

第1章 绪 论

尽管目前绝大多数专用集成电路都完全或部分采用基于硬件描述语言(HDL)的专用综合流程来进行设计,并采用现场可编程门阵列或单元库来实现,但是了解超大规模集成(VLSI)电路的电路设计和物理设计基本知识变得越来越重要了,原因至少有以下几点。首先,在解决更为困难的问题时需要了解电路和版图设计问题有基本的了解。其次,出色的工程师通常可凭借其物理直觉来快速估算出电路的行为特点,不需要完全依赖计算机辅助设计(CAD)工具。

为达到上述目标,本书将着力于从底向上建立对集成电路的了解,更多关注逻辑电路、逻辑设计以及系统设计。归根结底,我们一直相信学习VLSI设计的最佳方法是动手设计。

为了让读者熟悉一个数字系统的全定制实现过程,在本章中将介绍集成电路的发展简史、VLSI设计在当前及不久的将来所面临的挑战、VLSI设计透视、VLSI设计与制造、互补金属氧化物半导体(CMOS)逻辑设计的规则与模式,以及数字系统的设计和实现方法。

1.1 VLSI 简介

本节首先简要回顾一下VLSI技术的历史。接下来介绍一种硅平面工艺,现代CMOS工艺正是在此基础上建立起来的,还要介绍特征尺寸的最终限制。然后介绍的是VLSI电路的分类,采用VLSI电路的好处,制造VLSI电路的合适工艺,并简要介绍了按比例缩小理论。最后,详细讨论了深亚微米(DSM)器件和连线的设计挑战。文中还简单探讨了VLSI技术的经济学问题和未来发展。

1.1.1 简介

在本小节中将介绍VLSI技术的历史,一种基本的硅平面工艺以及特征尺寸的最终限制。

1.1.1.1 简史

VLSI的历史可追溯至晶体管的发明。1947年,John Bardeen, Valter Brattain 和 William Shockley(都在贝尔实验室)发明了第一只点接触晶体管。由于这一突出贡献,他们在1956年获得了物理学诺贝尔奖。继这一发明之后,贝尔实验室致力于开发双极型晶体管(BJT)。这一努力建立了现代双极型晶体管的基础。双极型晶体管迅速替代了真空管,成为电子系统的主流,因为它们相比而言更为可靠、噪声更小且功耗更低。

晶体管发明十年之后,德州仪器(TI)的Jack Kilby制作了第一枚集成电路,探索了在单个硅片上制作多个晶体管微型化的可能性。这项工作奠定了晶体管-晶体管逻辑(TTL)系列的基石,而TTL系列在20世纪70年代到20世纪80年代非常流行。它们当中的一部分在今天依然存在且被广为使用,尽管这些电路现在的制造可能与原先的设计有所不同。由于发明了这款集成电路,Jack Kilby在2000年荣获了物理学诺贝尔奖。

尽管金属氧化物半导体场效应晶体管(MOSFET)或简称金属氧化物半导体(MOS)晶体管

的发明要早于双极型晶体管，但是它的使用却是直到 20 世纪 70 年代才被推广开来。Julius Lilien-field(德国)(美国专利 1745175)和 Oskar Heil(英国专利 439457)分别在 1925 年和 1935 年发布了他们的专利。1963 年，仙童公司的 Frank Wanlass 利用分立元件搭建了第一个互补金属氧化物半导体(CMOS)逻辑门，展示了 CMOS 技术利用两种不同 MOS 管——n 型和 p 型管表现出的超低等待功率特征。这一电路是当今 CMOS 领域的奠基石。

1.1.1.2 硅平面工艺

随着硅平面工艺的出现，MOS 集成电路由于其低成本而变得流行开来，每个晶体管所占的面积要比双极型晶体管小，且制造工艺要更为简单。在开始阶段采用的是 p 型 MOS(pMOS)管，但很快被 n 型 MOS(nMOS)管所取代，因为 nMOS 管的性能、可靠性和成品率更高。nMOS 逻辑电路的一个优点是它们所需的面积比 pMOS 电路要小。这一优点可借助英特尔公司采用 nMOS 工艺开发的 1101——256 位静态随机存储器(RAM)和 4004——第一款 4 位微处理器来说明。尽管如此，在 20 世纪 80 年代，对集成度的需求急速上升；nMOS 逻辑电路较高的等待功率极大地限制了集成度的提升。于是出现了 CMOS 工艺，并迅速取代了 nMOS 工艺，成为 VLSI 技术的主流，尽管采用 CMOS 工艺的逻辑电路所需的面积要远大于采用 nMOS 工艺实现的逻辑电路。现在 CMOS 工艺已经成为了 VLSI 设计最成熟和最流行的技术。

硅平面工艺的基本特征就是在硅表面的特定区域内淀积选定掺杂原子以改变或修改这些区域的电学特性的能力。为达到这一目标，如图 1.1 所示，需要进行如下基本步骤：(1)在硅表面形成氧化层(二氧化硅， SiO_2)；(2)去除选定区域的氧化层；(3)在硅表面和氧化层上淀积所需的掺杂原子；(4)选定区域上的掺杂原子扩散到硅当中去。重复采用这 4 个步骤，就可以制作出所需的集成电路(IC)。

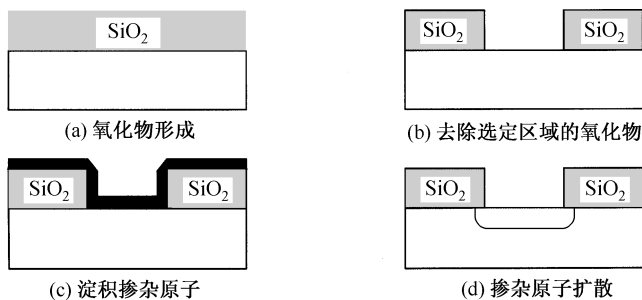


图 1.1 硅平面工艺的基本步骤

从 20 世纪 60 年代以来，基本硅平面工艺经过了一次又一次的完善。特征尺寸，即集成电路特定工艺中所允许的最小线宽或两条线之间的最小间距持续地快速变化，从 1969 年的 $8\ \mu\text{m}$ 变化到了现在(2010 年)的 $32\ \text{nm}$ 。随着这一变化，芯片上的晶体管数目急速增加，如图 1.2 所示，其中展示的是英特尔微处理器的发展。从图中可以看出，晶体管的数目几乎是每两年翻一番。芯片上晶体管数目的指数增长首先是被摩尔在 1965 年发现的，由他提出了著名的摩尔定律，摩尔定律表示芯片上晶体管的数目每 18 至 24 个月要翻一番。三年之后，Robert Noyce 和摩尔创立了英特尔公司，当今已经成为了全球最大的芯片制造商。

摩尔定律实际受两个主要力量推动。首先，随着集成电路工艺设备的改良，特征尺寸大幅

减小。其次,管芯^①面积逐步增加从而可以容纳更多的晶体管,因为晶圆的缺陷密度已经显著降低。现在,面积为 $2 \times 2 \text{ cm}^2$ 的管芯已并不少见。然而我们要看到的是管芯的成品率和每块管芯的成本与管芯面积息息相关。为了获得较好的管芯成品率,管芯的面积必须要控制在一定的范围内。有关管芯成品率和更普通的话题细节,即 VLSI 经济学方面的内容将在本节的稍后部分讨论。

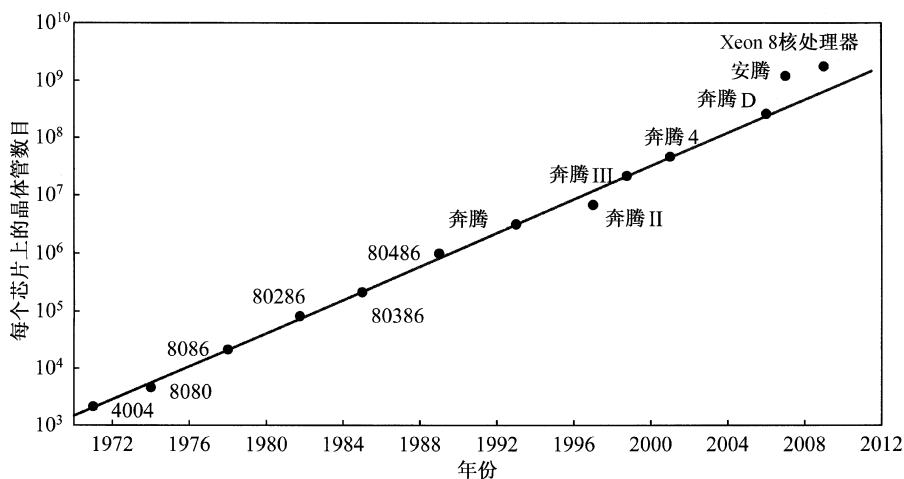


图 1.2 以英特尔微处理器来展示摩尔定律

1.1.1.3 特征尺寸的最终限制

特征尺寸可以减小到多少存在一个限制。至少有两个因素会限制特征尺寸无限制地减小下去。首先,随着器件尺寸的减小,数字的统计波动($\approx \sqrt{n}$)会成为限制电路性能的一个重要因素,不光是对模拟电路而言,最终对数字电路也是如此。这使得电路的设计相比以前变得更为困难。最终,每个器件中可能只包含几个电子,电路设计的整个概念会变得与今天所用的概念截然不同。

限制特征尺寸无限制减小的另一个因素是光刻设备的分辨率。尽管诸如光学邻近校正(OPC)、相移掩膜、双重曝光和浸润式光刻等这些分辨率提升技术已经使得目前分辨率可达到 45 nm 甚至更低,但是随着特征尺寸的降低,在不远的将来还需要更为昂贵的光刻设备以及其他制造设备。而设备的开发和制造成本很有可能会限制可用于制造集成电路的特征尺寸。

复习题

- Q1-1 描述硅平面工艺的基本步骤。
- Q1-2 描述摩尔定律及其含义。
- Q1-3 特征尺寸的最终限制是什么?

1.1.1.2 VLSI 电路的基本特征

在本节中将讨论 VLSI 电路的含义、采用 VLSI 电路的好处、制造 VLSI 电路的合适工艺以及按比例缩小理论。

^① 管芯通常指的是晶圆上的一个 IC, 而芯片表示的是已经从晶圆上切割下来 IC。有时候它们二者可以互用。

1.1.2.1 VLSI 电路的分类

粗略地讲,“VLSI”这个词是指所有包含很多元件,例如晶体管、电阻、电容等的集成电路。更为确切地说,集成电路(IC)可根据所包含元件的数目划分为以下几类。

- 小规模集成(SSI)指的是包含少于 10 个门的 IC。
- 中等规模集成(MSI)指的是包含多至 100 个门的 IC。
- 大规模集成(LSI)指的是包含多至 1000 个门的 IC。
- 超大规模集成(VLSI)指的是包含超过 1000 个门的 IC。

在此“门”指的是一个二输入基本门,例如一个与门或与非门。另外一种衡量一个 IC 复杂度的常用方法是晶体管的数目。采用如下经验系数,FPGA 器件、基于单元的设计和双极全定制设计的门数目和晶体管数目都可以相互转换。

- 现场可编程门阵列(FPGA):在 FPGA 器件中,由于一些不可避免的开销,1 个门近似等于 7~10 个晶体管,这取决于所选用 FPGA 器件的类型。
- 基于单元的 CMOS 设计:由于一个基本的二输入与非门有两个 pMOS 管和两个 nMOS 管组成,在基于单元的设计中一个基本门通常可算做 4 个晶体管。
- 双极全定制设计:在双极型晶体管逻辑中,例如 TTL,每个基本门大致都由 10 个基本元件组成,包括晶体管和电阻。

常见的一些 VLSI 电路有:微处理器,微控制器(嵌入式系统),存储器器件(静态随机存取存储器/SRAM,动态随机存取存储器/DRAM,闪存),不同的 FPGA 器件以及专用处理器,例如数字信号处理(DSP)器件和图形处理单元(GPU)器件。

近年来,有时会用超大规模(ULSI)这一术语来表示集成了亿万个元件的单个电路。不过在本书中不会采用这一术语。取而代之的是书中只采用 VLSI 这一术语来表示晶体管及/或其他元件高度集成的 IC。

1.1.2.2 采用 VLSI 电路的好处

完成集成电路的分类之后,现在来讨论采用 VLSI 电路的好处。由于几乎无须采用人工组装,集成可以降低制造成本,除此之外集成还可以给设计带来很大改善。这一点从集成带来的以下几点影响可以看出。首先,集成减少了寄生,包括电容,电阻甚至电感,因而刻蚀的最终电路可以高速工作。其次,集成降低了功耗,从而产生的热量更少。集成电路的大部分功耗是由 I/O 电路造成的,当信号转换时 I/O 电路中的大量电容需要被充电或者是放电。通过适当的集成,这些 I/O 电路中的绝大部分可以去除掉。第三,集成系统的物理尺寸更小一些,因为芯片占据的面积比初始系统的面积要小。因此,采用 VLSI 技术来设计一个电子系统可使得系统具有更高性能、消耗更少功率并占据更小的面积。这些因素反映到最终的产品中就是更小的物理尺寸、更低的功耗和更低的成本。

1.1.2.3 VLSI 技术

现在可以采用多种技术来设计和实现一个集成电路。其中最为常用的技术为 CMOS 工艺,双极型晶体管和砷化镓(GaAs)。在这三种技术中,CMOS 工艺是 VLSI 领域的主导技术,因为它有最低的功耗和最高的集成度,因而极具吸引力。双极型晶体管的特有特征是比较 CMOS 电路具有更高的工作频率。因此,在射频(RF)应用中常采用双极型晶体管。现在

CMOS 工艺中也可以嵌入双极型晶体管的制造,因而造就了一种混合工艺,即双极 CMOS (BiCMOS) 工艺。为提高双极型晶体管的性能,绝大多数现代 BiCMOS 工艺还提供一种新型的晶体管,称为锗硅 (SiGe) 晶体管,以用于特殊高频应用中,例如 RF 收发器,需要工作在 5 ~ 10 GHz 的频率之下。第三种技术是砷化镓,是专用于微波应用的一种技术,需要工作频率高至 100 GHz。

有两种标准可用来评估一种技术是否可以提供高度集成。这两种标准包括功耗和每个晶体管(或逻辑门)的面积。截至目前,只有 CMOS 工艺可同时满足这两个标准,因此在业界广为采用,在通信、计算机、消费类产品、多媒体等领域出现了大量不同的芯片。不过,功耗问题仍然是设计此类 CMOS 芯片面临的主要挑战,因为单个芯片的功耗可高达 100 ~ 120 瓦。除此之外,散热机制会加大整个系统设计的复杂程度。因此,现今可用的商用器件通常会通过低功耗设计和/或复杂的片内和/或片外电源管理模块来将其功耗限制在这一值以内。

1.1.2.4 按比例缩小理论

在特征尺寸的驱动下,Dennard 在 1974 年提出了恒定场按比例缩小理论。根据这一理论,器件的每一个方向的尺寸,即 x , y 和 z 以系数 k 按比例缩小,其中 $0 < k < 1$ 。为了维持器件中的恒定场,所有工作电压都要求按同一系数 k 缩小,而电荷密度需按系数 $1/k$ 放大。

图 1.3 对这一按比例缩小理论进行了说明。原始器件的沟道长度和宽度按系数 k 进行了缩小,即 $L' = k \cdot L$, $W' = k \cdot W$; 二氧化硅厚度同样按同样比例缩小,即 $t'_{ox} = k \cdot t_{ox}$ 。最大耗尽深度 x_d 同样遵循这一规则,即 $x'_d = k \cdot x_d$ 。

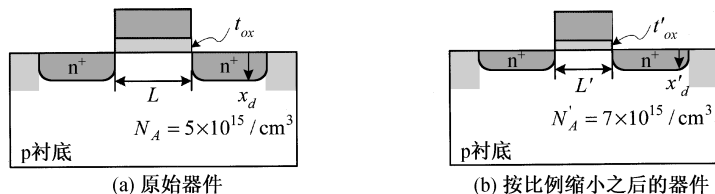


图 1.3 按比例缩小理论说明(图中尺寸为按比例绘制)

图 1.4 展示的是历史上从 0.5 μm 变化到今天的 32 nm 的历史缩小记录。从图中可以看到,每一代的按比例缩小系数 k 大约为 0.7。这意味着每一代芯片的面积减小的系数为 2,因此当芯片面积维持不变时,芯片上的晶体管数目会加倍。

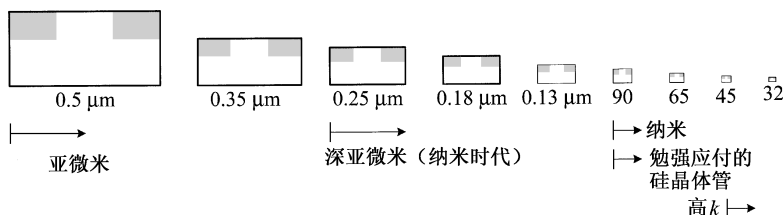


图 1.4 CMOS 工艺的特征尺寸发展趋势

按比例缩小器件的优点有以下几点。首先,器件密度将按系数 $1/k^2$ 增加。其次,电路延迟将按系数 k 减小。第三,每个器件的功耗将按系数 k^2 减小。因此,按比例缩小的器件相比原始器件具有更小的面积,更高的性能,并具有更低的功耗。

复习题

Q1-4 什么是 VLSI?

Q1-5 解释 CMOS 技术成为当今 VLSI 领域主流的原因。

Q1-6 采用 VLSI 电路有何优点?

1.1.1.3 VLSI 电路设计中存在的问题

当特征尺寸小于 $1\ \mu\text{m}$ 时, VLSI 制造工艺被称为亚微米(SM)工艺, 而当特征尺寸大致小于 $0.25\ \mu\text{m}$ 时, VLSI 制造工艺被称为深亚微米(DSM)工艺^①。由这两种工艺制作的器件分别被称为 SM 器件和 DSM 器件。目前, 在大规模系统设计中 DSM 器件较为常见, 因为它们可以提供一种更为经济的方式, 可在单个芯片上集成一个更为复杂的系统。所得到的芯片通常称为片上系统(SoC)器件。

尽管 DSM 工艺使得我们可以设计非常复杂的大规模系统, 但这其中仍然存在许多设计挑战, 尤其是在特征尺寸低于 $0.13\ \mu\text{m}$ 的时候。相关的设计问题可分为两个主要的大类: DSM 器件和 DSM 互连线^②。在下面的章节中将对这类问题做简单讨论。

1.1.1.3.1 DSM 器件设计中存在的问题

DSM 器件的设计问题包括薄氧化层(栅氧化层)隧通/击穿、栅极泄漏漏电流、亚阈值电流、速度饱和、短沟道效应对 V_T 的影响、热载流子效应以及由漏极引起的势垒降低(DIBL)效应。

典型 DSM 工艺制作的器件特征总结如表 1.1 所示。从表中可以看到, 薄氧化层(栅氧化层, 即二氧化硅, SiO_2)厚度从 $0.25\ \mu\text{m}$ 工艺的 $5.7\ \text{nm}$ 降低至了 $32\ \text{nm}$ 工艺的 $1.65\ \text{nm}$ 。这一降低带来的侧面影响是薄氧化层的隧通和击穿。薄氧化层的隧通可导致非常大的栅极泄漏电流。为避免薄氧化层击穿, 必须要降低栅极的工作电压。这意味着噪声容限要相应降低, 亚阈值电流不能再被忽略。为减小栅极泄漏电流, 自 $45\ \text{nm}$ 工艺开始广泛采用高 k MOS 晶体管。在高 k MOS 晶体管中, 栅氧化层被一种高 k 电介质所取代。如此一来, 栅电介质厚度可能会显著增加, 因而可大大降低栅极泄漏电流。实际的栅电介质厚度取决于栅电介质材料的相对介电常数, 更多细节请参考 3.4.1.2 节。

表 1.1 典型 DSM 工艺中的器件特征

工 艺	0.25 μm	0.18 μm	0.13 μm	90 nm	65 nm	45 nm	32 nm
	1998	1999	2000	2002	2006	2008	2010
$t_{\text{ox}}(\text{nm})$	5.7	4.1	3.1	2.5	1.85	1.75	1.65
$V_{\text{DD}}(\text{V})$ ^③	2.5	1.8	1.2	1.0	0.80	0.80	0.80
$V_T(\text{V})$	0.55	0.4	0.35	0.35	0.32	0.32	0.32

除此之外, 随着器件沟道长度的减小, 已经不能像在长沟道器件中那样忽略速度饱和、短沟

① 有时候特征尺寸小于 $100\ \text{nm}$ 的工艺被称为纳米(nm)工艺。

② 将晶体管、电路、单元、模块和系统连接在一起的线称为互连线。

③ 本书原文版中下标斜体字符表示混乱, 因而在翻译版中仅对不统一的地方进行了规范, 其他表示方法与原文版保持一致——编者注。

道效应对 V_T 的影响以及热载流子效应。当电场处于某一临界值以下时, 沟道或者硅体中的电子和空穴的速度与所施加的电场成正比。然而, 在 400 K 时二者的速度会饱和在 8×10^6 cm/s, 而与掺杂浓度无关, 对应的电子和空穴的电场强度分别为 6×10^4 V/cm 和 2.4×10^5 V/cm。当速度达到饱和时, MOS 晶体管的漏电流与所施加的栅源电压将呈线性关系而非二次方关系。

当沟道长度与漏极的耗尽层厚度相当时, 此类器件被称为短沟道器件。短沟道效应 (SCE) 是指短沟道器件的 V_T 降低, 这可能是由以下几种原因综合造成的: 电荷共享, DIBL 和次表面穿通。电荷共享是指形成沟道的电荷并不完全是由栅极电压提供的, 还有可能是由其他原因提供的。DIBL 是指漏极电压对阈值电压的影响。次表面穿通指的是逻辑电压对源端 pn 结电子势垒的影响。

热电子是指动能高于其热能的电子。由于沟道末端的电场足够高, 在漏极结的空间电荷区通过热电子的碰撞电离可能会产生电子空穴对。产生的这些电子和空穴载流子可能会带来多个影响。首先, 它们可能会陷入栅氧化层并逐步改变器件的阈值电压 V_T 。其次, 它们可能会注入栅极当中形成栅电流。第三, 它们可能会进入衬底当中形成衬底电流。第四, 它们可能会使得寄生双极型晶体管正向偏置, 从而使器件失效。

1.1.3.2 DSM 互连线设计中存在的问题

DSM 互连线的设计问题来自 RLS 寄生器件, 包括电阻压降、RC 延迟、电容耦合、电感耦合、 Ldi/dt 噪声、电迁徙和天线效应。下面将对其一一进行简要介绍。

VLSI 芯片中连线的功能是充当传递信号、电源和时钟的导体。由于连线自身的固有电阻和电容的存在, 每根连线都有一定的电阻压降。典型 DSM 工艺的金属 1 连线的特征如表 1.2 中所给出。从表中可以看到, 连线的厚度和宽度随着特征尺寸的发展而减小。因此, 连线的方块电阻, 即每个方块的电阻值大大增加。

表 1.2 典型 DSM 工艺的金属 1 连线特征

工 艺	0.25 μm	0.18 μm	0.13 μm	90 nm	65 nm	45 nm	32 nm
厚度	0.61 μm	0.48 μm	0.40 μm	150 nm	170 nm	144 nm	95 nm
宽度/间距	0.3 μm	0.23 μm	0.17 μm	110 nm	105 nm	80 nm	56 nm
电阻值/方块 ($\text{m}\Omega/\square$)	44	56	68	112	100	118	178

连线电阻导致电阻压降, 这样会影响逻辑电路的性能, 甚至会引起逻辑电路的误操作。在 DSM 工艺中, 电阻压降更应引起重视, 原因有如下几点。首先, 连线电阻随着特征尺寸的减小而增加。其次, 流经连线的电流随着电源电压的降低而增加。第三, 由于所执行的功能更加复杂, 导致现代 DSM 芯片中的功耗增加。这意味着电流也会相应增加。因此, 在设计一个现代 DSM 芯片时, 一个重要的系统设计问题就是降低芯片的电阻压降。

连线的电容和电阻共同导致 RC 延迟, 这可能会影响逻辑电路的性能。随着特征尺寸的减小, 连线的 RC 延迟大幅度增加, 但门延迟却显著减小, 每一代变化的系数约为 $k \approx 0.7$ 。因此, 通过一根连线的信号延迟很容易地会大于经过驱动这一连线的栅极的延迟。如图 1.5 所示, 对于一根 1 mm 连线, 互连线延迟要小于门延迟, 当特征尺寸小于 130 nm 时互连线延迟可以忽略。然后, 当特征尺寸小于 90 nm 时, 互连线延迟要比门延迟大得多, 甚至成为整个系统延迟的决定力量。

此外, 在 DSM 工艺中, 两条相邻互连线的间距很窄, 甚至远小于连线的厚度, 如图 1.6(a) 所示。由此带来的直接效应是这两条连线之间的电容不可再被忽略, 可能会导致电容耦合现

象的发生,这表明其中一条连线中的信号耦合进了另外一条连线当中,而耦合进来的信号可能会破坏原有的信号。电容耦合可能会导致信号完整性问题。

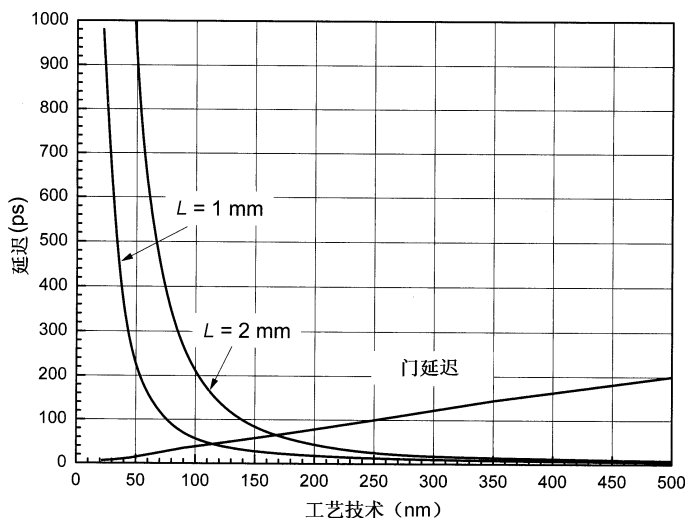
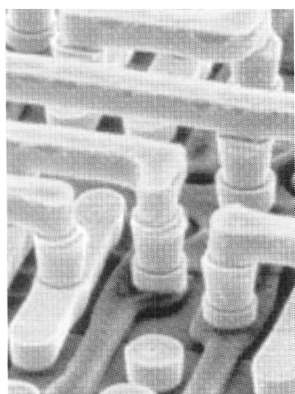


图 1.5 不同工艺中门延迟与连线延迟的关系

除了电容耦合,DSM 工艺中的连线电感可能还会引起耦合效应。法拉第感应定律表明,当一个导体处于邻近一个载流导体所产生的时变磁场当中时,该导体中会产生一个感应电压。这就是所谓的电感耦合。像电容耦合一样,电感耦合也可能会给邻近的连线带来噪声干扰。然而不同的是,电容耦合只限制在两个导体之间,而电感耦合可能会以一种不可预知的方式影响电路中的很大一部分面积。这是因为感性效应是在一个闭环中形成的,而这当中对于单个的信号路径可能会存在很多可能的返回路径,这些路径具有不同的路径长度。



(a) 0.35 μm 工艺中具有钨通孔的四层铝连线 (经国际商务公司授权引用。未经允许禁止使用)

工艺(nm)	金属层
500	3 (Al)
350	4 (Al)
250	5 (Al)
180	6 (Al, 低 k 值)
130	7/8/9 (Cu, 低 k 值)
90	7/8/9 (Cu, 低 k 值)
65	8/9 (Cu, 低 k 值)
45	8/9 (Cu, 低 k 值)
32	8/9 (Cu, 低 k 值)

(b) 典型工艺中按比例缩小的金属层

图 1.6 现代 CMOS 工艺中连线的变迁

由带有自感 L 的连线的时变电流引起的感应电压可表示为 Ldi/dt 。 Ldi/dt 效应通常由信号的转换尤其是时钟转换引起的。尽管连线的自感非常小,但如果连线中的电流变化非常快, Ldi/dt 效应也会变得很显著。 Ldi/dt 效应会给位于连线下方的电路带来电源尖峰脉冲,当这些尖峰脉冲足够大时会导致 Ldi/dt 噪声,并导致电源完整性问题。

电迁徙现象是指导体中的电子进行迁徙并移出了导体的晶格。这可能会使得导线断裂。为避免这一现象的发生,通过一条导线的电流密度必须要控制在一个临界值之内,而这一临界值由导线决定。铝中的电迁徙要比铜中的电迁徙严重得多。因此,在采用铝为导线的实际工艺中,铝通常要与少量的铜制成合金以削弱这一效应。

有一种在整个制造工艺完成之前就可能令器件失效的现象被称为天线效应。天线效应是指在制造过程中金属线上存储电荷的现象。当存储在金属 1 上的电荷数量足够多,形成一个较强的电场时,可能会对栅氧化层造成直接损坏。因此,必须要限制直接与多晶硅相接的金属 1 的面积。

1.1.3.3 VLSI 系统设计中存在的问题

随着 CMOS 工艺的发展,特征尺寸减小的速度要比我们所想象的快得多。随着这一发展,VLSI 系统设计中的许多挑战就随机出现了。这些挑战主要体现在以下几方面:电源分布网络、电源管理、时钟分配网络以及设计和测试方法。

一个完整的电源分布网络需要考虑电阻压降、热点、 Ldi/dt 噪声、地电位上跳和电迁徙。除了电阻压降和 Ldi/dt 噪声之外,热点问题也是设计 VLSI 系统时必须要考虑的问题。热点是指芯片上的某些局部区域的温度远高于芯片的平均温度。这些热点可能会损害芯片的性能,甚至最终使得整个芯片失效。

为了将功耗控制在一个可接受的限制范围之内,必须要在所有模块的设计中对功耗进行细致分析。由于在一个数字系统中并不是所有模块在所有时间段都需要处于激活状态,因此可能的话可以在某一时刻只为需要工作的模块供电。这样一来,一个 VLSI 芯片的功耗就可控制在一个额定界限之内。实际上,电源管理已经成为现代 VLSI 芯片设计中的一个重要问题。它甚至可以决定最终产品是否可以成功上市。

VLSI 系统设计的另外一个挑战是时钟分配网络。因为在这些系统中,为了执行复杂的逻辑功能,需要较大的面积和较高的工作频率来实现,时钟偏移必须要控制在非常窄的范围内。否则系统可能无法正常工作。因此在设计时钟分配网络时必须多加小心。此外,由于时钟分配网络需要驱动较大的电容,因而时钟分配网络的功耗不能再被忽视,必须要谨慎处理。

随着 VLSI 芯片高度集成的出现,系统复杂度的增加使得相关的设计变得更为困难。一个实用而有效的设计方法是采用分而治之的模式来限制可同时处理的元件的数目。不过,随着所需设计系统复杂度的提升,将多个不同模块组合成所需系统的难度也變得越来越大。而且,对组合系统进行测试也显得更为重要且更具有挑战性。

复习题

-
- Q1-7 解释图 1.5 的含义。
 - Q1-8 什么是热点问题?
 - Q1-9 DSM 器件有什么设计问题?
 - Q1-10 DSM 互连线有什么设计问题?
-

1.1.4 VLSI 经济学

粗略地讲,一个集成电路的成本主要由两个因素组成:固定成本和可变成本。固定成本,又称一次性工程(NRE)成本,是与销售量无关的。它主要是指从项目启动到获得第一次成功

样品所花费的启动成本。更为准确地说,固定成本包括直接成本和间接成本。直接成本包括研发(R&D)成本、掩膜制造成本以及市场和销售成本;间接成本包括制造设备的投资、CAD工具的投资、基础设施建设成本,等等。可变成本与产品数量成正比,主要是指晶圆制造成本,即晶圆的价格,对于一片 300 mm 的晶圆而言价格大概为 1200 ~ 1600 美元之间。

根据上述讨论,每个集成电路的成本可用下式来表示:

$$\text{每个 IC 的成本} = \text{IC 的可变成本} + \frac{\text{固定成本}}{\text{产品数量}} \quad (1.1)$$

每个集成电路的可变成本可用下面的等式来表示:

$$\text{IC 的可变成本} = \frac{\text{管芯的成本} + \text{管芯测试的成本} + \text{封装和最终测试的成本}}{\text{最终测试产量} \times \text{每片晶圆上的管芯数}} \quad (1.2)$$

一个管芯的成本等于晶圆价格除以好管芯的数目,可用如下公式来表示:

$$\text{管芯的成本} = \frac{\text{晶圆价格}}{\text{每片晶圆上的管芯数} \times \text{管芯产量}} \quad (1.3)$$

除去边缘处碎裂的管芯,每片晶圆上的管芯数目可近似用如下等式来表示:

$$\text{每片晶圆上的管芯数} = \frac{3}{4} \frac{d^2}{A} - \frac{1}{2\sqrt{A}} d \quad (1.4)$$

其中 d 为晶圆的直径, A 为方块区域管芯的面积。本式的推导留给读者当做练习。

管芯产量可通过以下广为采用的函数来估算:

$$\text{管芯产量} = \left(1 + \frac{D_0 A}{\alpha}\right)^{-\alpha} \quad (1.5)$$

其中 D_0 为缺陷密度,即每单位面积上的缺陷,单位为缺陷数/ cm^2 , α 为评估制造复杂度的参数。 D_0 和 α 的典型值分别为 0.3 ~ 1.3 和 4.0。从这一等式可以清晰地看到,管芯产量与管芯面积成反比。

下面的两个例子举例说明了上述有关集成电路成本的概念。在这两个例子中有意忽略了固定成本,在计算管芯成本时只考虑了晶圆价格。

例 1-1 管芯面积为 1 cm^2 。假设晶圆的直径为 30 cm,管芯面积为 1 cm^2 。缺陷 D_0 密度为 0.6 缺陷数/ cm^2 ,制造复杂度 α 为 4。假设每片晶圆的价格为 1500 美元,在不考虑固定成本的情况下计算每个管芯的成本。

解: 利用式(1.4)可估算出每个晶圆上的管芯数目。

$$\begin{aligned} \text{每片晶圆上的管芯数} &= \frac{3}{4} \frac{d^2}{A} - \frac{1}{2\sqrt{A}} d \\ &= \frac{3}{4} \left(\frac{30^2}{1}\right) - \frac{1}{2\sqrt{1}} 30 = 660 \end{aligned}$$

管芯产量可通过式(1.5)计算,如下所示:

$$\begin{aligned} \text{管芯产量} &= \left(1 + \frac{D_0 A}{\alpha}\right)^{-\alpha} \\ &= \left(1 + \frac{0.6 \times 1}{4}\right)^{-4} = 0.57 \end{aligned}$$

利用式(1.3)可求得每个管芯的成本为:

$$\text{管芯的成本} = \frac{\text{晶圆价格}}{\text{每片晶圆上的管芯数} \times \text{管芯产量}} = \frac{1500}{660 \times 0.57} = 3.98 (\text{美元})$$

例 1-2 管芯面积为 $2.5\text{ mm} \times 2.5\text{ mm}$ 。假设晶圆的直径为 30 cm ，管芯面积为 $2.5\text{ mm} \times 2.5\text{ mm}$ 。缺陷 D_0 密度为 $0.6\text{ 缺陷数}/\text{cm}^2$ ，制造复杂度 α 为 4。假设每片晶圆的价格为 1500 美元，在不考虑固定成本的情况下计算每个管芯的成本。

解：利用式(1.4)可估算出每个晶圆上的管芯数目。

$$\begin{aligned}\text{每片晶圆上的管芯数} &= \frac{3}{4} \frac{d^2}{A} - \frac{1}{2\sqrt{A}} d \\ &= \frac{3}{4} \left(\frac{30^2}{0.0625} \right) - \frac{1}{2\sqrt{0.0625}} 30 = 10\,740\end{aligned}$$

管芯产量可通过式(1.5)计算，如下所示：

$$\begin{aligned}\text{管芯产量} &= \left(1 + \frac{D_0 A}{\alpha} \right)^{-\alpha} \\ &= \left(1 + \frac{0.6 \times 0.0625}{4} \right)^{-4} = 0.96\end{aligned}$$

利用式(1.3)可求得每个管芯的成本为：

$$\text{管芯的成本} = \frac{\text{晶圆价格}}{\text{每片晶圆上的管芯数} \times \text{管芯产量}} = \frac{1500}{10\,740 \times 0.96} = 0.15 (\text{美元})$$

通过上面两个例子可以很明显地看出，管芯产量是管芯面积的强函数，即与管芯面积成反比。因此在实际应用中，有必要限制管芯的面积以控制管芯产量，从而将管芯成本控制在可接受的范围之内。需指出的是，在实际的芯片设计项目中，管芯面积、功耗和性能通常是考虑和权衡的三个主要因素。

复习题

Q1-11 决定一个集成电路成本的两个主要因素什么？

Q1-12 描述一次性工程(NRE)成本的含义。

Q1-13 一个集成电路的可变成本的要素是什么？

1.2 开关 MOS 晶体管

在本节中将开始研究 nMOS 晶体管和 pMOS 晶体管的基本工作原理，并将这两类晶体管当做开关来看待，分别称为 nMOS 开关和 pMOS 开关。然后考虑将 nMOS 晶体管和 pMOS 晶体管用做组合开关的组合效应，这一组合开关称为传输门(TG)或 CMOS 开关。同时还将讨论这三种开关在逻辑电路中的使用。最后介绍的是开关逻辑的基本设计规则。这种逻辑又称为可控逻辑，因为在某些特殊信号的控制下，数据输入可直接指向输出。

1.2.1 nMOS 晶体管

一个 nMOS 晶体管的物理结构基本由一个金属-氧化物-硅(MOS)系统和位于 p 型硅衬底表面的两个 n^+ 区构成，如图 1.7(a) 所示。MOS 系统是一个夹层结构，在两层金属或者一层多晶硅和 p 型衬底之间插入一层电介质(绝缘体)。金属或多晶硅被称为栅。衬底表面的两个 n^+ 区分别被称为漏和源。

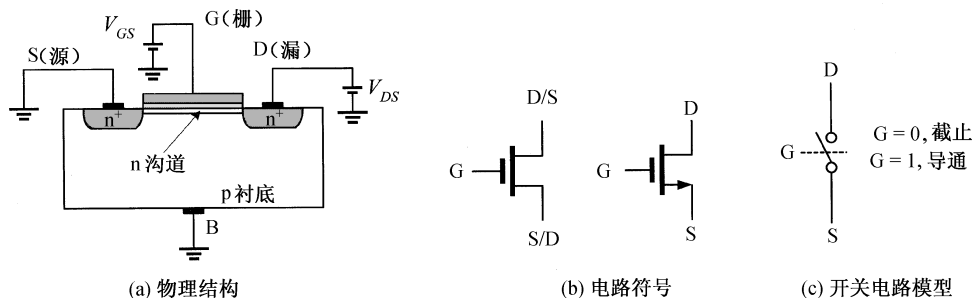


图 1.7 nMOS 晶体管

nMOS 晶体管的工作原理可通过图 1.7(a) 来说明。当在栅(电极)上施加足够大的正向电压 V_{GS} 时, 栅电压在硅表面建立的正向电场会将电子从 p 型衬底吸附到硅表面。这些电子在漏和源之间形成一个沟道。形成沟道的最小电压 V_{GS} 被定义为阈值电压, nMOS 的阈值电压用 V_{Tn} 表示。对于目前的亚微米和深亚微米工艺而言, V_{Tn} 的值为 $0.3 \sim 0.7 \text{ V}$, 具体取决于特定工艺情况。

对于数字应用, nMOS 晶体管可被当成一个简单的开关元件。当栅极电压大于或等于其阈值电压时, 开关导通, 否则开关截止。由于 nMOS 晶体管结构对称, 任意一个 n^+ 区都可被用做源或者漏, 具体取决于工作电压的施加方式。正向电压更高的是漏极, 另一端是源极, 因为 nMOS 晶体管的载流子是电子。

图 1.7(b) 给出的是在电路设计中经常使用的电路符号。在源极带有标示性箭头的符号通常在模拟应用中使用, 因为在模拟电路中的源极和漏极角色是固定的。另外一个不带标示性箭头的符号通常在数字应用中使用, 因为漏极和源极的角色会根据电路的实际工作条件动态变化。开关电路模型如图 1.7(c) 所示。

1.2.2 pMOS 晶体管

类似地, pMOS 晶体管的物理结构基本由一个金属-氧化物-硅(MOS)系统和位于 n 型硅衬底表面的两个 p^+ 区构成, 如图 1.8(a) 所示。MOS 系统是一个夹层结构, 在两层金属或者一层多晶硅和 n 型衬底之间插入一层电介质(绝缘体)。金属或多晶硅被称为栅。衬底表面的两个 p^+ 区分别被称为漏和源。

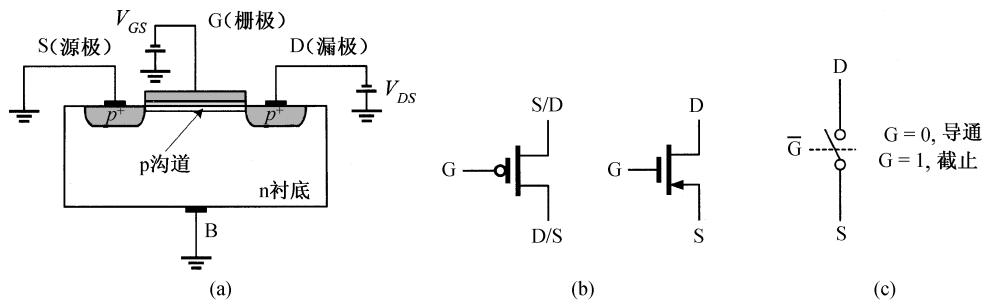


图 1.8 (a) 物理结构; (b) 电路符号; (c) pMOS 晶体管的开关电路模型

pMOS 晶体管的工作原理可通过图 1.8(a) 来说明。当在栅(电极)上施加足够大的负电压 V_{GS} 时, 栅电压在硅表面建立的反向电场会将空穴从 n 型衬底吸附到硅表面。这些空穴在漏和

源之间形成一个沟道。形成沟道的最小电压 $|V_{GS}|$ 被定义为阈值电压, pMOS 的阈值电压用 V_{Tp} 表示。对于目前的亚微米和深亚微米工艺而言, V_{Tp} 的值为 $-0.3 \sim -0.7$ V, 具体取决于特定工艺情况。

与 nMOS 晶体管一样, 对于数字应用, pMOS 晶体管可被当成一个简单的开关元件。当栅极电压小于或等于其阈值电压时, 开关导通, 否则开关截止。由于 pMOS 晶体管结构对称, 任意一个 p^+ 区都可被用做源或者漏, 具体取决于工作电压的施加方式。正向电压更高的是源极, 另一端是漏极, 因为 pMOS 晶体管的载流子是空穴。

图 1.8(b) 给出的是在电路设计中经常使用的电路符号。pMOS 晶体管的符号规则与 nMOS 晶体管的符号规则完全一样。在源极带有标示性箭头的符号通常在模拟应用中使用, 因为在模拟电路中的源极和漏极角色是固定的。另外一个不带标示性箭头但是在栅极上有一个小圆圈的符号通常在数字应用中使用, 因为漏极和源极的角色会根据电路的实际工作条件动态变化。小圆圈用于与 nMOS 晶体管区分, 表明 pMOS 晶体管是低电平有效的。开关电路模型如图 1.8(c) 所示。

1.2.3 CMOS 传输门

由于 nMOS 晶体管导通要求栅极和源极之间的电压幅值达到 V_{Tn} , 假设栅极和漏极同时接 V_{DD} , 因此 nMOS 开关的最大输出电压为 $V_{DD} - V_{Tn}$ 。同样地, 假设 pMOS 晶体管的栅极和漏极同时接 0 V, 那么 pMOS 开关的最小输出电压为 $|V_{Tp}|$ 。上面两句话还可以通过将 0 V 表示成逻辑 0, 将 V_{DD} 表示成逻辑 1 来描述成信息传输的形式, 如下所述。nMOS 晶体管可传输完美的 0, 但无法传输无损的 1; pMOS 晶体管可以传输完美的 1, 但无法传输无损的 0。

之前提及的 nMOS 晶体管和 pMOS 晶体管的缺点可通过将一个 nMOS 晶体管与一个 pMOS 晶体管并联以形成组合开关的形式来克服, 即所谓的传输门(TG)或 CMOS 开关, 如图 1.9 所示。由于 nMOS 晶体管和 pMOS 晶体管都是并联连接的, 因此其中一个晶体管的缺陷可由另一个晶体管来弥补。图 1.9(a) 给出了一个 TG 开关的路结构, 图 1.9(b) 给出了逻辑图标中常用的逻辑符号。

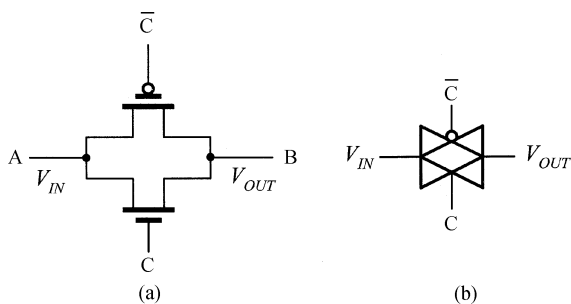


图 1.9 (a) 电路结构; (b) TG 开关的逻辑符号

尽管采用 TG 可以克服信号传输的损失, 但是每个 TG 开关需要两个晶体管, 即一个 nMOS 晶体管和一个 pMOS 晶体管。这意味着采用 TG 开关比单独采用 nMOS 开关或 pMOS 开关需要占据更多的芯片面积。在实际当中, 当芯片面积受到限制时, nMOS 晶体管要比 pMOS 晶体管更受欢迎, 因为电子迁移率要远大于空穴迁移率。因此, nMOS 晶体管的性能要比 pMOS 晶体管的性能更好。

复习题

- Q1-14 描述 nMOS 晶体管的工作原理。
- Q1-15 描述 pMOS 晶体管的工作原理。
- Q1-16 描述 CMOS 开关的工作原理。

- Q1-17 采用 nMOS 晶体管当开关时有什么缺点?
- Q1-18 采用 pMOS 晶体管当开关时有什么缺点?

1.2.4 简单开关逻辑设计

如前所述, 三种开关, 即 nMOS 开关, pMOS 开关和 TG 开关中的任意一种都可用做开关来控制两点之间的闭合(导通)或者断开(截止)状态。将这些开关进行合理组合, 就可构建出开关逻辑电路。在此首先将讨论组合开关, 然后介绍从一个给定的开关函数开始构建一个开关逻辑电路的系统设计方法。

1.2.4.1 组合开关

在很多应用中, 通常会将两个或多个开关以串联、并联或者组合的方式组合起来形成一个组合开关。例如, 图 1.10 中给出了两种将两个开关串联起来形成一个组合开关的情况。所得开关的工作由两个控制信号控制: S_1 和 S_2 。当两个控制信号 S_1 和 S_2 都有效时组合开关才导通, 否则都处于截止状态。

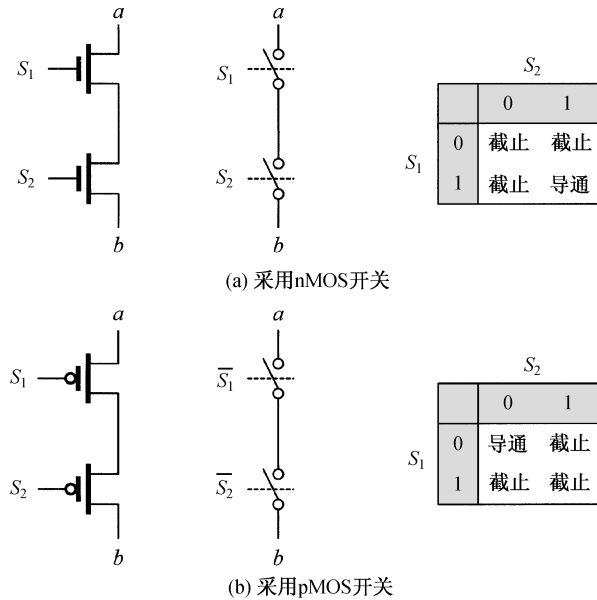


图 1.10 串联开关工作原理

之前讲过, 为了激活一个 nMOS 开关需要在 nMOS 晶体管的栅极施加一个高电平电压, 激活一个 pMOS 开关需要在 pMOS 晶体管的栅极施加一个低电平电压。因此, 对于图 1.10(a) 中的组合 nMOS 开关, 只有当两个控制信号 S_1 和 S_2 都为高电平电压(通常为 V_{DD})时才导通, 而其他时候都处于截止状态。对于图 1.10(b) 中的组合 pMOS 开关, 只有当两个控制信号 S_1 和 S_2 都为低电平电压(通常为地电位)时才导通, 而其他时候都处于截止状态。

图 1.11 给出了两种将两个开关并联起来形成一个组合开关的情况。所得开关的工作由两个控制信号控制: S_1 和 S_2 。只要其中的任意一个开关导通, 组合开关就是导通的。因此, 只有当两个控制信号 S_1 和 S_2 都无效时组合开关才截止, 否则都处于导通状态。

对于图 1.11(a), 当控制信号 S_1 和 S_2 中任意一个为高电平电压(通常为 V_{DD})时组合

nMOS 开关都处于导通状态,只有当两个控制信号都为地电位时开关才处于截止状态。在图 1.11(b)中,当控制信号 S_1 和 S_2 中任意一个为低电平电压(通常为地电位)时组合 pMOS 开关都处于导通状态,只有当两个控制信号都为高电平电压时开关才处于截止状态。

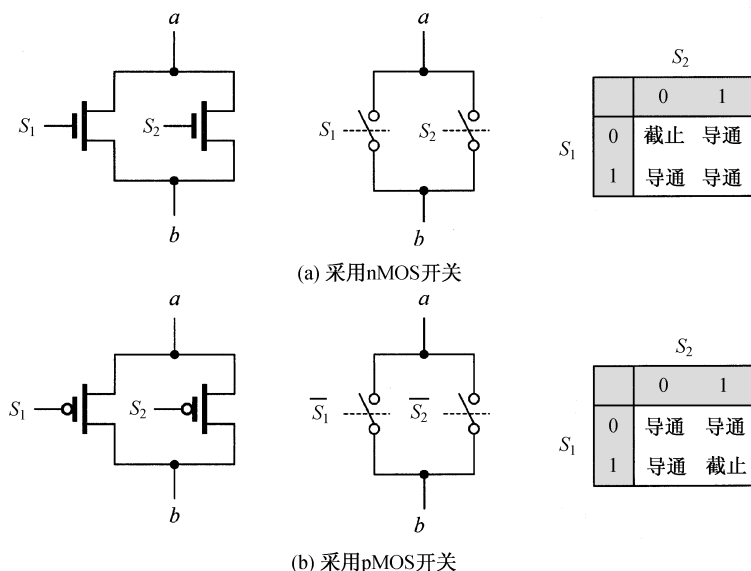


图 1.11 并联开关工作原理

1.2.4.2 $f/\bar{f}$ 模式

利用之前所述 nMOS 开关和 pMOS 开关特性来实现一个开关函数的一种简单逻辑图被称为 $f/\bar{f}$ 模式,完全 CMOS(FCMOS)逻辑,或简称 CMOS 逻辑。 $f/\bar{f}$ 模式背后的基本原理是根据下述观察结果得到的,即 pMOS 开关可以传输完美的逻辑 1 信号, nMOS 开关可以传输完美的逻辑 0 信号。因此,在实现真值开关函数 f 时采用 pMOS 开关,而实现互补开关函数 $\bar{f}$ 采用 nMOS 开关。

$f/\bar{f}$ 模式的框图如图 1.12 所示,其中两个模块分别指 f 模块和 $\bar{f}$ 模块。 f 模块实现真值开关函数,当其导通时将输出与 V_{DD} (即逻辑 1)相接,而 $\bar{f}$ 模块实现互补开关函数,当其导通时将输出与地(即逻辑 0)相接。输出节点的电容代表输出节点的内在寄生电容。

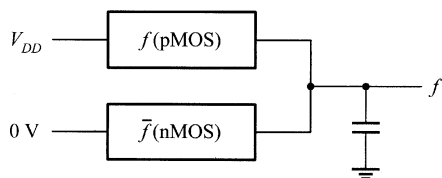


图 1.12 $f/\bar{f}$ 模式框图

接下来将列举两个例子来说明如何利用 $f/\bar{f}$ 模式来设计实际的开关逻辑电路。第一个例子是设计一个二输入与非门。

例 1-3 二输入与非门。二输入与非门的开关函数为 $f(x, y) = \overline{x \cdot y}$ 。利用 $f/\bar{f}$ 模式来设计并实现一个二输入与非门。

解: 根据德·摩根定律,有 $f(x, y) = \overline{x \cdot y} = \bar{x} + \bar{y}$ 。因此, f 模块由两个并联的 pMOS 晶体管实现。 $\bar{f}$ 模块由两个串联的 nMOS 晶体管实现,因为 f 的互补函数为 $\bar{f}(x, y) = x \cdot y$ 。最终的开关逻辑电路和 CMOS 逻辑电路分别如图 1.13(a)和(b)所示。

从图 1.13 可以很明显地看出 f 模块和 $\bar{f}$ 模块的功能是相互对偶的。下面的例子将进一步说明如何利用 $f/\bar{f}$ 模式来设计一个更为复杂的逻辑电路, 即与或非 (AOI) 门, 并同时采用 nMOS 开关和 pMOS 开关来实现。

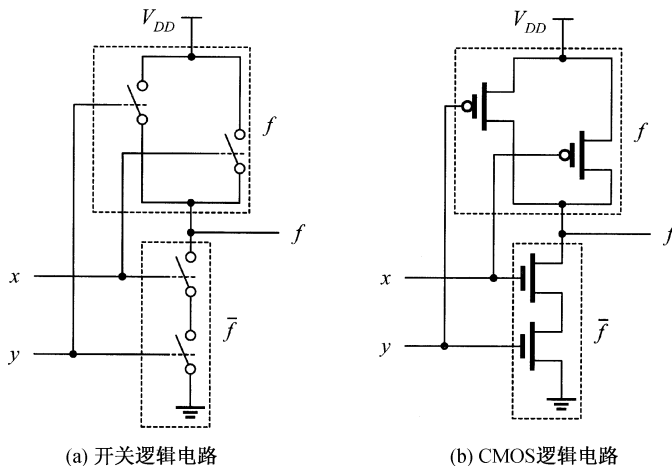


图 1.13 二输入与非门

例 1-4 与或非门。 利用 $f/\bar{f}$ 模式来实现如下开关函数: $f(w, x, y, z) = \overline{w \cdot x + y \cdot z}$ 。

解: 根据德·摩根定律, 有 $f(w, x, y, z) = \overline{w \cdot x} \cdot \overline{y \cdot z} = (\bar{w} + \bar{x}) \cdot (\bar{y} + \bar{z})$ 。因此, f 模块由两对并联连接的 pMOS 晶体管串联之后实现。 $\bar{f}$ 模块由两对串联连接的 nMOS 晶体管并联之后实现, 因为 f 的互补函数为 $\bar{f}(w, x, y, z) = w \cdot x + y \cdot z$ 。最终的开关逻辑电路和 CMOS 逻辑电路分别如图 1.14 (a) 和 (b) 所示。

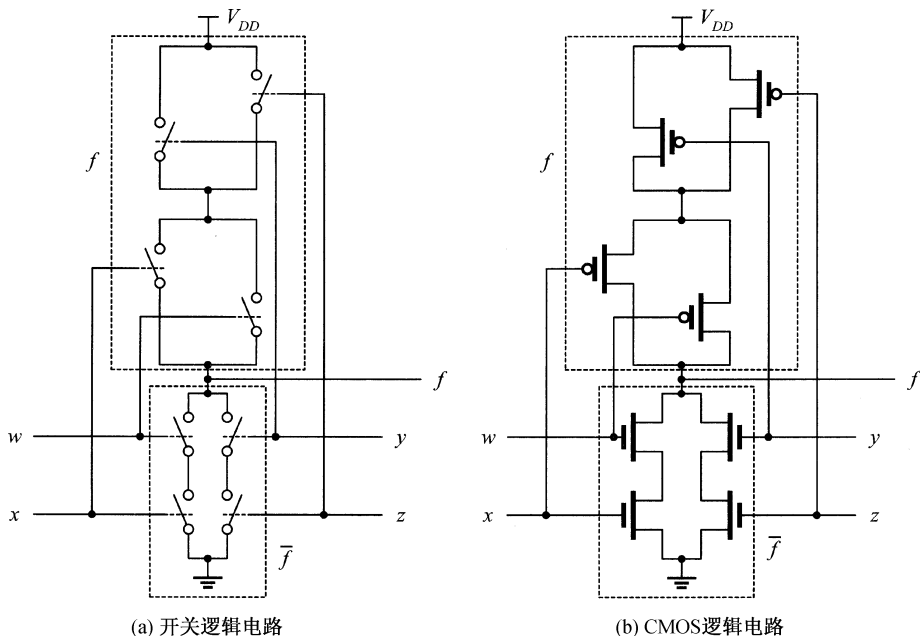


图 1.14 AOI 门

复习题

Q1-19 利用 $f/\bar{f}$ 模式设计一个 CMOS 二输入或非门。

Q1-20 利用 $f/\bar{f}$ 模式设计一个 CMOS 三输入与非门。

Q1-21 利用 $f/\bar{f}$ 模式设计一个 CMOS 三输入或非门。

1.2.5 CMOS 逻辑设计规则

CMOS 逻辑的基本设计规则可通过研究图 1.15(a) 和 (b) 中分别展示的门逻辑电路和开关逻辑电路之间的区别来说明。在图 1.15(a) 中, 如果两个输入 x 和 y 都为 1, 则输出 f 为 1, 否则输出 f 为 0; 在图 1.15(b) 中, 如果输入 x 和 y 都为 1, 则输出 f 为 1, 否则输出 f 为不确定状态。因此, 图 1.15(b) 中所示的开关逻辑电路无法准确地实现图 1.15(a) 所示的功能, 因为它只能实现 x 和 y 都为 1 时的功能, 而无法实现其他时候的功能。

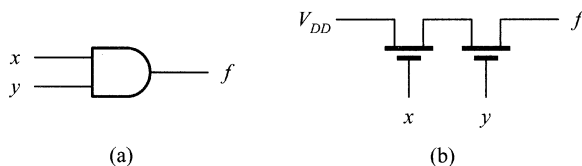


图 1.15 (a) 门逻辑电路和 (b) 开关逻辑电路之间的区别

1.2.5.1 两个基本规则

一个普通的逻辑电路有两个确定值: 逻辑 0(地) 和逻辑 1 (V_{DD})。为了让开关电路能够在所有的输入变量组合下都能输出一个确定的值, 需要遵循如下两个规则来完全利用 CMOS 开关正确地实现一个开关函数。

- 规则 1 (节点值规则): 在任何时间, 信号求和点 (例如 f) 必须一直接 0(地) 或 1 (V_{DD})。
- 规则 2 (节点无竞争规则): 信号求和点 (例如 f) 必须在任何时候都不能同时接 0(地) 或 1 (V_{DD})。

任何逻辑电路必须一直遵循规则 1 以确保正常工作。规则 2 将无比逻辑电路与有比逻辑电路区分开来。如果一个电路同时遵循规则 1 和规则 2, 则将其称为无比逻辑电路, 如果违反规则 2 则将其称为有比逻辑电路, 但是要合理设置上拉和下拉路径的尺寸。

无论 nMOS 开关、pMOS 开关或 TG 开关的相对尺寸如何, 无比逻辑电路总是能正确地执行所设计的开关函数。相反, 对于一个有比逻辑电路而言, 为了实现正确的函数, 电路中所用的 nMOS 开关、pMOS 开关或 TG 开关必须要进行合理设置。

1.2.5.2 开关函数的留数

CMOS 开关逻辑设计的规则遵循 Shannon 展开定理, 该定理如下所述。

定理 1.2.1 (Shannon 展开定理)。令 $f(x_{n-1}, \dots, x_{i+1}, x_i, x_{i-1}, \dots, x_0)$ 为一个具有 n 个变量的开关函数。则 f 可针对变量 x_i 分解成如下形式, 其中 $0 \leq i < n$:

$$f(x_{n-1}, \dots, x_{i+1}, x_i, x_{i-1}, \dots, x_0) = x_i \cdot f(x_{n-1}, \dots, x_{i+1}, 1, x_{i-1}, \dots, x_0) + \bar{x}_i \cdot f(x_{n-1}, \dots, x_{i+1}, 0, x_{i-1}, \dots, x_0) \quad (1.6)$$

Shannon 展开定理的证明非常琐碎，在此就忽略了。为方便起见，将 x_i 看做控制变量，将其他变量看做函数变量。

例 1-5 Shannon 展开定理。试考虑开关函数 $f(x, y, z) = xy + yz + xz$ ，将其按照变量 y 分解。

$$\begin{aligned} f(x, y, z) &= xy + yz + xz \\ &= y \cdot f(x, 1, z) + \bar{y} \cdot f(x, 0, z) \\ &= y \cdot (x + z + xz) + \bar{y} \cdot (xz) \\ &= xy + yz + x\bar{y}z \end{aligned}$$

容易证明，最后一行与原始开关函数其实是一样的。因而在此已经证明了 Shannon 展开定理的有效性。

为方便起见，通常会在字体上对一个变量加以区分。变量 x 是一个可具有两个值，即 1 和 0 的对象，它可以以两种形式出现，即 x 和 $\bar{x}$ 。一个字符表示一个变量，而无论这个字符是真值或补值形式。换句话说， x 和 $\bar{x}$ 是两个不同的字符，但却是同一个变量 x 。有了上述说明之后，可将开关函数 $f(X)$ 留数定义成集合 $X = \{x_{n-1}, x_{n-2}, \dots, x_1, x_0\}$ 的子集形式如下所述。

开关函数的留数

令 $X = \{x_{n-1}, x_{n-2}, \dots, x_1, x_0\}$ 为所有 n 个变量的集合，而 $Y = \{y_{m-1}, y_{m-2}, \dots, y_1, y_0\}$ 为 X 的一个子集，其中 $y_i \in X$ 且 $m \leq n$ 。 $f(X)$ 有关 Y 的留数用 $f_Y(X)$ 表示，它被定义为当 Y 中的所有补值字符变量都被设为 0、所有真值字符被设为 1 时的函数值。

基于这一定义，Shannon 展开定理可被看做是开关函数的两个留数的组合，即 $f(X) = x_i f_{x_i}(X) + \bar{x}_i f_{\bar{x}_i}(X)$ 。

例 1-6 开关函数的留数。计算出如下开关函数关于变量 x 的留数：

$$f(x, y, z) = xy + yz + xz$$

解：根据定义， f 关于变量 x 的留数可通过将 x 设为 1 并计算开关函数 f 的值来获得。即

$$f_x(1, y, z) = 1 \cdot y + yz + 1 \cdot z = x + y$$

因此， f 关于变量 x 的留数为 $x + y$ 。

根据留数的定义，一个开关函数可用两个留数的组合来表示。例如， $f(x, y, z) = x \cdot f_x(1, y, z) + \bar{x} \cdot f_{\bar{x}}(0, y, z)$ 。可以邀请有兴趣的读者对此进行验证。通过多次应用 Shannon 展开定理进一步分解留数函数，可形成一棵二叉树。这样一棵树通常被称为分解树。图 1.16 中给出了一个四变量开关函数 $f(w, x, y, z)$ 的分解树例子，这棵分解树是通过按顺序进行关于变量 w 和 x 的分解而得到的。于是开关函数 f 可以表示成 4 个留数的组合形式，如下所示。

$$\begin{aligned} f(w, x, y, z) &= w \cdot f_w(1, x, y, z) + \bar{w} \cdot f_{\bar{w}}(0, x, y, z) \\ &= w \cdot [x \cdot f_{wx}(1, 1, y, z) + \bar{x} \cdot f_{w\bar{x}}(1, 0, y, z)] + \\ &\quad \bar{w} \cdot [x \cdot f_{\bar{w}x}(0, 1, y, z) + \bar{x} \cdot f_{\bar{w}\bar{x}}(0, 0, y, z)] \\ &= wx f_{wx}(1, 1, y, z) + w\bar{x} f_{w\bar{x}}(1, 0, y, z) + \\ &\quad \bar{w}x f_{\bar{w}x}(0, 1, y, z) + \bar{w}\bar{x} f_{\bar{w}\bar{x}}(0, 0, y, z) \end{aligned} \quad (1.7)$$

分解树中的每一个内部节点都可用一个二选一的多路选择器来实现，其中将控制变量当做源选择变量，留数当做它的两个输入。最终所得电路被称为树网络。