redéclenchables ("non-retriggerable"), pour lesquels toute action sur l'entrée est inhibée pendant l'état métastable.

Les monostables redéclenchables pour lequel tout flanc actif d'entrée réinitialise l'état métastable à son début. On réalise ainsi aisément une détection de seuil de fréquence : si la période du signal d'entrée est inférieure à l'état métastable, la sortie du monostable se maintient à 1.

On peut transformer un monostable redéclenchable en non-redéclenchable en ajoutant une porte de blocage de l'entrée de déclenchement commandée par une rétroaction de la sortie.

A titre d'exercice, on montrera que le 4528 est redéclenchable.

Monostables : emploi

- ▶ valeurs de R
 - ◆ TTL : $5\text{k}\Omega$ à $50\text{k}\Omega$
 - ◆ CMOS : qq $\text{k}\Omega$ à qq $\text{M}\Omega$
 - grâce aux impédances d'entrée élevée
 - possibilité de longs délais
- ▶ précautions
 - ◆ sensibilité aux parasites
 - câblage court
 - éviter les couplages parasites entre sortie et bornes de R et C
 - ◆ découpler l'alimentation V_{DD} , car le seuil de décision en dépend

Pour réaliser des monostables, les CMOS sont nettement plus avantageux que les TTL :

- on peut utiliser des résistances de pull-up et de pull-down jusqu'à quelques dizaines de $\text{M}\Omega$, compte tenu des impédances d'entrée élevées. (En TTL une résistance de pull-down doit être inférieure à $1\text{k}\Omega$ et un pull-up à $50\text{k}\Omega$)
- on peut donc atteindre des durées d'état métastable T_m sensiblement plus élevées
- le seuil des CMOS est mieux défini, peu dépendant de la température, et proportionnel à la tension d'alimentation ce qui rend les T_m peu sensibles aux fluctuations de l'alimentation

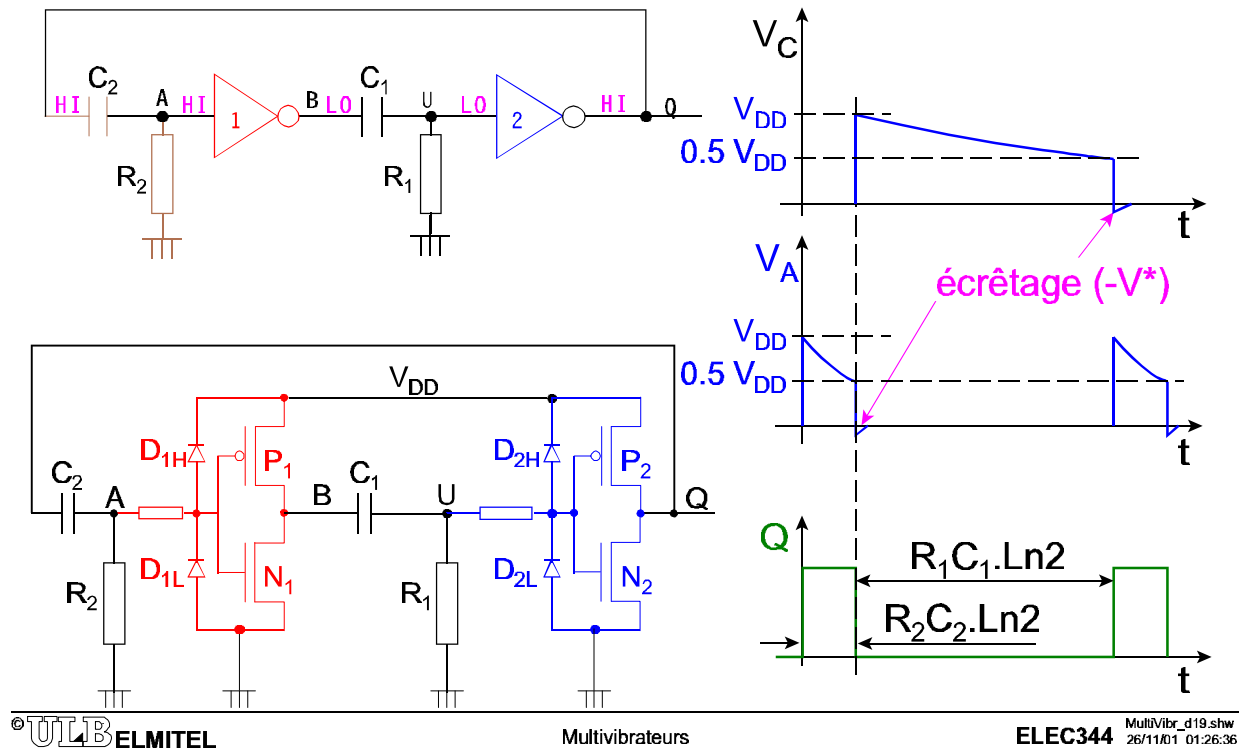
Comme tous les circuits pilotés par flanc ou impulsion courte, les monostables sont sensibles au bruit sur les entrées.

Vu différence de pente parfois très faible entre l'exponentielle du RC et le seuil du comparateur qui le mesure, un bruit sur ce signal, ou sur l'alimentation, peut influencer sensiblement la durée métastable. On évitera donc les couplage parasites et on soignera le découplage de l'alimentation.

Multivibrateurs : plan

- définitions
- bistables
- monostables
- **astables**
 - ◆ à deux constantes de temps
 - ◆ à une constante de temps
 - ◆ le 555

Astable à deux RC



85

L'astable présenté ici est composé de deux monostables simples obtenus, comme on l'a vu précédemment par mise en cascade de 2 inverseurs CMOS et deux circuits RC. Chacun des monostables est responsable d'un des 2 états métastables.

Pour expliquer le fonctionnement du circuit, il faut partir d'un état plausible qui soit métastable.

Supposons l'état initial représenté sur la figure avec $Q=HI$ et les deux condensateurs déchargés.

Le circuit R_1-C_1 est au repos, car C_1 est déchargé et il n'y a pas de courant dans R_1 .

Par contre, R_2-C_2 n'est pas en régime; C_2 se charge à travers le transistor P_2 et la résistance R_2 ; le potentiel V_A décroît exponentiellement.

Lorsque V_A devient $\leq V_{DD}/2$

- l'inverseur 1 change d'état et sa sortie présente un flanc montant
- ce flanc se répercute via C_1 (circuit dérivateur) sur l'entrée de l'inverseur 2
- la sortie Q passe de HI en LO
- ce flanc descendant est transmis par le condensateur C_2 à l'entrée A qui tend à passer de $(V_{DD}/2)$ à $(-V_{DD}/2)$
- la diode d'écrêtage D_{1L} va entrer en conduction et maintenir V_A à la tension de seuil $(-V^*)$;
- C_2 se décharge rapidement à travers une basse impédance formée par le transistor N_2 et la diode d'écrêtage D_{1L} ; le circuit R_2-C_2 est alors au repos.
- C_1 se charge au travers de R_1 et le potentiel du point U descend exponentiellement

Lorsque V_U devient $\leq V_{DD}/2$, le scénario précédent se reproduit en permutant les rôles des R-C et des inverseurs; Q revient à l'état initial et ainsi de suite

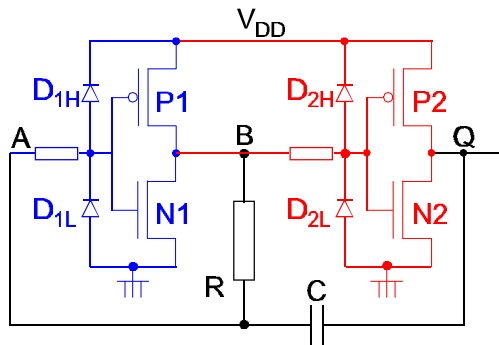
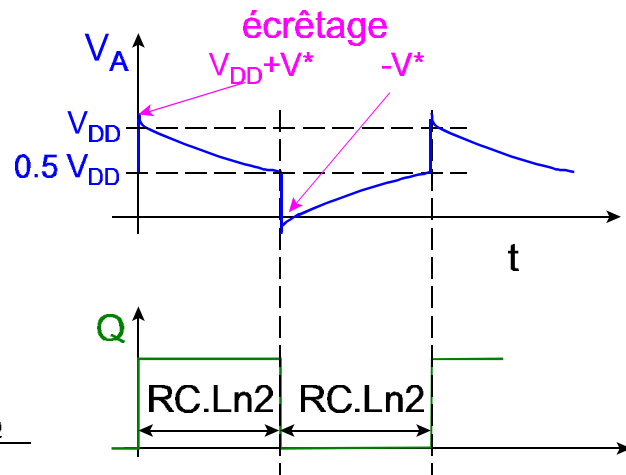
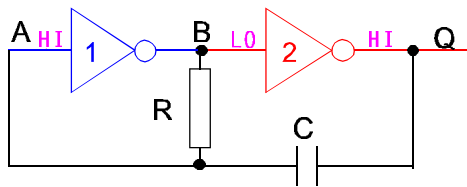
La sortie Q est donc une onde rectangulaire dont les durées des états HI et LO sont respectivement voisines de $(R_2C_2)Ln2$ et $(R_1C_1)Ln2$

Ce montage démarre spontanément, mais les tensions d'alimentation doivent avoir un temps de montée court, faute de quoi le montage se verrouille dans un état stable.

A titre d'exercice, on recherchera la présence de cet état stable.

86

Astable à un RC



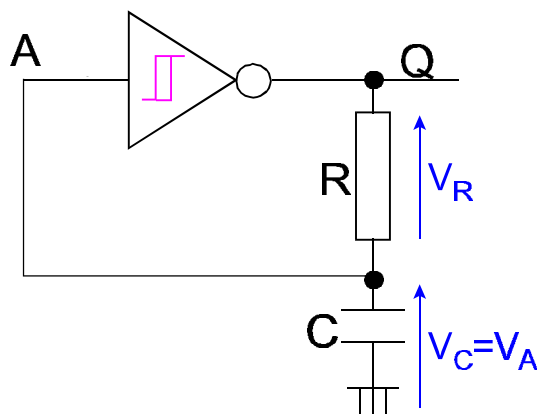
L'astable précédent permettait de régler séparément les durées des état HI et LO. Si l'on se contente d'un signal de sortie approximativement carré, on peut se limiter à une seule constante de temps RC.

Trouvons un état plausible qui soit métastable : soit $V_A = \text{HI}$, $V_B = \text{LO}$ et $Q = \text{HI}$. Le condensateur C est déchargé.

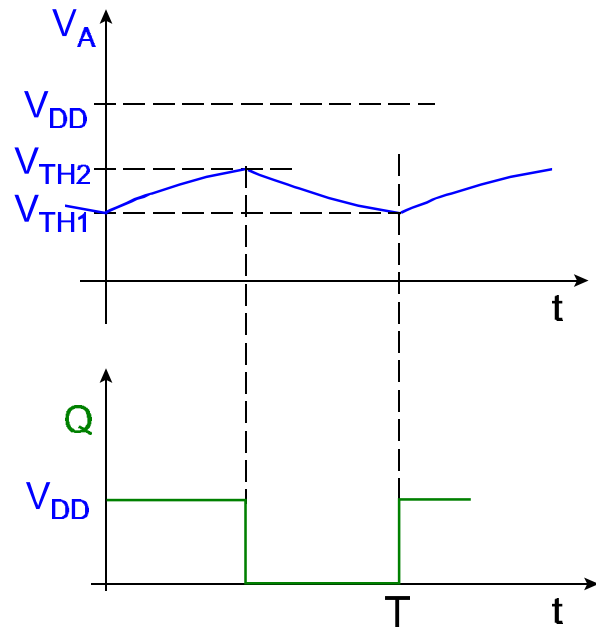
Le RC n'est pas en régime, puisque la résistance voit une tension 0 à sa borne supérieure et une tension V_{DD} à sa borne inférieure.

- C se charge à travers le circuit formé par le transistor "pull-up" P_2 , R et le transistor pull-down" N_1 . V_A décroît exponentiellement.
- lorsque V_A atteint le seuil $V_{DD}/2$, l'inverseur 1 change d'état, ce qui fait basculer l'inverseur 2; la sortie passe en LO
- ce flanc descendant est transmis par le condensateur C à l'entrée A qui tend à passer de $(V_{DD}/2)$ à $(-V_{DD}/2)$
- la diode d'écrtage inférieure D_{1L} va entrer en conduction et maintenir V_A à $(-V^*)$;
- C se décharge rapidement à travers une basse impédance formée par le transistor "pull-down" N_2 et la diode d'écrtage D_{1L} ;
- lorsque la tension en A redevient nulle, la diode d'écrtage D_{1L} se bloque; le condensateur est déchargé, mais le circuit RC n'est pas en régime, car la résistance R voit une tension V_{DD} à sa borne supérieure et 0 à sa borne inférieure
- C se charge à travers le circuit formé par le transistor "pull-up" P_1 , la résistance R et le transistor "pull-down" N_2 . Le potentiel du point A croît exponentiellement.
- lorsque V_A atteint le seuil $V_{DD}/2$, l'inverseur 1 change d'état, ce qui fait basculer l'inverseur 2 et la sortie passe en HI
- C se décharge rapidement à travers le "pull-up" P_2 et la diode d'écrtage D_{1H} en A
- on est revenu à l'état de départ

Astable à trigger de SCHMITT



$$T \approx 2.RC. \left(1 - \frac{V_{TH1}}{V_{TH2}}\right)$$



Il est conseillé d'utiliser des portes avec entrée "trigger de Schmitt" chaque fois que les temps de montée d'entrée sont trop lents par rapport aux spécifications.

On peut aussi se servir de l'hystérèse entre les deux seuils d'entrée pour fabriquer un astable. Le RC est ici monté en intégrateur.

Le principe du montage est simple:

- au départ du montage (non-représenté sur la figure), le condensateur est déchargé, l'entrée A est à LO, la sortie Q à HI.
- le RC n'est pas en régime puisque la résistance voit une tension positive égale à V_{DD} . C se charge et la tension du point A croît en $V_{DD}(1-\exp(-t/RC))$
- lorsque V_A atteint le seuil supérieur V_{TH2} , l'inverseur bascule et la sortie Q passe en LO
- la tension sur la résistance est alors négative et vaut $(-V_{TH2})$; le courant s'inverse et décharge le condensateur
- lorsque V_A descend au seuil inférieur V_{TH1} , l'inverseur bascule et la sortie Q passe en HI
- la tension sur la résistance est alors positive et vaut $(V_{DD}-V_{TH1})$; le courant s'inverse et charge le condensateur
-

Si les seuils V_{TH1} et V_{TH2} sont symétriques par rapport à $V_{DD}/2$, le rapport cyclique de la sortie Q est de 50%.

Exercice : démontrer que, dans ce cas, et en supposant que les exponentielles peuvent être assimilées à leur tangente, la période vaut approximativement :

$$T = 2.RC.(1-V_{TH1}/V_{TH2})$$

Le 555

- ▶ montage universel
 - ◆ monostable
 - ◆ astable
- ▶ montage ancien (1970)
 - ◆ toujours très populaire
 - ◆ a été reconçu en CMOS
- ▶ large gamme de périodes
 - ◆ de $10\mu\text{s}$ à 10s
 - ◆ gamme de C de 1nF à $100\mu\text{F}$
 - ◆ gamme de R de $1\text{k}\Omega$ à $10\text{M}\Omega$
- ▶ bon marché
- ▶ cf labo

Parmi les circuits multivibrateurs , le "Timer 555" est l'un des plus populaires depuis 1970.
Il permet à la fois de réaliser des bistables et des astables, dans de très larges plages de constantes de temps .
Il sera étudié en détail au laboratoire.