

706003

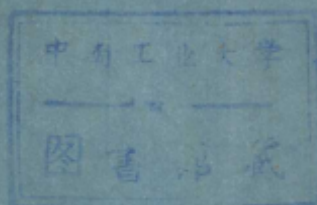
清

高等学校教材

激光全息检测技术

王永保 编

JI GUANG QUAN XI JIAN CE JI SHU



73

B

西北工业大学出版社

封面设计：杨君明

版式设计：潘玉浩

插图描绘：安 红



ISBN 7-5612-0182-6/O·17(课)

定 价：1.97 元

71
V

高等学校教材

激光全息检测技术

王永保 编

西北工业大学出版社

1989年10月 西安

内 容 简 介

本书是叙述激光全息技术应用的高等工科院校教材。全书共分六章。第一章介绍光波和原子结构的基础知识；第二章介绍激光形成的基本原理及特性；第三章介绍产生激光的器件及用于激光全息照相的主要激光器种类；第四章介绍激光全息照相的原理、装置和照相系统；第五、第六两章分别介绍激光全息照相的两个主要用途：无损检验和干涉计量术。其目的是让读者对激光全息照相的原理、器件、方法以及在检测技术领域中的应用有一个较全面和深入的了解。

本书主要作为高等工科院校各有关专业选修课教材；也可供从事激光检测工作的技术人员和管理人员学习和参考。

高等学校教材 激光全息检测技术

编 者 王永保
责任编辑 蒋相宗
责任校对 钱伟峰

西北工业大学出版社出版
(西安市友谊西路127号)
陕西省新华书店发行
西北工业大学出版社印刷厂印装
ISBN 7-5612-0182-6/O·17(课)

开本 787×1092毫米 1/16 9.5印张 229千字
1989年10月第1版 1989年10月第1次印刷
印数 1—4500册 定价: 1.97元

前 言

自1960年激光问世以来,激光技术的应用得到迅速发展。其中激光全息检测技术在整个激光技术应用领域中发展最早、最快,也是较为成熟的技术之一。这种技术已广泛应用于各个学科和工业部门。所以编写这本《激光全息检测技术》,可作为高等工业学校有关专业的选修课程教材,也可供从事激光检测工作的技术人员和管理人员学习和参考。

鉴于这一技术涉及的知识面较广,且处在不断发展阶段,所以在内容安排上力求简明、系统、完整,并尽量结合实际进行介绍。本书共分六章,第一章介绍光的二象性,这是了解激光产生原理和全息照相原理所必须具备的基础知识;第二、第三章介绍激光形成的基本概念和器件;第四章是介绍利用激光作为光源进行全息照相的基本原理、方法和具体照相过程及其装置;第五、第六章是介绍激光全息照相应用领域中两个主要应用方面。

本书是一本技术性较强的书籍,所以尽可能多地结合实际进行编写。书中所编资料及列举的实例,不少为我国有关单位及部门多年来进行研究的实际成果,在此谨向为本书提供宝贵资料及照片的有关单位和个人表示感谢。

本书由西安电子科技大学安毓英副教授进行审阅,并对本书提出不少宝贵意见,在此表示感谢。

由于编者水平有限,时间又较仓促,尤其是编写这种尚处在迅速发展之中的新技术方面的教材,困难较多。书中难免有许多缺点和错误,恳切期望读者批评指正。

编 者

1988年9月

目 录

第一章 基础知识	1
§ 1-1 光的波、粒二象性	1
§ 1-2 光的干涉	4
§ 1-3 光的衍射	9
§ 1-4 光的偏振	13
§ 1-5 物质的原子结构	16
第二章 激光产生的基本原理及特性	21
§ 2-1 光和物质的相互作用	21
§ 2-2 光学谐振腔	33
§ 2-3 激光振荡条件	38
§ 2-4 激光的模式	41
§ 2-5 激光的特性	45
第三章 激光器	53
§ 3-1 氦-氖激光器.....	53
§ 3-2 红宝石激光器	65
§ 3-3 氩离子激光器	75
§ 3-4 氦-镉激光器.....	79
第四章 激光全息照相	83
§ 4-1 基本原理	83
§ 4-2 激光全息照相的装置及要求	87
第五章 激光全息无损检验	103
§ 5-1 检验方法	103
§ 5-2 加载方法	110
§ 5-3 激光全息无损检验的应用	112
第六章 激光全息干涉计量	134
§ 6-1 材料的泊桑比测定	134
§ 6-2 全息照相测量机床结构特性	136
§ 6-3 全息干涉术在力学方面的应用	139
§ 6-4 全息干涉术测量风洞流场	143
参考文献	147

第一章 基础知识

§ 1-1 光的波、粒二象性

谈到光的重要性，人们会列举出光的很多用途。早在三千年前，人们对光就有所认识和研究。但是，要问光究竟是什么？光的本性是什么？要回答这个问题倒不是一件十分容易的事情。人们对光的本性认识，过去就经历过一个很长的时期，有过激烈地争论，而现在对光的本性认识也还在继续不断地深入和完善。

早在 17 世纪 60 年代，人们对光的认识就有两种截然不同的观点。一种是以英国物理学家牛顿 (I. Newton) 为代表的微粒说，认为光是从发光体发出来的，而且是以一定速度向空间传播的微粒；另一种观点是以荷兰科学家惠更斯 (C. Huygens) 为代表的波动说，认为光是某种振动，通过介质，以波动形式向四周传播。两种观点各有根据，都可解释一些光的现象。但在解释折射现象时，却出现了严重分歧。实验表明，当光线从光疏介质进入光密介质（例如从空气介质进入水介质）时，光线是折向法线的。为了解释这一现象，微粒说需要假设水中的光速大于空气中的光速，而波动说则需要假设水中的光速小于空气中的光速。但当时人们还不能准确地用实验方法来确定光速，因此无法从折射现象中判断两种学说的优劣。但由于微粒说能较自然地说明光的直线传播现象，因此，光的微粒说在当时占统治地位。

19 世纪初，杨氏 (T. Young) 用实验方法研究了光的干涉现象，证实了光的波动性质。1862 年，傅科 (J. L. Foucault) 用旋转镜法测定了光在不同介质中的传播速度，证明水中的光速小于空气中的光速，这是波动说的一个重要的论证。从而使惠更斯提出的光的波动说逐渐处于主导地位。但惠更斯没有说清波动过程的特性，没有指出光现象的周期性，也没有提到波长的概念。他把光看成是一种机械波，为了解决波的传播介质，就臆造出一种“以太”介质，它是一种所谓弹性媒质，充满于整个宇宙空间，光的传播取决于“以太”的弹性和密度，而实际上这种假想的“以太”是不存在的。

19 世纪 70 年代，麦克斯韦 (L. C. Maxwell) 认为光波是电磁波的一种，创造了光的电磁波理论来代替光的机械波理论。他直接从频率和波长来测定电磁波的传播速度，发现恰好等于光速，都是每秒三十万公里。根据电磁波理论，光波照射到物体表面时，物体表面要受到压力。1901 年，列别捷夫 (H. Лeбeдeв) 用实验方法测定了光压，其结果与理论十分符合，从而巩固了光的电磁波理论。

在电磁波中，通常用电矢量 E 来描述电场强度（简称电场），用磁矢量 B 来描述磁感应强度（简称磁场）； E 和 B 互相垂直，且都和电磁波传播方向 Z 垂直（图 1-1）。由于许多物理现象都是由光波的电场 E 所引起的，对人眼或感光仪器起作用的通常

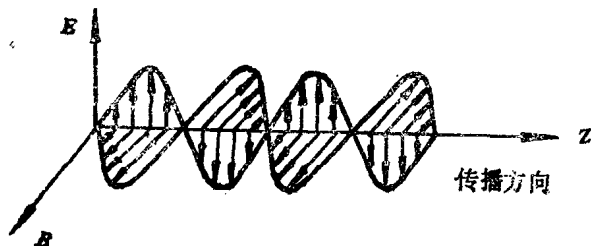


图 1-1 电磁波

是电场 E 。故以后凡是提到光波中的振动矢量时，都是指电矢量 E 。

当光波沿 Z 方向传播时，其两个相邻的波峰或波谷间的距离称为波长，以 λ 符号表示。光波的波长以微米 (μm) 作单位，一微米为 10^{-6} 米。人们的眼睛能够看到的光，其波长大约在 $0.39 \sim 0.76 \mu\text{m}$ 之间，在这个波长范围内的光称为可见光，它包含着红、橙、黄、绿、青、蓝、紫七种颜色。任何一定波长的光称为单色光。

与波长有密切联系的一个物理量是频率 ν ， ν 是在单位时间内通过一定点的波长的数目。因为波长就是波的空间周期，故通常以每秒若干周来表示频率的大小，频率的单位以赫兹表示（即每秒一周）。频率和波长相乘就是光传播的速度。如以 c 表示光速：

$$c = \lambda \cdot \nu \quad (1-1)$$

光速为 3×10^8 米/秒，对于红光来说，其波长为 $0.76 \sim 0.63$ 微米，故红光的频率为 $3.95 \times 10^{14} \sim 4.76 \times 10^{14}$ 赫兹。

频率为 ν 的电磁场在空间传播时，电场 E 随时间 t 的变化规律为：

$$E = A_0 \cos \omega t \quad \omega = 2\pi \nu \quad (1-2)$$

其中 A_0 是光波的振幅矢量，它代表电场 E 的最大值。振幅是描述光波振动的一个重要物理量，振幅的大小决定光的强弱，即光能的大小。对于任何波动，当它在介质中传播时，介质中各质点总要发生振动，因而具有动能；同时介质还要产生形变，因而具有位能。由此可见，波动的传播总是伴随着能量的传递。这种过程一般用能流密度来描述。所谓能流密度，就是光波在单位时间内通过与光波传播方向垂直的单位面积所传递的能量（或称光功率）。光波所传递的平均能流密度（即光的强度或辐照度 I ）是与光波振幅的平方成正比：

$$I \propto A_0^2$$

通常我们关心的是光强的相对分布，因此式子中比例关系的系数并不重要，为简单起见，可以把系数认为等于 1，于是将上式写成等式：

$$I = A_0^2 \quad (1-3)$$

这里的光强 I 应理解为相对强度。

描述光波振动的另一个重要的物理量是位相 ωt 。它表示在一个波长范围内各点的电场大小，位相都是以角度来表示的。图 1-2 表示光波的电场 E 与位相角的关系。在时刻 1 处的位相角为零， $E_1 = A$ ；时刻 2 处的位相角为 90 度， $E_2 = 0$ ；时刻 3 处的位相角为 180 度， $E_3 = -A$ ；时刻 4 处的位相角为 270 度， $E_4 = 0$ 。可见，当知道了某一时刻的位相角时，就可求出该处的电场大小。但是，位相这个物理量更主要的还是用来分析两个光波在叠加时的情形。例如有两个波长相同的光波叠加时，它们具有相同的偏振方向和相同的位相，也就是说两个光波的波峰和波谷是一一对应的，如图 1-3(a) 所示那样，当它们叠加在一起后，所合成的光波振幅就互相叠加而增强；如果两个光波的位相相反，也就是说一个光波的波峰落在另一个光波的波谷中，如图 1-3(b) 所示那样，则合成的光波振幅就互相抵消而减弱。因此，根据两个光波的位相是相同还是相反，就可得到合成后的光波是加强还是减弱。

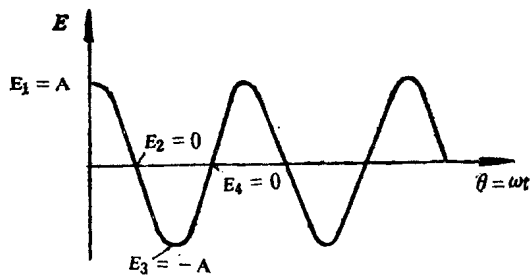


图 1-2 电场与位相角的关系

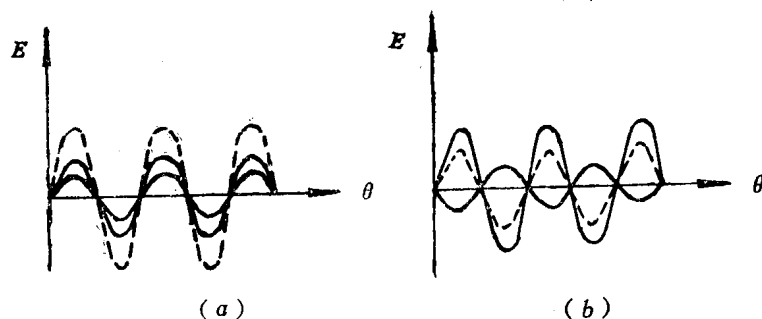


图 1-3 光波的叠加

从光波的波动方程中可以看出，振幅和位相是表征光波在传播中状态变化的两个基本物理量。当光波照射到物体上时，物体表面对于光波的反射就会使光波的传播状态发生变化，也就是使光波的振幅和位相发生各不相同的变化。而这种不同的变化就反映着物体各部分不同的特点，例如物体的前后位置、表面的凹凸不平以及反光本领的强弱等等。

从以上的叙述中，我们把光看成是一种波动——电磁波，这是从牛顿到麦克斯韦，人类对光的本性在认识上的一次飞跃，光作为一种波动已确定无疑了。

可是，人们的认识并没有到此为止。随着人们对于光的实践不断地加深加广，人们发现在研究光与物质的相互作用时，许多现象就无法用电磁波理论来解释。例如，光电效应便是一个突出的事例。

为了便于说明问题，我们先介绍一个实验。图 1-4 中有一个光电管，阴极上涂有感光层，并通过电流计、电池组与阳极串联起来，使阳极和阴极之间保持一个电位差。当阴极上没有光照射时，电路中就没有电流通过，电流计不偏转。当阴极被光照射时，电流计立刻偏转，表明电路中有电流通过，这是因为阴极上释放出来的电子，跑到电位较高的阳极上去而形成的。物体被光照射而释放出电子的这种现象称为光电效应，被释放出来的电子就叫做光电子，所产生的电流称为光电流。如果改变入射光的强度，光电流也改变，入射光强，光电流就大，因为被释放出来的光电子数目多，入射光弱，光电流就小，因为被释放出来的光电子数目少。但是，入射光的强度却不能改变光电子的速度和动能，也就是说光电子的能量与入射光的强弱无关。此外，如果用不同频率的光去照射阴极，发现被释放出来的电子能量是不一样的，说明其能量与入射光的频率有关，当入射光的频率越高时，光电子的能量也越大，而且，当入射光的频率低于某极限值（其值随不同的阴极材料而异）时，无论光的强度多大、照射时间多长，都没有光电子产生。

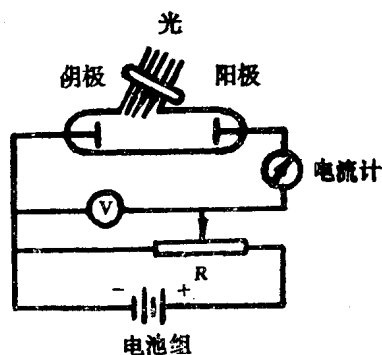


图 1-4 光电效应

上述这些现象，都无法用电磁波的概念来解释。因为按照电磁波的理论，从阴极释放出来的光电子，其能量应当与入射光的强度（振幅）有关，也就是说，入射光越强，光的能量越强，这些能量被阴极表面的电子所吸收，则电子脱离阴极以后就应该具有较大的能量才对，然而实验结果表明，光电子的能量只与入射光的频率有关，而与入射光的强度无关。显然，光电效应无法用电磁波理论来解释。

基于上述这些难以解释的实验结果,1905年爱因斯坦提出了光的量子理论。他认为物质的原子所发射和吸收的光,并不是连续的波,而是由特殊物质组成的一个个的微粒,也就是说,在讨论原子发射或吸收光的时候,必须把光看成是由大量以光速运动的粒子组成的粒子流。这些粒子叫做光子。每个光子具有的能量用 ϵ 表示, $\epsilon = h\nu$, h 为普朗克常数,实验测出, $h = 6.626 \times 10^{-34}$ 焦耳·秒。

不同频率的光,其光子的能量不同。频率越高,光子的能量越大。当光与物质(原子、电子等)相互作用交换能量时,光子只能整个地被原子、电子吸收或发射。每个光子的能量 $h\nu$ 是频率为 ν 的光子的最小能量单位。光的强弱取决于光子的数目。在光电效应中,一个光子只能打出一个光电子,频率越高的光子,其能量就大,打出来的光电子的能量也越大,如果光的频率过低,一个光子的能量不足以将一个光电子拉出阴极表面,那么,这种光子再多也不能打出光电子来,这就是为什么频率过低的光,即使光强再大,照射的时间再长也产生不了光电子的原因。这样,爱因斯坦利用光子假设成功地解释了光电效应的规律。这就是光子论的主要概念。

现在,人们已经接受了光子论,而且认识到光的本性是既具有波动性又具有微粒性,即具有波、粒二象性。这是人类对光的认识的第二个飞跃。当然这种认识还只是具有相对真理的意义。因为要问光究竟是什么?一个光子究竟是一个粒子还是一个波?仍然是一个无法回答的问题。近代物理实验已发现波长约为百分之一埃的光子(γ 射线)在强磁场中可转化为电子和正电子对的现象(正电子有和电子相同的质量和电量,但电荷是正的),这一现象揭示了光子和电子之间存在着深刻的联系。但这种联系在经典理论中无法作出解释,有些问题还需要进一步去探讨。因此,光究竟是“什么”的最终解答,只能通过人们不断地认识来获得。然而,我们却不可由此而否定今日关于光具有波粒二象性的认识。因为,这种认识确实是在一定条件下具有客观真理性的知识。一般地说来,在光的干涉现象中,光的波动性表现得较为明显,这时,往往把光看成是由一列一列的光波组成;而在原子发射或吸收光等现象中,光的微粒性表现得较为明显,这时,往往把光看成是由一个一个的光子组成的。下面在介绍激光是如何产生时,就用光的微粒性来解释;而在介绍激光全息照相术时,就用光的波动性来解释。

§1-2 光的干涉

光是电磁波,波的重要特征之一就是干涉现象。例如在太阳光下观察肥皂泡,泡上呈现出彩色条纹,水面上漂浮的油膜,当被太阳光照射时,也能观察到类似的现象,这些现象都是由于光的干涉现象所形成的。历史上最先为光的波动理论提供实验基础的就是光的干涉现象。

当两列光波在空间传播时,在它们相交的区域,各处的光强同各个光波单独作用所产生的光强之和可能是极不相同的,有些地方光强接近于零,而另一些地方的光强却大得多。我们把这种光波在空间叠加而形成明暗相间的稳定分布现象叫做光的干涉。光的干涉现象虽然不难被人们观察到,特别是激光出现以后,由于激光具有极好的相干性,所以光的干涉现象变得十分容易被观察到。但值得注意的是,并非任意两列光波相遇都能产生干涉现象。从两个完全独立的光源(如两盏白炽灯)发出的光波即使相遇,我们能看到的只是各处光强的增加,不发生明暗相间的条纹,不能产生光的干涉现象。仔细研究发光体的发光机构之后,人

们认识到, 只有设法从同一光源取得两个光束, 然后让这两束光在空间重叠, 才能出现光的干涉现象。能产生干涉现象的光称为相干光。两列光波要产生干涉现象必须满足一定的相干条件。

一、频率相同的两束光在相遇时有相同的振动方向和固定的位相差。这是产生干涉现象的必要条件。

下面分析双缝干涉, 从中可引导出光波干涉的规律性。

如图 1-5, 先使单色光通过狭缝 s , 再到达与 s 平行的两个狭缝 s_1 和 s_2 。为简便起见, 假定 s_1 和 s_2 的初位相都为 0, 且设 P 点距 s_1 为 r_1 , 若波速为 c , 则光波从 s_1 传播到 P 点所需时间为 r_1/c 。显然, 由于在 t 时刻处振动的位相为 ωt , 则此时在 P 点的振动将比 s_1 处

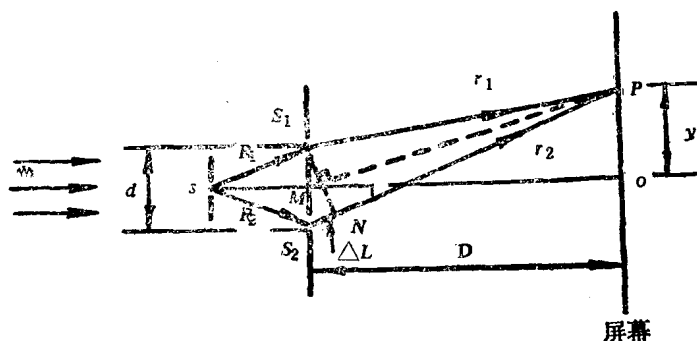


图 1-5 双缝干涉

落后 r_1/c 的时间, 也就是说, 此时 P 点的位相将是 $\omega\left(t - \frac{r_1}{c}\right)$ 。同理, 由于 t 时刻在 s_2 处振动的位相也为 ωt , 所以由 s_2 处传播到 P 点的振动在此时刻的位相就将是 $\omega\left(t - \frac{r_2}{c}\right)$, 此处 r_2 是 P 点到 s_2 的距离。可见, 在 P 点两光波的位相差 $\Delta\varphi$ 应为:

$$\begin{aligned}\Delta\varphi &= \omega\left(t - \frac{r_1}{c}\right) - \omega\left(t - \frac{r_2}{c}\right) \\ &= \omega\left(\frac{r_2 - r_1}{c}\right)\end{aligned}$$

因为

$$\omega = 2\pi\gamma$$

$$c = \gamma\lambda$$

所以

$$\Delta\varphi = 2\pi\left(\frac{r_2 - r_1}{\lambda}\right) \quad (1-4)$$

对一定的光波来说, 由于其频率相同, 即波长 λ 是一个常数, 所以位相差 $\Delta\varphi$ 完全取决于两个光波到达 P 点的距离之差 $\Delta L = (r_2 - r_1)$ 。 ΔL 通常被称为两列光波的光程差。

当光程差为波长整数倍时, 即

$$\Delta L = K\lambda \quad (K = 0, \pm 1, \pm 2, \dots)$$

时, 两列光波在 P 点处互相加强, 出现明亮的条纹。显然在屏幕中点 O 处, $\Delta L = 0$, 即 $K = 0$, 也将出现明亮条纹, 这叫做中央亮纹。其它在 $K = \pm 1, \pm 2, \dots$ 的明条纹, 分别称为第一级

明条纹、第二级明条纹……。

而当光程差为半波长的奇数倍时，即

$$\Delta L = \left(K\lambda + \frac{\lambda}{2} \right) \quad (K = 0, \pm 1, \pm 2, \dots)$$

时，两列光波互相抵消，出现暗条纹。由此可见，在屏幕上将出现一些明暗相间的干涉条纹。

从上面讨论说明，在分析干涉现象时，波长 λ 可以当作一把尺子，用它去度量两列光波到达某点的光程差，就可以确定两列光波在该点是加强(形成亮纹)，还是抵消(形成暗纹)。下面具体分析一下双光束干涉时的光强分布。

由于 s_1 和 s_2 相对于 s 是对称分布。所以 $R_1 = R_2$ 。另外， s_1 和 s_2 的电场不同于 s 处的电场。若记 s 处电场为

$$E_0 = A_0 \cos(2\pi\gamma t + \varphi_0) \quad (1-5)$$

则 s_1 和 s_2 处电场 E_{10} 和 E_{20} 分别为

$$E_{10} = A_{10} \cos(2\pi\gamma t + \varphi_{10}') \quad (1-6)$$

$$E_{20} = A_{20} \cos(2\pi\gamma t + \varphi_{20}') \quad (1-7)$$

其中

$$\varphi_{10}' = \varphi_0 - \frac{2\pi R_1}{\lambda}$$

$$\varphi_{20}' = \varphi_0 - \frac{2\pi R_2}{\lambda}$$

因为

$$R_1 = R_2$$

故可记

$$\varphi_0' = \varphi_{10}' = \varphi_{20}'$$

从 s_1 和 s_2 这两个次波源发出的光波，传播到屏幕上 P 点时，在该点产生的电场 E_1 和 E_2 分

别为

$$E_1 = A_1 \cos\left(2\pi\gamma t + \varphi_0' - \frac{2\pi r_1}{\lambda}\right) \quad (1-8)$$

$$E_2 = A_2 \cos\left(2\pi\gamma t + \varphi_0' - \frac{2\pi r_2}{\lambda}\right) \quad (1-9)$$

则 P 点合成的电场 E 为

$$\begin{aligned} E &= E_1 + E_2 \\ &= A_1 \cos\left(2\pi\gamma t + \varphi_0' - \frac{2\pi r_1}{\lambda}\right) \\ &\quad + A_2 \cos\left(2\pi\gamma t + \varphi_0' - \frac{2\pi r_2}{\lambda}\right) \\ &= \cos 2\pi\gamma t \left[A_1 \cos\left(\varphi_0' - \frac{2\pi r_1}{\lambda}\right) + A_2 \cos\left(\varphi_0' - \frac{2\pi r_2}{\lambda}\right) \right] \\ &\quad - \sin 2\pi\gamma t \left[A_1 \sin\left(\varphi_0' - \frac{2\pi r_1}{\lambda}\right) + A_2 \sin\left(\varphi_0' - \frac{2\pi r_2}{\lambda}\right) \right] \end{aligned} \quad (1-10)$$

记

$$\begin{aligned}\delta &= \left(\varphi_0' - \frac{2\pi r_1}{\lambda} \right) - \left(\varphi_0' - \frac{2\pi r_2}{\lambda} \right) \\ &= \frac{2\pi}{\lambda} (r_2 - r_1)\end{aligned}$$

并记

$$\cos \theta = \frac{A_1 \cos \left(\varphi_0' - \frac{2\pi r_1}{\lambda} \right) + A_2 \cos \left(\varphi_0' - \frac{2\pi r_2}{\lambda} \right)}{(A_1^2 + A_2^2 + 2A_1 A_2 \cos \delta)^{1/2}} \quad (1-11)$$

则有

$$\sin \theta = \frac{A_1 \sin \left(\varphi_0' - \frac{2\pi r_1}{\lambda} \right) + A_2 \sin \left(\varphi_0' - \frac{2\pi r_2}{\lambda} \right)}{(A_1^2 + A_2^2 + 2A_1 A_2 \cos \delta)^{1/2}} \quad (1-12)$$

将式(1-11)和(1-12)代入式(1-10)后得

$$E = [A_1^2 + A_2^2 + 2A_1 A_2 \cos \delta]^{1/2} \cos(2\pi \nu t + \theta) \quad (1-13)$$

即合成的电场 E 的振幅 A 为

$$A = [A_1^2 + A_2^2 + 2A_1 A_2 \cos \delta]^{1/2} \quad (1-14)$$

当单独照明 s_1 和 s_2 时, P 点的光强分别为 I_1 和 I_2

$$\begin{aligned}I_1 &= A_1^2 \\ I_2 &= A_2^2\end{aligned}$$

而叠加后的光强 I

$$\begin{aligned}I &= A^2 \\ &= I_1 + I_2 + 2\sqrt{I_1 I_2} \cos \delta\end{aligned} \quad (1-15)$$

式(1-15)表示由于存在位相差 δ , 光强度一般不是简单地相加 (即 $I \neq I_1 + I_2$), 这里多了一项 $2\sqrt{I_1 I_2} \cos \delta$ 。因为这一项与 δ 有关, 所以它的数值因 P 点位置而异, 且可正可负。这就是说屏幕上各点的光强重新分布了, 有些地方加强 ($I > I_1 + I_2$); 有些地方减弱 ($I < I_1 + I_2$), 也就是说发生了“干涉”。式中反应干涉效应的一项 $2\sqrt{I_1 I_2} \cos \delta$ 称为“干涉项”。

当 $A_1 = A_2$ 时, $I = 4I_1 \cos^2 \frac{\delta}{2}$ 。 I 随 δ 变化情况示于图 1-6 中。

从图中可以看出在屏幕中央 $\nu = 0$ 处, $r_1 = r_2$, 光程差为零, 故 $\delta = 0$, 所以出现亮条纹。随着 ν 值 (绝对值) 的增大, 光程差也不断增大, 结果形成明暗相间的条纹。一般讲:

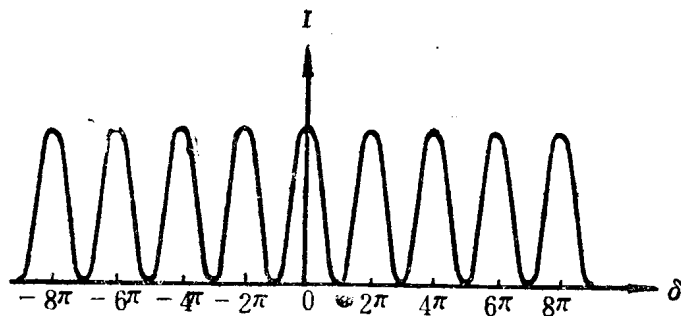


图 1-6 干涉光强分布

$$\Delta L = K\lambda, \quad \text{即 } \delta = 2K\pi \quad \text{亮纹 } (K = 0, \pm 1, \pm 2, \dots)$$

$$\Delta L = \left(K + \frac{1}{2}\right)\lambda, \quad \text{即 } \delta = (2K + 1)\pi \quad \text{暗纹 } (K = 0, \pm 1, \pm 2, \dots)$$

由图 1-5 中可看出

$$\Delta L = r_2 - r_1 = s_1 N$$

而

$$\Delta s_1 s_2 N \sim \Delta MPO$$

故

$$\Delta L : d = y : MP$$

$$\approx y : D$$

(1-16)

其中 d 为 s_1 s_2 之间距离。

由式(1-16)得出屏幕上亮暗条纹的位置是:

$$y = \frac{D}{d} K \lambda \quad \text{亮纹 } (K = 0, \pm 1, \pm 2, \dots)$$

$$y = \frac{D}{d} \left(K + \frac{1}{2}\right) \lambda \quad \text{暗纹 } (K = 0, \pm 1, \pm 2, \dots)$$

在普通光源中,光是由光源中各原子(或分子)所发射。各个原子(或分子)所发射的光互不关联,即各原子所发射的光的初位相 φ_0 不相同;而且每一原子所发射的光也是断续的,每次发光持续时间一般约 10^{-8} 秒;同一原子前后所发射的两列光波之间的初位相也彼此无关。所以普通光源所发射的光波为大量断续的、初位相完全无关的波列总和。

对两个初位相 φ_{10} 和 φ_{20} 不等的光波,其位相差 δ 应为

$$\delta = \varphi_{10} - \varphi_{20} + \frac{2\pi}{\lambda} \Delta L \quad (1-17)$$

因为 φ_{10} 、 φ_{20} 在普通光源中是在不断变化的,所以位相差 δ 也在发光持续时间(约 10^{-8} 秒)内发生一次变化;因而屏幕上的干涉强度 I 也以同样速度变化。而人眼或探测仪器一般都不能记录或察觉如此迅速的光强变化。记录或探测到的只是光强在一段时间中的平均值。

$$\text{即} \quad \overline{I} = I_1 + I_2 + 2\sqrt{I_1 I_2} \overline{\cos \delta} \quad (1-18)$$

而 δ 可以取 $(-\pi)$ 到 π 间的任意值,所以有

$$\overline{\cos \delta} = 0$$

$$I = I_1 + I_2 \quad (1-19)$$

即在屏幕上观察到的合成光强 I 是单个光源的光强简单相加,没有干涉现象。这种叠加叫非相干叠加。

但在上述双缝实验中,两束光波来自同一光源,虽然初位相不断在变化,但始终有 $\varphi_{10}' = \varphi_{20}'$, 所以位相差 δ 不随时间变化,故有干涉现象出现。这种叠加叫相干叠加。

通常将没有固定位相关系的光波叫非相干光,非相干光叠加时服从公式(1-19);而有固定位相关系的光波叫相干光,相干光叠加时服从公式(1-15)。

若 I_1 和 I_2 相等时,则 I 最大值和最小值为

$$\begin{aligned} I_{\text{最大}} &= I_1 + I_2 + 2\sqrt{I_1 I_2} \\ &= 4I_1 \end{aligned} \quad (1-20)$$

$$\begin{aligned} I_{\text{最小}} &= I_1 + I_2 - 2\sqrt{I_1 I_2} \\ &= 0 \end{aligned} \quad (1-21)$$

从上面讨论中说明, 两束相干光波在空间相遇时, 叠加后的光强取决于两束光波在该处的位相差。如果位相差为 π (即半个光波波长) 的偶数倍时, 光强为 $4I_1$ (当 $I_1 = I_2$ 时); 如果位相差为 π 的奇数倍时, 光强为零 (当 $I_1 = I_2$ 时)。但是, 当位相差为其它值时, 该点的光强可按公式(1-15)确定。

二、两束光波在相遇处所产生的振幅差不悬殊。

若两束光波的振幅相差悬殊, 例如 $A_1 \gg A_2$, 则该处合成振动的振幅 A 将同单一光波在该处的振幅 A_1 没有多大差别, 因此, 观察不出明显的干涉现象。

三、两束光波在相遇处的光程差不能太大。

因为实际光源所发出的光波, 并非是一个无限长的正弦波, 而是由一系列有限长的波列组成的。当两束光波在相遇点的光程差不大时, 则两束光波中有固定位相差的波列同时作用于一点, 能产生清晰的干涉现象。而当光程差很大时, 一束光波的波列已通过, 而另一束光波的相应波列尚未到达, 两个相应的波列间没有重叠, 因此, 产生不了干涉现象。

§ 1-3 光的衍射

衍射是波遇到障碍物时, 能绕过障碍物而到达它后面的现象。例如人在墙后, 照样能听到墙前传来的声音, 这就是声波的衍射效应。光是电磁波, 所以它和其它波一样会有衍射效应。只是由于光波的波长很短, 只有障碍很小时 (同波长可比较), 才有显著的衍射效应。

一、衍射现象

用单色光照射一个宽度小于 0.1 毫米的狭缝, 在离狭缝较远处用一屏幕接收。按几何光学原理, 在屏上看到的应是狭缝投影而成的一亮条, 如图 1-7 所示。但观察到的却是一组明暗交替的条纹。条纹近似为等距, 中心亮纹亮度最高, 逐渐依次减弱。这种现象称为光的衍射现象, 图 1-8 示出狭缝的衍射图案。在此, 光已绕到狭缝的几何阴影以外了, 这种光强度的

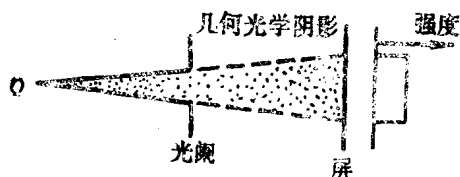


图 1-7 光阑后面的几何光学阴影区

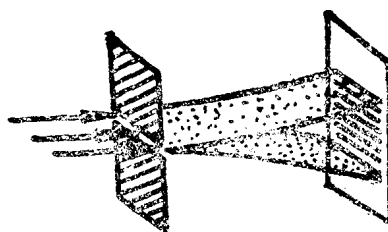


图 1-8 狭缝的衍射图案

分布是无法用直线光成像原理来解释的。由此可见光的“直线传播”概念只是在光通过宽度比波长大得多的缝时, 才能成立, 而当缝的几何尺寸接近波长的时候, 光的衍射效应就变得明显起来。

这种衍射现象可用惠更斯原理来解释。按照惠更斯原理, 光波波阵面上的每一点都是一个新的波源, 向外发射次波, 波阵面上各点的次波是相关的、同位相的。图 1-9 中点光源 Q 位于透镜焦点上, 通过透镜形成一个平面光波并照射在狭缝上。图 1-10 表示, 缝上平面光波形成的各次波之间还存在干涉效应。光屏上光强分布就是由狭缝上发出的次波相干叠加所决定的。因为光屏上不同的点, 各次波位相不同, 所以它们产生的干涉强度也不同。一些地方

是干涉极大（亮条纹），一些地方是干涉极小（暗条纹）。这就产生如图 1-8 中的衍射现象。不难看出，干涉和衍射本质上都是振动叠加的结果。不同的是前者为两个波或多个波的叠加，而后者是一个波的独自表演。

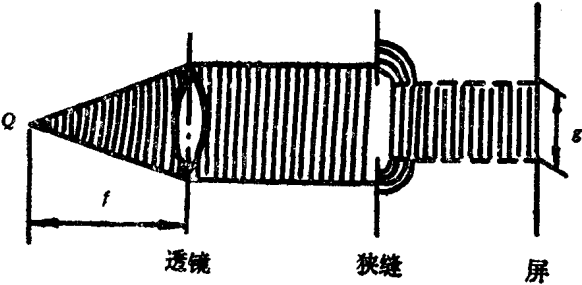


图 1-9 平面波的狭缝衍射

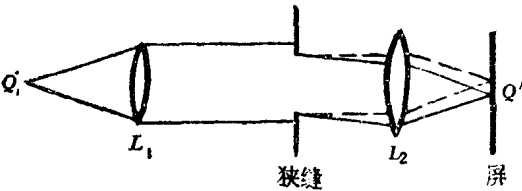


图 1-10 夫琅和费衍射装置

二、菲涅耳衍射和夫琅和费衍射

衍射问题分为两类。一类是光源和光屏（或两者之一）距离狭缝或光阑为有限远，如图 1-11(a)所示。在这种情况下，入射光线和衍射光线（或两者之一）不是平行光，这叫菲涅耳衍射。另一类是光源和光屏距狭缝或光阑都是无限远，因而入射光和衍射光都是平行光，这种情况称为夫琅和费衍射，在实际光路布置中，夫琅和费衍射条件可以利用两个会聚透镜来实现，如图1-11(b)所示。位于透镜前焦点上的光源Q所发出的光经透镜 L_1 形成平行光，而经过衍射后，透镜 L_2 把平行的衍射光线会聚在光屏上。在数学处理上，菲涅耳衍射比夫琅和费衍射麻烦得多。所以下面只讨论夫琅和费衍射。

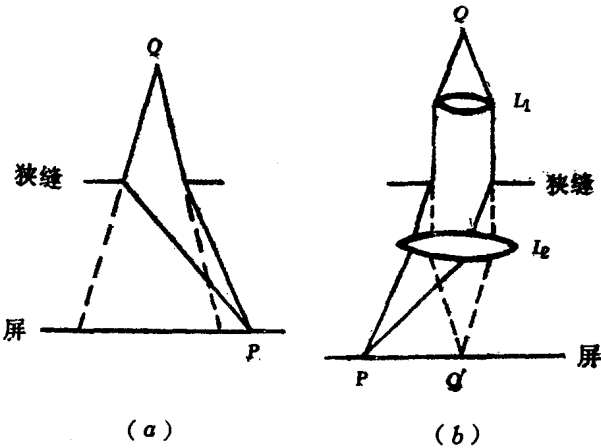


图 1-11 菲涅耳和夫琅和费衍射

三、单缝夫琅和费衍射

在计算单缝夫琅和费衍射花样时，设入射的平行光垂直照射在狭缝 BB' 上，缝宽为 b （图1-12）。由于狭缝高度比光波波长大得多，所以只需要考虑缝宽方面的次波叠加情况。

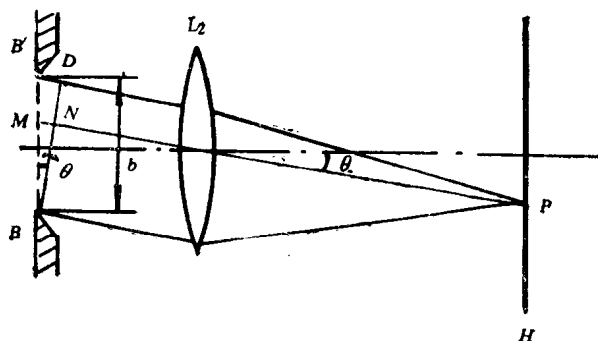


图 1-12 单缝夫琅和费衍射

为了计算 P 点的干涉强度，将缝的面积分为若干条平行的窄带，从每一条这样的窄带发出次波，其振幅正比于窄带的宽度 dx 。为简化计算起见，设光波的初位相为零； A_0 为通过整个狭缝光波的振幅。则狭缝处各窄带所发出次波的振动可用下式表示

$$dE = A_0 \frac{dx}{b} \cos \omega t$$

这些次波都可认为是球面波，各自向前传播。现只对其中沿图面与原入射方向成 θ 角的方向传播的各个次波进行讨论， θ 角为光屏 H 上 P 点和透镜中心连线与光轴的夹角，亦称衍射角。

令

$$\overline{BM} = x$$

则

$$\overline{MN} = x \sin \theta$$

$$\begin{aligned} dE &= \left(\frac{A_0}{b} dx \right) \cos(\omega t - kx \sin \theta) \\ &= \left(\frac{A_0}{b} dx \right) \cos(kx \sin \theta - \omega t) \end{aligned} \quad (1-22)$$

式中 $k = \frac{2\pi}{\lambda}$ 是波数

如果从 N 到 P 的光程为 Δ ，那么 P 点的振动可表示为

$$dE = \left(\frac{A_0}{b} dx \right) \cos[k(x \sin \theta + \Delta) - \omega t] \quad (1-23)$$

用复数形式表达，则为

$$dE = \frac{A_0}{b} dx e^{ikx \sin \theta} \cdot e^{ik\Delta} \cdot e^{-i\omega t} \quad (1-24)$$

从狭缝平面所有各点发出的次波到达 P 点而叠加，得到合振动，其值取决于上式 x 从 0 到 b 的积分。

$$E = \int_0^b dE = \frac{A_0}{b} e^{ik\Delta} e^{-i\omega t} \int_0^b e^{ikx \sin \theta} dx \quad (1-25)$$

因从 BD 平面上各点到达 P 点（经过透镜不同部分）的光程是相等的，都等于 Δ ，而与