

UNIT I OVER VIEW OF NDT

1.Introduction

In **destructive testing**, or (Destructive Physical Analysis DPA) tests are carried out to the specimens failure, in order to understand a specimens performance or material behaviour under different loads. These tests are generally much easier to carry out, yield more information, and are easier to interpret than nondestructive testing. Destructive testing is most suitable, and economic, for objects which will be mass-produced, as the cost of destroying a small number of specimens is negligible. It is usually not economical to do destructive testing where only one or very few items are to be produced (for example, in the case of a building). Analyzing and documenting the destructive failure mode is often accomplished using a high-speed camera recording continuously (movie-loop) until the failure is detected. Detecting the failure can be accomplished using a sound detector or stress gauge which produces a signal to trigger the high-speed camera. These high-speed cameras have advanced recording modes to capture almost any type of destructive failure.^[2] After the failure the high-speed camera will stop recording. The capture images can be played back in slow motion showing precisely what happen before, during and after the destructive event, image by image.

Some types of destructive testing:

- Stress tests
- Crash tests
- Hardness tests
- Metallographic tests

1.1 Materials Testing

Prior to manufacturing, many material, design, and production decisions are made to ensure product reliability and proper performance. To validate these decisions, a variety of testing methods are employed. The methods are grouped into two major categories:

☐ Mechanical Testing

☐ Non-Destructive Testing (NDT)

Mechanical testing, which is also known as destructive testing, is accomplished by forcing a part to fail by the application of various load factors. In contrast, non-destructive testing does not affect the part's future usefulness and leaves the part and its component materials in tact.

1.2 Mechanical Testing

Typically mechanical testing involves such attributes as hardness, strength, and impact toughness. Additionally, materials can be subjected to various types of loads such as tension or compression. Mechanical testing can occur at room temperatures or in either high or low temperature extremes.

Hardness

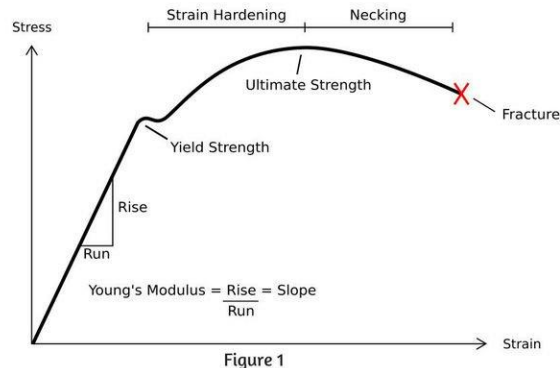
The resistance to indentation and to scratching or abrasion. The two most common hardness tests are the Brinell test and the Rockwell test.

In the Brinell hardness test, a known load is applied for a given period of time to a specimen surface using a hardened steel or tungsten-carbide ball, causing a permanent indentation. Standard ball diameter is 10 millimeters, or approximately four-tenths of an inch. The diameter of the resulting permanent indentation is then measured and converted to a Brinell hardness number. The Rockwell hardness test involves the use of an indenter for penetrating the surface of a material first by applying a minor, or initial load, and then applying a major, or final load under specific conditions. The difference between the minor and major penetration depths is then noted as a hardness value directly from a dial or digital readout. The harder the material the higher the number.

Tensile Test

Tensile – Force is applied perpendicular to the cross sectional area of the test item. Two of the primary material properties that tensile tests determine are:

Yield Strength, which is the stress required to permanently elongate, or deform, a material a specific amount, commonly 0.2% of total elongation.



Ultimate Tensile Strength, which is the maximum stress a material can withstand just prior to fracturing.

Compression – Compressive loads are applied to a point just beyond the yield strength of the material and measured at that point or continued to the point of failure if required.

Impact – Impact tests measure resistance to shock loading or impact by determining the amount of energy absorbed by the test specimen. There are two basic types of impact tests:

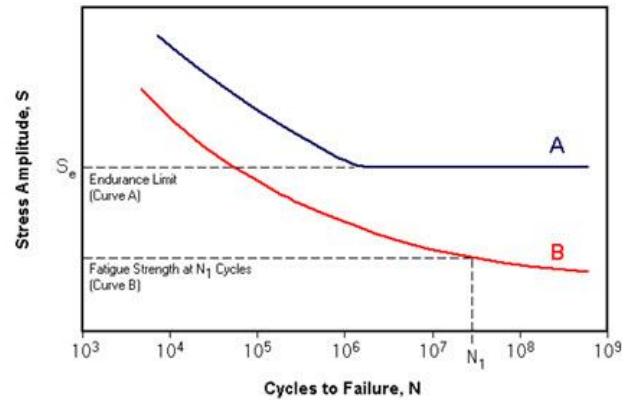
☐ ☐ Pendulum

☐ ☐ Drop Weight

Most common pendulum impact tests are the Charpy notched-bar impact test and the Izod notched-bar impact test. In both tests, the specimen is fractured and the energy absorbed is documented. The chief differences between these two impact tests are the way the test specimen is held and in the pendulum hammer design. In the dropped weight test, a known weight is dropped from a specified height. Such tests have advantages in that the impact is unidirectional with failure beginning at the weakest point and propagating from there. A principle advantage over the pendulum impact tests is that the drop weight impact test can define failure by either deformation, crack initiation, or complete failure of the specimen.

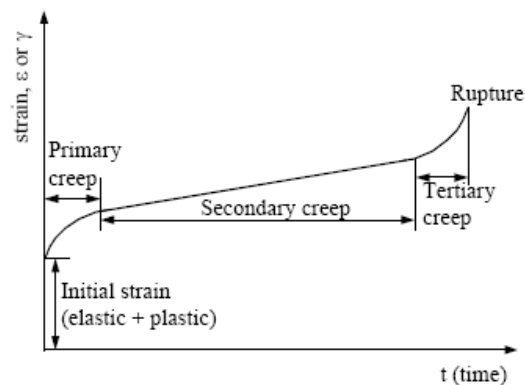
Fracture Toughness - Measures a material's resistance to brittle fracture and can be quantified by linear elastic fracture mechanics.

Fatigue – Measures material failure under repeated loading below the yield strength. Stresses measured below failure is referred to as the 'endurance limit' while the number of repeating cycles the material can withstand above this limit is known as 'fatigue life.'



Creep test

Measures a material's continuing dimensional change while under timed stress load. Creep tests are usually performed at elevated temperatures and can last for a thousand hours or longer.



1.3 Properties

- **Brittleness:** Ability of a material to break or shatter without significant deformation when under stress; opposite of plasticity
- **Compressive strength:** Maximum stress a material can withstand before compressive failure (MPa)
- **Creep:** The slow and gradual deformation of an object with respect to time
- **Ductility:** Ability of a material to deform under tensile load (% elongation)
- **Elasticity:** Ability of a body to resist a distorting influence or stress and to return to its original size and shape when the stress is removed
- **Fatigue limit:** Maximum stress a material can withstand under repeated loading (MPa)
- **Fracture toughness:** Ability of a material containing a crack to resist fracture (J/m^2)
- **Hardness:** Ability to withstand surface indentation and scratching (e.g. Brinnell hardness number)
- **Plasticity:** Ability of a material to undergo irreversible or permanent deformations without breaking or rupturing; opposite of brittleness
- **Poisson's ratio:** Ratio of lateral strain to axial strain (no units)
- **Resilience:** Ability of a material to absorb energy when it is deformed elastically (MPa); combination of strength and elasticity

- **Stiffness:** Ability of an object resists deformation in response to an applied force; rigidity; complementary to flexibility
- **Tensile strength:** Maximum tensile stress a material can withstand before failure (MPa)
- **Toughness:** Ability of a material to absorb energy (or withstand shock) and plastically deform without fracturing (or rupturing); a material's resistance to fracture when stressed; combination of strength and plasticity

1.4 NON-DESTRUCTIVE TESTING

Up to this point we have learnt various testing methods that somehow destruct the test specimens. These were, tensile testing, hardness testing, etc. In certain applications, the evaluation of engineering materials or structures without impairing their properties is very important, such as the quality control of the products, failure analysis or prevention of the engineered systems in service. This kind of evaluations can be carried out with Non destructive test (NDT) methods. It is possible to inspect and/or measure the materials or structures without destroying their surface texture, product integrity and future usefulness. The field of NDT is a very broad, interdisciplinary field that plays a critical role in inspecting that structural component and systems perform their function in a reliable fashion. Certain standards has been also implemented to assure the reliability of the NDT tests and prevent certain errors due to either the fault in the equipment used, the miss-application of the methods or the skill and the knowledge of the inspectors. Successful NDT tests allow locating and characterizing material conditions and flaws that might otherwise cause planes to crash, reactors to fail, trains to derail, pipelines to burst, and variety of less visible, but equally troubling events. However, these techniques generally require considerable operator skill and interpreting test results accurately may be difficult because the results can be subjective. These methods can be performed on metals, plastics, ceramics, composites, cermets, and coatings in order to detect cracks, internal voids, surface cavities, delamination, incomplete crack defective welds and any type of flaw that could lead to premature failure. Commonly used

Visual inspection:

VI is particularly effective detecting macroscopic flaws, such as poor welds. Many welding flaws are macroscopic: crater cracking, undercutting, slag inclusion, incomplete penetration welds, and the like. Like wise, VI is also suitable for detecting flaws in composite structures and piping of all types. Essentially, visual inspection should be performed the way that one would inspect a new car prior to delivery, etc. Bad welds or joints, missing fasteners or components, poor fits, wrong dimensions, improper surface finish, delaminations in coatings, large cracks, cavities, dents, inadequate size, wrong parts, lack of code approval stamps and similar proofs of testing.

Radiography:

Radiography has an advantage over some of the other processes in that the radiography provides a permanent reference for the internal soundness of the object that is radiographed. The x-ray emitted from a source has an ability to penetrate metals as a function of the accelerating voltage in the x-ray emitting tube. If a void present in the object being radiographed, more x-rays will pass in that area and the film under the part in turn will have more exposure than in the non-void areas. The sensitivity of x-rays is nominally 2% of the materials thickness. Thus for a piece of steel with a 25mm thickness, the smallest void that could be detected would be 0.5mm in dimension. For this reason, parts are often radiographed in different planes. A thin crack does not show up unless the x-rays ran parallel to the plane 0 the crack. Gamma radiography is identical to x-ray radiography in function. The difference is the source of the penetrating electromagnetic radiation which is a radioactive material such m Co 60. However this method is less popular because of the hazards of handling radioactive materials.

Liquid (Dye) penetrant method:

Liquid

penetrant inspection (LPI) is one of the most widely used nondestructive evaluation (NDE) methods. Its popularity can be attributed to two main factors, which are its relative ease of use and its flexibility. The technique is based on the ability of a liquid to be drawn into a "clean" surface breaking flaw by capillary action. This method is an inexpensive and convenient technique for surface defect inspection. The limitations of the liquid penetrant technique include the inability to inspect subsurface flaws and a loss of resolution on porous materials. Liquid penetrant testing is largely used on nonmagnetic materials for which magnetic particle inspection is not possible. Materials that are commonly inspected using LPI include the following; metals (aluminum, copper, steel, titanium, etc.), glass, many ceramic materials, rubber, plastics. Liquid penetrant inspection is used to inspect of flaws that break the surface of the sample. Some of these flaws are listed below; fatigue cracks, quench cracks grinding cracks, overload and impact fractures, porosity, laps seams, pin holes in welds, lack of fusion or braising along the edge of the bond line.

Magnetic particles:

Magnetic particle inspection is one of the simple, fast and traditional nondestructive testing methods widely used because of its convenience and low cost. This method uses magnetic fields and small magnetic particles, such as iron filings to detect flaws in components. The only requirement from an inspect ability standpoint is that the component being inspected must be made of a ferromagnetic material such iron, nickel, cobalt, or some of their alloys, since these materials are materials that can be magnetized to a level that will allow the inspection to be effective. On the other hand, an enormous volume of structural steels used in engineering is magnetic. In its simplest application, an electromagnet yoke is placed on the surface of the part to be examined, a kerosene-iron filling suspension is poured on the surface and the electromagnet is energized. If there is a discontinuity such as a crack or a flaw on the surface of the part, magnetic flux will be broken and a new south and north pole will form at each edge of the discontinuity. Then just like if iron particles are scattered on a cracked magnet, the particles will be attracted to and cluster at the pole ends of the magnet, the iron particles will also be attracted at the edges of the crack behaving poles of the magnet. This cluster of particles is much easier to see than the actual crack and this is the basis for magnetic particle inspection. For the best sensitivity, the lines of magnetic force should be perpendicular to the defect.

Eddy current testing:

Eddy currents are created through a process called electromagnetic induction. When alternating current is applied to the conductor, such as copper wire, a magnetic field develops in and around the conductor. This magnetic field expands as the alternating current rises to maximum and collapses as the current is reduced to zero. If another electrical conductor is brought into the close proximity to this changing magnetic field, current will be induced in this second conductor. These currents are influenced by the nature of the material such as voids, cracks, changes in grain size, as well as physical distance between coil and material. These currents form an impedance on a second coil which is used to as a sensor. In practice a probe is placed on the surface of the part to be inspected, and electronic equipment monitors the eddy current in the work piece through the same probe. The sensing circuit is a part of the sending coil. Eddy currents can be used for crack detection, material thickness measurements, coating thickness measurements, conductivity measurements for material identification, heat damage detection, case depth determination, heat treatment monitoring. Some of the advantages of eddy current inspection include; sensitivity to small cracks and other defects, ability to detect surface and near surface defects, immediate results, portable equipment, suitability for many different applications, minimum part preparation, no necessity to contact the part under inspection, ability to inspect complex shapes and sizes of conductive materials. Some limitation of eddy current inspection; applicability just on conductive materials, necessity for an accessible surface to the probe, skillful and trained personal, possible interference of surface finish and roughness, necessity for reference standards for setup, limited depth of

penetration, inability to detect of the flaws lying parallel to the probe coil winding and probe scan direction.

Ultrasonic Inspection:

Ultrasonic Testing (UT) uses a high frequency sound energy to conduct examinations and make measurements. Ultrasonic inspection can be used for flaw detection I evaluation, dimensional measurements, material characterization, and more. A typical UT inspection system consists of several functional units, such as the pulser/receiver, transducer, and display devices. A pulser/receiver is an electronic device that can produce high voltage electrical pulse. Driven by the pulser, the transducer of various types and shapes generates high frequency ultrasonic energy operating based on the piezoelectricity technology with using quartz, lithium sulfate, or various ceramics. Most inspections are carried out in the frequency rang of 1 to 25MHz. Couplants are used to transmit the ultrasonic waves from the transducer to the test piece; typical couplants are water, oil, glycerin and grease. The sound energy is introduced and propagates through the materials in the form of waves and reflected from the opposing surface. An internal defect such as crack or void interrupts the waves' propagation and reflects back a portion of the ultrasonic wave. The amplitude of the energy and the time required for return indicate the presence and location of any flaws in the work-piece. The ultrasonic inspection method has high penetrating power and sensitivity. It can be used from various directions to inspect flaws in large parts, such as rail road wheels pressure vessels and die blocks. This method requires experienced personnel to properly conduct the inspection and to correctly interpret the results. As a very useful and versatile NDT method, ultrasonic inspection method has the following advantages; sensitivity to both surface and subsurface discontinuities, superior depth of penetration for flaw detection or measurement, ability to single-sided access for pulse-echo technique, high accuracy in determining reflector position and estimating size and shape, minimal part preparation, instantaneous results with electronic equipment, detailed imaging with automated systems, possibility for other uses such as thickness measurements. Its limitations; necessity for an accessible surface to transmit ultrasound, extensive skill and training, requirement for a coupling medium to promote transfer of sound energy into test specimen, limits for roughness, shape irregularity, smallness, thickness or not homogeneity, difficulty to inspect of coarse grained materials due to low sound transmission and high signal noise, necessity for the linear defects to be oriented parallel to the sound beam, necessity for reference standards for both equipment calibration, and characterization of flaws.

Acoustic Method:

There are two different kind of acoustic methods: (a) acoustic emission; (b) acoustic impact technique.

Acoustic emission:

This technique is typically performed by elastically stressing the part or structure, for example, bending a beam, applying torque to a shaft, or pressurizing a vessel and monitoring the acoustic responses emitted from the material. During the structural changes the material such as plastic deformation, crack initiation, and propagation, phase transformation, abrupt reorientation of grain boundaries, bubble formation during boiling in cavitation, friction and wear of sliding interfaces, are the source of acoustic signals. Acoustic emissions are detected with sensors consisting of piezoelectric ceramic elements. This method is particularly effective for continuous surveillance of load-bearing structures.

Acoustic impact technique:

This technique consists of tapping the surface of an object and listening to and analyzing the signals to detect discontinuities and flaws. The principle is basically the same as when one taps walls, desktops or countertops in various locations with a finger or a hammer and listens to the sound emitted. Vitrified grinding wheels are tested in a similar manner to detect cracks in the wheel that may not be visible to the naked eye. This technique is easy to perform and can be instrumented and automated.

However, the results depend on the geometry and mass of the part so a reference standard is necessary for identifying flaws.

Test	Application	Limitation
Visual Inspection	Macroscopic surface flaws Small flaws are difficult to detect, no subsurface flaws.	Microscopy surface flaws Not applicable to larger structures; no subsurface flaws.
Dye penetrate	Surface flaws No subsurface	flaws not for porous materials
Magnetic Particle	Surface / near surface and layer flaws.	Limited subsurface capability, only for ferromagnetic materials
Ultrasonic	Subsurface flaws Material must be good conductor of sound.	
Eddy Current	Surface and near surface flaws	Difficult to interpret in some applications; only for metals
Radiography	Subsurface flaws Smallest defect detectable is 2% of the thickness;	radiation protection No subsurface flaws not for porous materials
Acoustic emission	Can analyze entire structure	Difficult to interpret, expensive equipments.

1.5 NON-DESTRUCTIVE TESTING – VISUAL TESTING – GENERAL PRINCIPLES

Direct visual testing

Direct visual testing may usually be made for local visual testing when access is sufficient to place the eye within 600 mm of the surface to be tested and at an angle not less than 30° to the surface to be tested. Mirrors may be used to improve the angle of vision, and aids such as a magnifying lens, endoscope and fibre optic may be used to assist testing.

Direct visual testing may also be made at greater distances than 600 mm specifically for general visual testing. A viewing distance appropriate to the test shall be used. The specific part, component, vessel, or section thereof, under immediate test, shall be illuminated, if necessary, with auxiliary lighting, to attain a minimum of 160 lx for general visual testing and a minimum of 500 lx for local visual testing.

Consideration shall be given to the application of illuminance to maximize the effectiveness of the test by:

- using the optimum direction of light with respect to the viewing point;
- avoiding glare;
- optimizing the colour temperature of the light source;
- using an illumination level compatible with the surface reflectivity.

Remote visual testing

When direct visual testing cannot be utilized, remote visual testing may have to be substituted. Remote visual testing uses visual aids such as endoscopes and fibre optics, coupled to cameras or other suitable instruments. The suitability of the remote visual testing system to perform the designated task shall be proven.

Personnel

Personnel who carry out tests according to this standard shall be shown to:

- a) be familiar with relevant standards, rules, specifications, equipment procedures/instructions;
 - b) be familiar with the relevant manufacturing procedure used and/or with the operating conditions of the component to be tested;
 - c) have satisfactory vision in accordance with EN 473. In addition, when performing general visual testing far vision shall be checked using the standard optotype in
- All visual tests shall be evaluated in terms of the acceptance criteria specified (e.g. product standard, contract).

Post-test documentation

When required (e.g. product standard, contract) a written test report shall be provided detailing the following:

- a) date and place of test;
- b) method used according to clauses 5 or 6;
- c) acceptance criteria and/or written procedure/instruction reference;
- d) equipment and/or system utilized including set-up;
- e) reference to customer's order;
- f) name of organization carrying out test;
- g) description and identification of test object;
- h) details of test findings with respect to the acceptance criteria (e.g. size, location);
- i) extent of test coverage;
- j) name and signature of person conducting test with date;
- k) name and signature of person supervising test with date, if required;
- l) marking of component tested, when appropriate;
- m) results.

This may be accomplished by referencing the visual testing written procedure and/or the instruction.

Records

Records shall be maintained as required (e.g. product standard, contract).

UNIT II SURFACE NDE METHODS

2.1 Liquid penetrant method:

In this method the surfaces to be inspected should be free from any coatings, paint, grease, dirt, dust, etc., therefore, should be cleaned with an appropriate way. Special care should be taken not to give additional damage to the surface to be inspected during the cleaning process. Otherwise, the original nature of surface could be disturbed and the results could be erroneous with the additional interferences of the surface features formed during the cleaning process. Surface cleaning can be performed with alcohol. Special chemicals like cleaner-remover can also be applied if needed. In the experiment, only cleaner-remover will be sufficient. Subsequent to surface cleaning, the surface is let to dry for 2 minutes. Commercially available cans of liquid penetrant dyes with different colors are used to reveal the surface defects.

2.2 Steps used in the experiment

:

- Clean the surface with alcohol and let surface dry for 5 min.
- Apply the liquid penetrant spray (red can) to the surface and brush for further penetration. Then, wait for 20 min.
- Wipe the surface with a clean textile and subsequently apply remover spray (blue can) to remove excess residues on the surface and wait for a few min.
- Apply the developer spray (yellow can) at a distance of about 30cm from the surface. The developer will absorb the penetrant that infiltrated to the surface features such as cracks, splits, etc., and then react with it to form a geometric shape which is the negative of the
- geometry of the surface features from which the penetrant is sucked.
- The polymerized material may be collected on a sticky paper for future evaluation and related documentation, if needed

2.3 Penetrants

Penetrants are carefully formulated to produce the level of sensitivity desired by the inspector. The penetrant must possess a number of important characteristics:

- spread easily over the surface of the material being inspected to provide complete and even coverage.
- be drawn into surface breaking defects by capillary action.
- remain in the defect but remove easily from the surface of the part.
- remain fluid so it can be drawn back to the surface of the part through the drying and developing steps.
- be highly visible or fluoresce brightly to produce easy to see indications.
- not be harmful to the material being tested or the inspector.

Penetrant materials are not designed to perform the same. Penetrant manufacturers have developed different formulations to address a variety of inspection applications. Some applications call for the detection of the smallest defects possible while in other applications, the rejectable defect size may be larger. The penetrants that are used to detect the smallest defects will also produce the largest amount of irrelevant indications. Standard specifications classify penetrant materials according to their physical characteristics and their performance.

- ☐ Penetrant materials come in two basic types:

Type 1 - Fluorescent Penetrants: they contain a dye or several dyes that fluoresce when exposed to ultraviolet radiation.

Type 2 - Visible Penetrants: they contain a red dye that provides high contrast against the white developer background.

Fluorescent penetrant systems are more sensitive than visible penetrant systems because the eye is drawn to the glow of the fluorescing indication. However, visible penetrants do not require a darkened area and an ultraviolet light in order to make an inspection.

- ☐ Penetrants are then classified by the method used to remove the excess penetrant from the part. The four methods are:

Method A - Water Washable: penetrants can be removed from the part by rinsing with water alone. These penetrants contain an emulsifying agent (detergent) that makes it possible to wash the penetrant from the part surface with water alone. Water washable penetrants are sometimes referred to as self-emulsifying systems.

Method B - Post-Emulsifiable, Lipophilic: the penetrant is oil soluble and interacts with the oil-based emulsifier to make removal possible.

Method C - Solvent Removable: they require the use of a solvent to remove the penetrant from the part.

Method D - Post-Emulsifiable, Hydrophilic: they use an emulsifier that is a water soluble detergent which lifts the excess penetrant from the surface of the part with a water wash.

Penetrants are then classified based on the strength or detectability of the indication that is produced for a number of very small and tight fatigue cracks. The five sensitivity levels are:

Level ½ - Ultra Low Sensitivity

Level 1 - Low Sensitivity

Level 2 - Medium Sensitivity

Level 3 - High Sensitivity

Level 4 - Ultra-High Sensitivity

The procedure for classifying penetrants into one of the five sensitivity levels uses specimens with small surface fatigue cracks. The brightness of the indication produced is measured using a photometer.

2.4 Developers

The role of the developer is to pull the trapped penetrant material out of defects and spread it out on the surface of the part so it can be seen by an inspector. Developers used with visible penetrants create a white background so there is a greater degree of contrast between the indication and the surrounding background. On the other hand, developers used with fluorescent penetrants both reflect and refract the incident ultraviolet light, allowing more of it to interact with the penetrant, causing more efficient fluorescence.

According to standards, developers are classified based on the method that the developer is applied (*as a dry powder, or dissolved or suspended in a liquid carrier*). The six standard forms of developers are:

Form a - Dry Powder

Form b - Water Soluble

Form c - Water Suspendable

Form d - Nonaqueous Type 1: Fluorescent (Solvent Based)

Form e - Nonaqueous Type 2: Visible Dye (Solvent Based)

Form f - Special Applications

Dry Powder

Dry powder developers are generally considered to be the least sensitive but they are inexpensive to use and easy to apply. Dry developers are white, fluffy powders that can be applied to a thoroughly dry surface in a number of ways; by dipping parts in a container of developer, by using a puffer to dust parts with the developer, or placing parts in a dust cabinet where the developer is blown around. Since the powder only sticks to areas of indications since they are wet, powder developers are seldom used for visible inspections.

Water Soluble

As the name implies, water soluble developers consist of a group of chemicals that are dissolved in water and form a developer layer when the water is evaporated away. The best method for applying water soluble developers is by spraying it on the part. The part can be wet or dry. Dipping, pouring, or brushing the solution on to the surface is sometimes used but these methods are less desirable. Drying is achieved by placing the wet, but well drained part, in a recirculating warm air dryer with a temperature of 21°C. Properly developed parts will have an even, light white coating over the entire surface.

Water Suspendable

Water suspendable developers consist of insoluble developer particles suspended in water. Water suspendable developers require frequent stirring or agitation to keep the particles from settling out of suspension. Water suspendable developers are applied to parts in the same manner as water soluble developers then the parts are dried using warm air.

Nonaqueous

Nonaqueous developers suspend the developer in a volatile solvent and are typically applied with a spray gun. Nonaqueous developers are commonly distributed in aerosol spray cans for portability. The

solvent tends to pull penetrant from the indications by solvent action. Since the solvent is highly volatile, forced drying is not required

2.5 Steps of Liquid Penetrant Testing

The exact procedure for liquid penetrant testing can vary from case to case depending on several factors such as the penetrant system being used, the size and material of the component being inspected, the type of discontinuities being expected in the component and the condition and environment under which the inspection is performed. However, the general steps can be summarized as follows:

1. *Surface Preparation*: One of the most critical steps of a liquid penetrant testing is the surface preparation. The surface must be free of oil, grease, water, or other contaminants that may prevent penetrant from entering flaws. The sample may also require etching if mechanical operations such as machining, sanding, or grit blasting have been performed. These and other mechanical operations can smear metal over the flaw opening and prevent the penetrant from entering.
2. *Penetrant Application*: Once the surface has been thoroughly cleaned and dried, the penetrant material is applied by spraying, brushing, or immersing the part in a penetrant bath.
3. *Penetrant Dwell*: The penetrant is left on the surface for a sufficient time to allow as much penetrant as possible to be drawn or to seep into a defect. Penetrant dwell time is the total time that the penetrant is in contact with the part surface. Dwell times are usually recommended by the penetrant producers or required by the specification being followed. The times vary depending on the application, penetrant materials used, the material, the form of the material being inspected, and the type of discontinuity being inspected for. Minimum dwell times typically range from 5 to 60 minutes. Generally, there is no harm in using a longer penetrant dwell time as long as the penetrant is not allowed to dry. The ideal dwell time is often determined by experimentation and may be very specific to a particular application.
4. *Excess Penetrant Removal*: This is the most delicate step of the inspection procedure because the excess penetrant must be removed from the surface of the sample while removing as little penetrant as possible from defects. Depending on the penetrant system used, this step may involve cleaning with a solvent, direct rinsing with water, or first treating the part with an emulsifier and then rinsing with water.
5. *Developer Application*: A thin layer of developer is then applied to the sample to draw penetrant trapped in flaws back to the surface where it will be visible. Developers come in a variety of forms that may be applied by dusting (*dry powders*), dipping, or spraying (*wet developers*).
6. *Indication Development*: The developer is allowed to stand on the part surface for a period of time sufficient to permit the extraction of the trapped penetrant out of any surface flaws. This development time is usually a minimum of 10 minutes. Significantly longer times may be necessary for tight cracks.
7. *Inspection*: Inspection is then performed under appropriate lighting to detect indications from any flaws which may be present.
8. *Clean Surface*: The final step in the process is to thoroughly clean the part surface to remove the developer from the parts that were found to be acceptable.

Advantages

- High sensitivity (*small discontinuities can be detected*).
- Few material limitations (*metallic and nonmetallic, magnetic and nonmagnetic, and conductive and nonconductive materials may be inspected*).
- . Rapid inspection of large areas and volumes.
- Suitable for parts with complex shapes.
- Indications are produced directly on the surface of the part and constitute a visual representation of the flaw.
- Portable (materials are available in aerosol spray cans)

- Low cost (materials and associated equipment are relatively inexpensive)

Disadvantages

- Only surface breaking defects can be detected.
- Only materials with a relatively nonporous surface can be inspected.
- Pre-cleaning is critical since contaminants can mask defects.
- Metal smearing from machining, grinding, and grit or vapor blasting must be removed.
- The inspector must have direct access to the surface being inspected.
- Surface finish and roughness can affect inspection sensitivity.
- Multiple process operations must be performed and controlled.
- Post cleaning of acceptable parts or materials is required.
- Chemical handling and proper disposal is required.

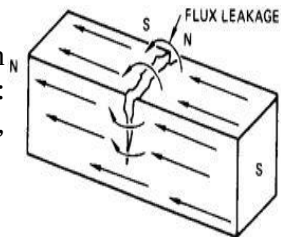
2.6 Magnetic Particle Testing

Magnetic particle testing is one of the most widely utilized NDT methods since it is fast and relatively easy to apply and part surface preparation is not as critical as it is for some other methods. This method uses magnetic fields and small magnetic particles (*i.e. iron filings*) to detect flaws in components. The only requirement from an inspectability standpoint is that the component being inspected must be made of a ferromagnetic material (*a materials that can be magnetized*) such as iron, nickel, cobalt, or some of their alloys. The method is used to inspect a variety of product forms including castings, forgings, and weldments. Many different industries use magnetic particle inspection such as structural steel, automotive, petrochemical, power generation, and aerospace industries. Underwater inspection is another area where magnetic particle inspection may be used to test items such as offshore structures and underwater pipelines.

2.6.1 Basic Principles

In theory, magnetic particle testing has a relatively simple concept. It can be considered as a combination of two nondestructive testing methods: magnetic flux leakage testing and visual testing. For the case of a bar magnet, the magnetic field is in and around the magnet.

Any place that a magnetic line of force exits or enters the magnet is called a “pole” (*magnetic lines of force exit the magnet from north pole and enter from the south pole*). When a bar magnet is broken in the center of its length, two complete bar magnets with magnetic poles on each end of each piece will result. If the magnet is just cracked but not broken completely in two, a north and south pole will form at each edge of the crack. The magnetic field exits the north pole and reenters at the south pole. The magnetic field spreads out when it encounters the small air gap created by the crack because the air cannot support as much magnetic field per unit volume as the magnet can. When the field spreads out, it appears to leak out of the material and, thus is called a flux leakage field.



If iron particles are sprinkled on a cracked magnet, the particles will be attracted to and cluster not only at the poles at the ends of the magnet, but also at the poles at the edges of the crack. This cluster of particles is much easier to see than the actual crack and this is the basis for magnetic particle inspection. The first step in a magnetic particle testing is to magnetize the component that is to be inspected. If any defects on or near the surface are present, the defects will create a leakage field. After the component has been magnetized, iron particles, either in a dry or wet suspended form, are applied to the surface of the magnetized part. The particles will be attracted and cluster at the flux leakage fields, thus forming a visible indication that the inspector can detect.

Advantages

- High sensitivity (*small discontinuities can be detected*).

- Indications are produced directly on the surface of the part and constitute a visual representation of flaw.
- Minimal surface preparation (*no need for paint removal*)
- Portable (*small portable equipment & materials available in spray cans*)
- Low cost (*materials and associated equipment are relatively inexpensive*)

Disadvantages

- Only surface and near surface defects can be detected.
- Only applicable to ferromagnetic materials.
- Relatively small area can be inspected at a time.
- Only materials with a relatively nonporous surface can be inspected.

2.7 Magnetism

The concept of magnetism centers around the magnetic field and what is known as a dipole. The term "*magnetic field*" simply describes a volume of space where there is a change in energy within that volume. The location where a magnetic field exits or enters a material is called a magnetic pole. Magnetic poles have never been detected in isolation but always occur in pairs, hence the name dipole. Therefore, a dipole is an object that has a magnetic pole on one end and a second, equal but opposite, magnetic pole on the other. A bar magnet is a dipole with a north pole at one end and south pole at the other. The source of magnetism lies in the basic building block of all matter, the atom. Atoms are composed of protons, neutrons and electrons. The protons and neutrons are located in the atom's nucleus and the electrons are in constant motion around the nucleus. Electrons carry a negative electrical charge and produce a magnetic field as they move through space. A magnetic field is produced whenever an electrical charge is in motion. The strength of this field is called the magnetic moment. When an electric current flows through a conductor, the movement of electrons through the conductor causes a magnetic field to form around the conductor. The magnetic field can be detected using a compass. Since all matter is comprised of atoms, all materials are affected in some way by a magnetic field; however, materials do not react the same way to the magnetic field.

2.8 Reaction of Materials to Magnetic Field

When a material is placed within a magnetic field, the magnetic forces of the material's electrons will be affected. This effect is known as Faraday's Law of Magnetic Induction. However, materials can react quite differently to the presence of an external magnetic field. The magnetic moments associated with atoms have three origins: the electron motion, the change in motion caused by an external magnetic field, and the spin of the electrons. In most atoms, electrons occur in pairs where these pairs spin in opposite directions. The opposite spin directions of electron pairs cause their magnetic fields to cancel each other. Therefore, no net magnetic field exists. Alternately, materials with some unpaired electrons will have a net magnetic field and will react more to an external field.

According to their interaction with a magnetic field, materials can be classified as:

Diamagnetic materials which have a weak, negative susceptibility to magnetic fields. Diamagnetic materials are slightly repelled by a magnetic field and the material does not retain the magnetic properties when the external field is removed. In diamagnetic materials all the electrons are paired so there is no permanent net magnetic moment per atom. Most elements in the periodic table, including copper, silver, and gold, are diamagnetic.

Paramagnetic materials which have a small, positive susceptibility to magnetic fields. These materials are slightly attracted by a magnetic field and the material does not retain the magnetic properties when the external field is removed. Paramagnetic materials have some unpaired electrons. Examples of paramagnetic materials include magnesium, molybdenum, and lithium.

Ferromagnetic materials have a large, positive susceptibility to an external magnetic field. They exhibit a strong attraction to magnetic fields and are able to retain their magnetic properties after the external field has been removed. Ferromagnetic materials have some unpaired electrons so their atoms have a net magnetic moment. They get their strong magnetic properties due to the presence of magnetic domains. In

these domains, large numbers of atom's moments are aligned parallel so that the magnetic force within the domain is strong (*this happens during the solidification of the material where the atom moments are aligned within each crystal "i.e., grain" causing a strong magnetic force in one direction*). When a ferromagnetic material is in the unmagnetized state, the domains are nearly randomly organized (*since the crystals are in arbitrary directions*) and the net magnetic field for the part as a whole is zero. When a magnetizing force is applied, the domains become aligned to produce a strong magnetic field within the part. Iron, nickel, and cobalt are examples of ferromagnetic materials. Components made of these materials are commonly inspected using the magnetic particle method.

2.9 Magnetic Field Characteristics

2.9.1 Magnetic Field In and Around a Bar Magnet

The magnetic field surrounding a bar magnet can be seen in the magnetograph below. A magnetograph can be created by placing a piece of paper over a magnet and sprinkling the paper with iron filings. The particles align themselves with the lines of magnetic force produced by the magnet. It can be seen in the magnetograph that there are poles all along the length of the magnet but that the poles are concentrated at the ends of the magnet (*the north and south poles*).

2.9.2 Magnetic Fields in and around Horseshoe and Ring Magnets

Magnets come in a variety of shapes and one of the more common is the horseshoe (U) magnet. The horseshoe magnet has north and south poles just like a bar magnet but the magnet is curved so the poles lie in the same plane. The magnetic lines of force flow from pole to pole just like in the bar magnet. However, since the poles are located closer together and a more direct path exists for the lines of flux to travel between the poles, the magnetic field is concentrated between the poles.

General Properties of Magnetic Lines of Force

Magnetic lines of force have a number of important properties, which include:

- ☐ They seek the path of least resistance between opposite magnetic poles (*in a single bar magnet shown, they attempt to form closed loops from pole to pole*).
- ☐ They never cross one another.
- ☐ They all have the same strength.
- ☐ Their density decreases with increasing distance from the poles.
- ☐ Their density decreases (*they spread out*) when they move from an area of higher permeability to an area of lower permeability.
- ☐ They are considered to have direction as if flowing, though no actual movement occurs.
- ☐ They flow from the south pole to the north pole within a material and north pole to south pole in air.

Electromagnetic Fields

Magnets are not the only source of magnetic fields. The flow of electric current through a conductor generates a magnetic field. When electric current flows in a long straight wire, a circular magnetic field is generated around the wire and the intensity of this magnetic field is directly proportional to the amount of current carried by the wire. The strength of the field is strongest next to the wire and diminishes with distance. In most conductors, the magnetic field exists only as long as the current is flowing. However, in ferromagnetic materials the electric current will cause some or all of the magnetic domains to align and a residual magnetic field will remain. Also, the direction of the magnetic field is dependent on the direction of the electrical current in the wire. The direction of the magnetic field around a conductor can be determined using a simple rule called the "*right-hand clasp rule*". If a person grasps a conductor in one's right hand with the thumb pointing in the direction of the current, the fingers will circle the conductor in the direction of the magnetic field.

Magnetic Field Produced by a Coil

When a current carrying wire is formed into several loops to form a coil, the magnetic field circling each loop combines with the fields from the other loops to produce a concentrated field through the center of the coil (*the field flows along the longitudinal axis and circles back around the outside of the coil*).

When the coil loops are tightly wound, a uniform magnetic field is developed throughout the length of the coil. The strength of the magnetic field increases not only with increasing current but also with each loop that is added to the coil. A long, straight coil of wire is called a *solenoid* and it can be used to generate a nearly uniform magnetic field similar to that of a bar magnet. The concentrated magnetic field inside a coil is very useful in magnetizing ferromagnetic materials for inspection using the magnetic particle testing method.

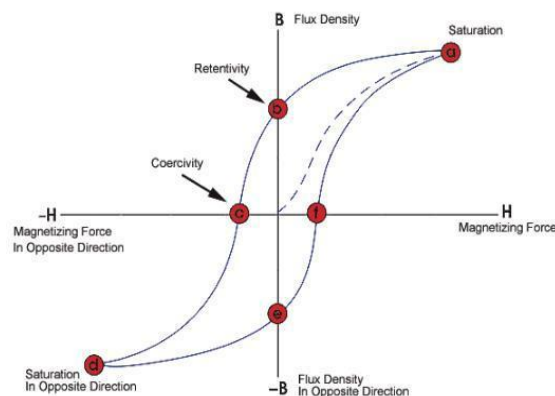
Quantifying Magnetic Properties

The various characteristics of magnetism can be measured and expressed quantitatively. Different systems of units can be used for quantifying magnetic properties. SI units will be used in this material. The advantage of using SI units is that they are traceable back to an agreed set of four base units; meter, kilogram, second, and Ampere.

- The unit for magnetic field strength **H** is *ampere/meter (A/m)*. A magnetic field strength of *1 A/m* is produced at the center of a single circular conductor with a *1 meter* diameter carrying a steady current of *1 ampere*.
- The number of magnetic lines of force cutting through a plane of a given area at a right angle is known as the magnetic *flux density*, **B**. The flux density or magnetic induction has the *Tesla* as its unit. One *Tesla* is equal to *1 Newton/(A/m)*. From these units, it can be seen that the flux density is a measure of the force applied to a particle by the magnetic field.
- The total number of lines of magnetic force in a material is called magnetic *flux*, **Φ**. The strength of the flux is determined by the number of magnetic domains that are aligned within a material. The *total flux* is simply the flux density applied over an area. Flux carries the unit of a *weber*, which is simply a *Tesla-meter²*.
- The *magnetization* **M** is a measure of the extent to which an object is magnetized. It is a measure of the magnetic dipole moment per unit volume of the object. Magnetization carries the same units as a magnetic field *A/m*.

2.10 The Hysteresis Loop and Magnetic Properties

A great deal of information can be learned about the magnetic properties of a material by studying its hysteresis loop. A hysteresis loop shows the relationship between the induced magnetic flux density (**B**) and the magnetizing force (**H**). It is often referred to as the *B-H* loop. An example hysteresis loop is shown below.



The loop is generated by measuring the magnetic flux of a ferromagnetic material while the magnetizing force is changed. A ferromagnetic material that has never been previously magnetized or has been thoroughly demagnetized will follow the dashed line as $\mathbf{H}$ is increased. As the line demonstrates, the greater the amount of current applied ($\mathbf{H}+$), the stronger the magnetic field in the component ($\mathbf{B}+$). At point "a" almost all of the magnetic domains are aligned and an additional increase in the magnetizing force will produce very little increase in magnetic flux. The material has reached the point of magnetic saturation. When $\mathbf{H}$ is reduced to zero, the curve will move from point "a" to point "b". At this point, it can be seen that some magnetic flux remains in the material even though the magnetizing force is zero. This is referred to as the point of retentivity on the graph and indicates the level of residual magnetism in the material (*Some of the magnetic domains remain aligned but some have lost their alignment*). As the

magnetizing force is reversed, the curve moves to point "c", where the flux has been reduced to zero. This is called the point of coercivity on the curve (*the reversed magnetizing force has flipped enough of the domains so that the net flux within the material is zero*). The force required to remove the residual magnetism from the material is called the coercive force or coercivity of the material. As the magnetizing force is increased in the negative direction, the material will again become magnetically saturated but in the opposite direction, point "d". Reducing **H** to zero brings the curve to point "e". It will have a level of residual magnetism equal to that achieved in the other direction. Increasing **H** back in the positive direction will return **B** to zero. Notice that the curve did not return to the origin of the graph because some force is required to remove the residual magnetism. The curve will take a different path from point "f" back to the saturation point where it will complete the loop.

From the hysteresis loop, a number of primary magnetic properties of a material can be determined:

1. **Retentivity** - A measure of the residual flux density corresponding to the saturation induction of a magnetic material. In other words, it is a material's ability to retain a certain amount of residual magnetic field when the magnetizing force is removed after achieving saturation (The value of **B** at point **b** on the hysteresis curve).
2. **Residual Magnetism or Residual Flux** - The magnetic flux density that remains in a material when the magnetizing force is zero. Note that residual magnetism and retentivity are the same when the material has been magnetized to the saturation point. However, the level of residual magnetism may be lower than the retentivity value when the magnetizing force did not reach the saturation level.
3. **Coercive Force** - The amount of reverse magnetic field which must be applied to a magnetic material to make the magnetic flux return to zero (The value of **H** at point **c** on the hysteresis curve).
4. **Permeability, μ** - A property of a material that describes the ease with which a magnetic flux is established in the material.
5. **Reluctance** - Is the opposition that a ferromagnetic material shows to the establishment of a magnetic field. Reluctance is analogous to the resistance in an electrical circuit.

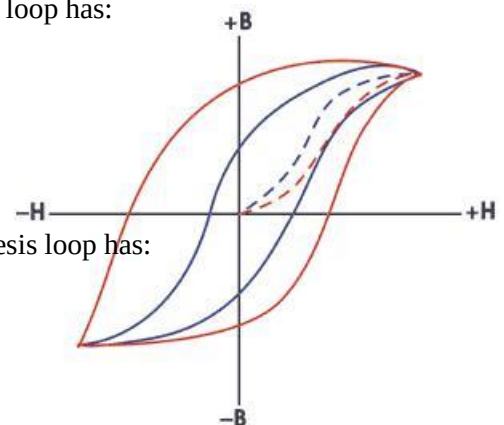
Permeability

As previously mentioned, permeability (μ) is a material property that describes the ease with which a magnetic flux is established in a component. It is the ratio of the flux density (**B**) created within a material to the magnetizing field (**H**) and it is represented by the following equation:

$$\mu = B/H$$

This equation describes the slope of the curve at any point on the hysteresis loop. The permeability value given in literature for materials is usually the maximum permeability or the maximum relative permeability. The maximum permeability is the point where the slope of the B/H curve for the unmagnetized material is the greatest. This point is often taken as the point where a straight line from the origin is tangent to the B/H curve. The shape of the hysteresis loop tells a great deal about the material being magnetized. The hysteresis curves of two different materials are shown in the graph.

- ☐ Relative to other materials, a material with a wider hysteresis loop has:
 - Lower Permeability
 - Higher Retentivity
 - Higher Coercivity
 - Higher Reluctance
 - Higher Residual Magnetism
- ☐ Relative to other materials, a material with a narrower hysteresis loop has:
 - Higher Permeability
 - Lower Retentivity
 - Lower Coercivity
 - Lower Reluctance
 - Lower Residual Magnetism



In magnetic particle testing, the level of residual magnetism is important. Residual magnetic fields are affected by the permeability, which can be related to the carbon content and alloying of the material. A

component with high carbon content will have low permeability and will retain more magnetic flux than a material with low carbon content.

2.11 Magnetic Field Orientation and Flaw Detectability

To properly inspect a component for cracks or other defects, it is important to understand that the orientation of the crack relative to the magnetic lines of force determines if the crack can or cannot be detected. There are two general types of magnetic fields that can be established within a component.

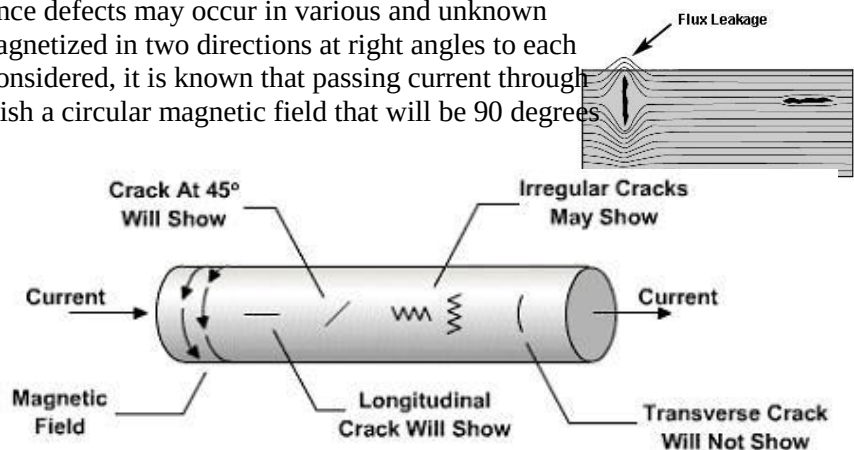
□ A *longitudinal magnetic field* has magnetic lines of force that run parallel to the long axis of the part. Longitudinal magnetization of a component can be accomplished using the longitudinal field set up by a coil or solenoid. It can also be accomplished using permanent magnets or electromagnets.

□ A *circular magnetic field* has magnetic lines of force that run circumferentially around the perimeter of a part. A circular magnetic field is induced in an article by either passing current through the component or by passing current through a conductor surrounded by the component.

The type of magnetic field established is determined by the method used to magnetize the specimen. Being able to magnetize the part in two directions is important because the best detection of defects occurs when the lines of magnetic force are established at right angles to the longest dimension of the defect. This orientation creates the largest disruption of the magnetic field within the part and the greatest flux leakage at the surface of the part. If the magnetic field is parallel to the defect, the field will see little disruption and no flux leakage field will be produced. An orientation of 45 to 90 degrees between the magnetic field and the defect is

necessary to form an indication. Since defects may occur in various and unknown directions, each part is normally magnetized in two directions at right angles to each other. If the component shown is considered, it is known that passing current through the part from end to end will establish a circular magnetic field that will be 90 degrees to the direction of the current.

Therefore, defects that have significant dimension in the direction of the current (*longitudinal defects*) should be detectable, while transverse-type defects will not be detectable with circular magnetization.



2.11.1 Magnetization of Ferromagnetic Materials

There are a variety of methods that can be used to establish a magnetic field in a component for evaluation using magnetic particle inspection. It is common to classify the magnetizing methods as either direct or indirect.

Magnetization Using Direct Induction (Direct Magnetization)

With direct magnetization, current is passed directly through the component. The flow of current causes a circular magnetic field to form in and around the conductor. When using the direct magnetization method, care must be taken to ensure that good electrical contact is established and maintained between the test equipment and the test component to avoid damage of the the component (*due to arcing or overheating at high resistance points*).

There are several ways that direct magnetization is commonly accomplished.

One way involves *clamping the component between two electrical contacts* in a special piece of equipment. Current is passed through the component and a circular magnetic field is established in and around the component. When the magnetizing current is stopped, a residual magnetic field will remain within the component. The strength of the induced magnetic field is proportional to the amount of current passed through the component

-A second technique involves using *clamps or prods*, which are attached or placed in contact with the component. Electrical current flows through the component from contact to contact. The current sets up a circular magnetic field around the path of the current.

Magnetization Using Indirect Induction (Indirect Magnetization)

Indirect magnetization is accomplished by using a strong external magnetic field to establish a magnetic field within the component. As with direct magnetization, there are several ways that indirect magnetization can be accomplished. The use of *permanent magnets* is a low cost method of establishing a magnetic field. However, their use is limited due to lack of control of the field strength and the difficulty of placing and removing strong permanent magnets from the component.

Electromagnets in the form of an adjustable horseshoe magnet (*called a yoke*) eliminate the problems associated with permanent magnets and are used extensively in industry. Electromagnets only exhibit a magnetic flux when electric current is flowing around the soft iron core. When the magnet is placed on the component, a magnetic field is established between the north and south poles of the magnet. Another way of indirectly inducing a magnetic field in a material is by using the magnetic field of a current carrying conductor. A circular magnetic field can be established in cylindrical components by using a *central conductor*. Typically, one or more cylindrical components are hung from a solid copper bar running through the inside diameter. Current is passed through the copper bar and the resulting circular magnetic field establishes a magnetic field within the test components. The use of *coils and solenoids* is a third method of indirect magnetization. When the length of a component is several times larger than its diameter, a longitudinal magnetic field can be established in the component. The component is placed longitudinally in the concentrated magnetic field that fills the center of a coil or solenoid. This magnetization technique is often referred to as a "*coil shot*".

Types of Magnetizing Current

As mentioned previously, electric current is often used to establish the magnetic field in components during magnetic particle inspection. Alternating current (AC) and direct current (DC) are the two basic types of current commonly used. The type of current used can have an effect on the inspection results, so the types of currents commonly used are briefly discussed here.

Direct Current

Direct current (DC) flows continuously in one direction at a constant voltage. A battery is the most common source of direct current. The current is said to flow from the positive to the negative terminal, though electrons flow in the opposite direction. DC is very desirable when inspecting for subsurface defects because DC generates a magnetic field that penetrates deeper into the material. In ferromagnetic materials, the magnetic field produced by DC generally penetrates the entire cross-section of the component.

Alternating Current

Alternating current (AC) reverses its direction at a rate of 50 or 60 cycles per second. Since AC is readily available in most facilities, it is convenient to make use of it for magnetic particle inspection. However, when AC is used to induce a magnetic field in ferromagnetic materials, the magnetic field will be limited to a thin layer at the surface of the component. This phenomenon is known as the "*skin effect*" and it occurs because the changing magnetic field generates eddy currents in the test object. The eddy currents produce a magnetic field that opposes the primary field, thus reducing the net magnetic flux below the surface. Therefore, it is recommended that AC be used only when the inspection is limited to surface defects.

Rectified Alternating Current

Clearly, the skin effect limits the use of AC since many inspection applications call for the detection of subsurface defects. Luckily, AC can be converted to current that is very much like DC through the process of rectification. With the use of rectifiers, the reversing AC can be converted to a one directional current. The three commonly used types of rectified current are described below.

Half Wave Rectified Alternating Current (HWAC)

When single phase alternating current is passed through a rectifier, current is allowed to flow in only one direction. The reverse half of each cycle is blocked out so that a one directional, pulsating current is produced. The current rises from zero to a maximum and then returns to zero. No current flows during the time when the reverse cycle is blocked out. The HWAC repeats at same rate as the unrectified current (50 or 60 Hz). Since half of the current is blocked out, the amperage is half of the unaltered AC. This type of current is often referred to as *half wave DC or pulsating DC*. The pulsation of the HWAC helps in forming magnetic particle indications by vibrating the particles and giving them added mobility where that is especially important when using dry particles. HWAC is most often used to power electromagnetic yokes.

Full Wave Rectified Alternating Current (FWAC) (Single Phase)

Full wave rectification inverts the negative current to positive current rather than blocking it out. This produces a pulsating DC with no interval between the pulses. Filtering is usually performed to soften the sharp polarity switching in the rectified current. While particle mobility is not as good as half-wave AC due to the reduction in pulsation, the depth of the subsurface magnetic field is improved.

Three Phase Full Wave Rectified Alternating Current

Three phase current is often used to power industrial equipment because it has more favorable power transmission and line loading characteristics. This type of electrical current is also highly desirable for magnetic particle testing because when it is rectified and filtered, the resulting current very closely resembles direct current. Stationary magnetic particle equipment wired with three phase AC will usually have the ability to magnetize with AC or DC (*three phase full wave rectified*), providing the inspector with the advantages of each current form

Magnetic Fields Distribution and Intensity

Longitudinal Fields

When a long component is magnetized using a solenoid having a shorter length, only the material within the solenoid and about the same length on each side of the solenoid will be strongly magnetized. This occurs because the magnetizing force diminishes with increasing distance from the solenoid. Therefore, a long component must be magnetized and inspected at several locations along its length for complete inspection coverage.

Circular Fields

When a circular magnetic field forms in and around a conductor due to the passage of electric current through it, the following can be said about the distribution and intensity of the magnetic field:

- The field strength varies from zero at the center of the component to a maximum at the surface.
- The field strength at the surface of the conductor decreases as the radius of the conductor increases (*when the current strength is held constant*).
- the field strength inside the conductor is dependent on the current strength, magnetic permeability of the material, and, if ferromagnetic, the location on the B-H curve.
- The field strength outside the conductor is directly proportional to the current strength and it decreases with distance from the conductor.

The images below show the magnetic field strength graphed versus distance from the center of the conductor when current passes through a solid circular conductor.

- ☐ In a nonmagnetic conductor carrying DC, the internal field strength rises from zero at the center to a maximum value at the surface of the conductor.
- ☐ In a magnetic conductor carrying DC, the field strength within the conductor is much greater than it is in the nonmagnetic conductor. This is due to the permeability of the magnetic material. The external field is exactly the same for the two materials provided the current level and conductor radius are the same
- ☐ When the magnetic conductor is carrying AC, the internal magnetic field will be concentrated in a thin layer near the surface of the conductor (*skin effect*). The external field decreases with increasing distance from the surface same as with DC.

a hollow circular conductor there is no magnetic field in the void area. The magnetic field is zero at the inner surface and rises until it reaches a maximum at the outer surface.

- ☐ Same as with a solid conductor, when DC current is passed through a magnetic conductor, the field strength within the conductor is much greater than in nonmagnetic conductor due to the permeability of the magnetic material. The external field strength decreases with distance from the surface of the conductor.

The external field is exactly the same for the two materials provided the current level and conductor radius are the same.

□ When AC current is passed through a hollow circular magnetic conductor, the skin effect concentrates the magnetic field at the outside diameter of the component. As can be seen from these three field distribution images, the field strength at the inside surface of hollow conductor is very low when a circular magnetic field is established by direct magnetization. Therefore, the direct method of magnetization is not recommended when inspecting the inside diameter wall of a hollow component for shallow defects (*if the defect has significant depth, it may be detectable using DC since the field strength increases rapidly as one moves from the inner towards the outer surface*).

□ A much better method of magnetizing hollow components for inspection of the ID and OD surfaces is with the use of a central conductor. As can be seen in the field distribution image, when current is passed through a nonmagnetic central conductor (copper bar), the magnetic field produced on the inside diameter surface of a magnetic tube is much greater and the field is still strong enough for defect detection on the OD surface.

Demagnetization

After conducting a magnetic particle inspection, it is usually necessary to demagnetize the component. Remanent magnetic fields can:

- affect machining by causing cuttings to cling to a component.
- interfere with electronic equipment such as a compass.
- create a condition known as "arc blow" in the welding process. Arc blow may cause the weld arc to wander or filler metal to be repelled from the weld
- cause abrasive particles to cling to bearing or faying surfaces and increase wear.

Removal of a field may be accomplished in several ways. The most effective way to demagnetize a material is by heating the material above its curie temperature (*for instance, the curie temperature for a low carbon steel is 770°C*). When steel is heated above its curie temperature then it is cooled back down, the the orientation of the magnetic domains of the individual grains will become randomized again and thus the component will contain no residual magnetic field. The material should also be placed with its long axis in an east-west orientation to avoid any influence of the Earth's magnetic field. However, it is often inconvenient to heat a material above its curie temperature to demagnetize it, so another method that returns the material to a nearly unmagnetized state is commonly used. Subjecting the component to a reversing and decreasing magnetic field will return the dipoles to a nearly random orientation throughout the material. This can be accomplished by pulling a component out and away from a coil with AC passing through it. With AC Yokes, demagnetization of local areas may be accomplished by placing the yoke contacts on the surface, moving them in circular patterns around the area, and slowly withdrawing the yoke while the current is applied. Also, many stationary magnetic particle inspection units come with a demagnetization feature that slowly reduces the AC in a coil in which the component is placed.

A field meter is often used to verify that the residual flux has been removed from a component. Industry standards usually require that the magnetic flux be reduced to less than 3 Gauss (3×10^{-4} Tesla) after completing a magnetic particle inspection.

Measuring Magnetic Fields

When performing a magnetic particle inspection, it is very important to be able to determine the direction and intensity of the magnetic field. The field intensity must be high enough to cause an indication to form, but not too high to cause nonrelevant indications to mask relevant indications. Also, after magnetic inspection it is often needed to measure the level of residual magnetism. Since it is impractical to measure the actual field strength within the material, all the devices measure the magnetic field that is outside of the material. The two devices commonly used for quantitative measurement of magnetic fields in magnetic particle inspection are the field indicator and the Hall-effect meter, which is also called a gauss meter.

Field Indicators

Field indicators are small mechanical devices that utilize a soft iron vane that is deflected by a magnetic field. The vane is attached to a needle that rotates and moves the pointer on the scale. Field indicators can be adjusted and calibrated so that quantitative information can be obtained. However, the measurement range of field indicators is usually small due to the mechanics of the device (*the one shown in the image has a range from plus 20 to minus 20 Gauss*). This limited range makes them best suited for measuring the residual magnetic field after demagnetization.

Hall-Effect (Gauss/Tesla) Meter

A Hall-effect meter is an electronic device that provides a digital readout of the magnetic field strength in Gauss or Tesla units. The meter uses a very small conductor or semiconductor element at the tip of the probe. Electric current is passed through the conductor. In a magnetic field, a force is exerted on the moving electrons which tends to push them to one side of the conductor. A buildup of charge at the sides of the conductors will balance this magnetic influence, producing a measurable voltage between the two sides of the conductor. The probe is placed in the magnetic field such that the magnetic lines of force intersect the major dimensions of the sensing element at a right angle.

Magnetization Equipment for Magnetic Particle Testing

To properly inspect a part for cracks or other defects, it is important to become familiar with the different types of magnetic fields and the equipment used to generate them. As discussed previously, one of the primary requirements for detecting a defect in a ferromagnetic material is that the magnetic field induced in the part must intercept the defect at a 45 to 90 degree angle. Flaws that are normal (90 degrees) to the magnetic field will produce the strongest indications because they disrupt more of the magnet flux. Therefore, for proper inspection of a component, it is important to be able to establish a magnetic field in at least two directions. A variety of equipment exists to establish the magnetic field for magnetic particle testing. One way to classify equipment is based on its portability. Some equipment is designed to be portable so that inspections can be made in the field and some is designed to be stationary for ease of inspection in the laboratory or manufacturing facility.

Portable Equipment

Permanent Magnets

Permanent magnets can be used for magnetic particle inspection as the source of magnetism (*bar magnets or horseshoe magnets*). The use of industrial magnets is not popular because they are very strong (*they require significant strength to remove them from the surface, about 250 N for some magnets*) and thus they are difficult and sometimes dangerous to handle. However, permanent magnets are sometimes used by divers for inspection in underwater environments or other areas, such as explosive environments, where electromagnets cannot be used. Permanent magnets can also be made small enough to fit into tight areas where electromagnets might not fit.

Electromagnetic Yokes

An electromagnetic yoke is a very common piece of equipment that is used to establish a magnetic field. A switch is included in the electrical circuit so that the current and, therefore, the magnetic field can be turned on

and off. They can be powered with AC from a wall socket or by DC from a battery pack. This type of magnet generates a very strong magnetic field in a local area where the poles of the magnet touch the part being inspected. Some yokes can lift weights in excess of 40 pounds.

Pros

Prods are handheld electrodes that are pressed against the surface of the component being inspected to make contact for passing electrical current (*AC or DC*) through the metal. Prods are typically made from copper and have an insulated handle to help protect the operator. One of the prods has a trigger switch so that the current can be quickly and easily turned on and off. Sometimes the two prods are connected by any insulator, as shown in the image, to facilitate one hand operation. This is referred to as a dual prod and is commonly used for weld inspections.

However, caution is required when using prods because electrical arcing can occur and cause damage to the component if proper contact is not maintained between the prods and the component surface. For this reason, the use of prods is not allowed when inspecting aerospace and other critical components. To help prevent arcing, the

prod tips should be inspected frequently to ensure that they are not oxidized, covered with scale or other contaminant, or damaged.

Portable Coils and Conductive Cables

Coils and conductive cables are used to establish a longitudinal magnetic field within a component. When a preformed coil is used, the component is placed against the inside surface on the coil. Coils typically have three or five turns of a copper cable within the molded frame. A foot switch is often used to energize the coil. Also, flexible conductive cables can be wrapped around a component to form a coil. The number of wraps is determined by the magnetizing force needed and of course, the length of the cable. Normally, the wraps are kept as close together as possible. When using a coil or cable wrapped into a coil, amperage is usually expressed in ampere-turns. Ampere-turns is the amperage shown on the amp meter times the number of turns in the coil.

Portable Power Supplies

Portable power supplies are used to provide the necessary electricity to the prods, coils or cables. Power supplies are commercially available in a variety of sizes. Small power supplies generally provide up to 1,500A of half-wave DC or AC. They are small and light enough to be carried and operate on either 120V or 240V electrical service. When more power is necessary, mobile power supplies can be used. These units come with wheels so that they can be rolled where needed. These units also operate on 120V or 240V electrical service and can provide up to 6,000A of AC or half-wave DC.

Stationery Equipment

Stationary stationary system is the wet horizontal (bench) unit. Wet horizontal units are designed to allow for batch inspections of a variety of components. The units have head and tail stocks (*similar to a lathe*) with electrical contact that the part can be clamped between. A circular magnetic field is produced with direct magnetization.

Most units also have a movable coil that can be moved into place so the indirect magnetization can be used to

produce a longitudinal magnetic field. Most coils have five turns and can be obtained in a variety of sizes. The

wet magnetic particle solution is collected and held in a tank. A pump and hose system is used to apply the particle solution to the components being inspected. Some of the systems offer a variety of options in electrical current used for magnetizing the component (*AC, half wave DC, or full wave DC*). In some units, a

demagnetization feature is built in, which uses the coil and decaying AC. magnetic particle inspection equipment is designed for use in laboratory or production environment. The most common

Magnetic Field Indicators

Determining whether a magnetic field is of adequate strength and in the proper direction is critical when performing magnetic particle testing. There is actually no easy-to-apply method that permits an exact measurement of field intensity at a given point within a material. Cutting a small slot or hole into the material and measuring the leakage field that crosses the air gap with a Hall-effect meter is probably the best way to get an estimate of the actual field strength within a part. However, since that is not practical, there are a number of tools and methods that are used to determine the presence and direction of the field surrounding a component.

Hall-Effect Meter (Gauss Meter)

As discussed earlier, a Gauss meter is commonly used to measure the tangential field strength on the surface of the part. By placing the probe next to the surface, the meter measures the intensity of the field in the air adjacent to the component when a magnetic field is applied. The advantages of this device are: it provides a quantitative measure of the strength of magnetizing force tangential to the surface of a test piece, it can be used for measurement of residual magnetic fields, and it can be used repetitively. The main disadvantage is that such devices must be periodically calibrated.

Quantitative Quality Indicator (QQI)

The Quantitative Quality Indicator (QQI) or *Artificial Flaw Standard* is often the preferred method of assuring proper field direction and adequate field strength (*it is used with the wet method only*). The QQI is a thin strip (*0.05 or 0.1 mm thick*) of AISI 1005 steel with a specific pattern, such as concentric circles or a plus sign, etched on it. The QQI is placed directly on the surface, with the etched side facing the surface, and it is usually fixed to the surface using a tape then the component is then magnetized and particles applied. When the field strength is adequate, the particles will adhere over the engraved pattern and provide information about the field direction.

Pie Gage

The pie gage is a disk of highly permeable material divided into four, six, or eight sections by non-ferromagnetic material (*such as copper*). The divisions serve as artificial defects that radiate out in different directions from the center. The sections are furnace brazed and copper plated. The gage is placed on the test piece copper side up and the test piece is magnetized. After particles are applied and the excess removed, the indications provide the inspector the orientation of the magnetic field. Pie gages are mainly used on flat surfaces such as weldments or steel castings where dry powder is used with a yoke or prods. The pie gage is not recommended for precision parts with complex shapes, for wet-method applications, or for proving field magnitude. The gage should be demagnetized between readings.

Slotted Strips

Slotted strips are pieces of highly permeable ferromagnetic material with slots of different widths. These strips can be used with the wet or dry method. They are placed on the test object as it is inspected. The indications produced on the strips give the inspector a general idea of the field strength in a particular area.

Magnetic Particles

As mentioned previously, the particles that are used for magnetic particle inspection are a key ingredient as they form the indications that alert the inspector to the presence of defects. Particles start out as tiny milled pieces of iron or iron oxide. A pigment (*somewhat like paint*) is bonded to their surfaces to give the particles color. The metal used for the particles has high magnetic permeability and low retentivity. High magnetic permeability is important because it makes the particles attract easily to small magnetic leakage fields from discontinuities, such as flaws. Low retentivity is important because the particles themselves never become strongly magnetized so they do not stick to each other or the surface of the part. Particles are available in a dry mix or a wet solution.

Dry Magnetic Particles

Dry magnetic particles can typically be purchased in red, black, gray, yellow and several other colors so that a high level of contrast between the particles and the part being inspected can be achieved. The size of the magnetic particles is also very important. Dry magnetic particle products are produced to include a range of particle sizes. The fine particles have a diameter of about 50 μm while the coarse particles have a diameter of 150 μm (*fine particles are more than 20 times lighter than the coarse particles*). This makes fine particles more sensitive to the leakage fields from very small discontinuities. However, dry testing particles cannot be made exclusively of the fine particles where coarser particles are needed to bridge large discontinuities and to reduce the powder's dusty nature. Additionally, small particles easily adhere to surface contamination, such as remnant dirt or moisture, and get trapped in surface roughness features. It should also be recognized that finer particles will be more easily blown away by the wind; therefore, windy conditions can reduce the sensitivity of an inspection. Also, reclaiming the dry particles is not recommended because the small particles are less likely to be recaptured and the "once used" mix will result in less sensitive inspections. The particle shape is also important. Long, slender particles tend to align themselves along the lines of magnetic force. However, if dry powder consists only of elongated particles, the application process would be less than desirable since long particles lack the ability to flow freely. Therefore, a mix of rounded and elongated particles is used since it results in a dry powder that

flows well and maintains good sensitivity. Most dry particle mixes have particles with L/D ratios between one and two.

Wet Magnetic Particles

Magnetic particles are also supplied in a wet suspension such as water or oil. The wet magnetic particle testing method is generally more sensitive than the dry because the suspension provides the particles with more mobility and makes it possible for smaller particles to be used (*the particles are typically 10 μm and smaller*) since dust and adherence to surface contamination is reduced or eliminated. The wet method also makes it easy to apply the particles uniformly to a relatively large area.

Wet method magnetic particles products differ from dry powder products in a number of ways. One way is that both visible and fluorescent particles are available. Most non-fluorescent particles are ferromagnetic iron oxides, which are either black or brown in color. Fluorescent particles are coated with pigments that fluoresce when exposed to ultraviolet light. Particles that fluoresce green-yellow are most common to take advantage of the peak color sensitivity of the eye but other fluorescent colors are also available. The carrier solutions can be water or oil-based. Water-based carriers form quicker indications, are generally less expensive, present little or no fire hazard, give off no petrochemical fumes, and are easier to clean from the part. Water-based solutions are usually formulated with a corrosion inhibitor to offer some corrosion protection. However, oil-based carrier solutions offer superior corrosion and hydrogen embrittlement protection to those materials that are prone to attack by these mechanisms. Also, both visible and fluorescent wet suspended particles are available in aerosol spray cans for increased portability and ease of application.

Dry Particle Inspection

In this magnetic particle testing technique, dry particles are dusted onto the surface of the test object as the item is magnetized. Dry particle inspection is well suited for the inspections conducted on rough surfaces. When an electromagnetic yoke is used, the AC current creates a pulsating magnetic field that provides mobility to the powder. Dry particle inspection is also used to detect shallow subsurface cracks. Dry particles with half wave DC is the best approach when inspecting for lack of root penetration in welds of thin materials.

Steps for performing dry particles inspection:

□ **Surface preparation** - The surface should be relatively clean but this is not as critical as it is with liquid penetrant inspection. The surface must be free of grease, oil or other moisture that could keep particles from moving freely. A thin layer of paint, rust or scale will reduce test sensitivity but can sometimes be left in place with adequate results. Specifications often allow up to 0.076 mm of a nonconductive coating (*such as paint*) or 0.025 mm of a ferromagnetic coating (*such as nickel*) to be left on the surface. Any loose dirt, paint, rust or scale must be removed. Some specifications require the surface to be coated with a thin layer of white paint in order to improve the contrast difference between the background and the particles (*especially when gray color particles are used*).

□ **Applying the magnetizing force** - Use permanent magnets, an electromagnetic yoke, prods, a coil or other means to establish the necessary magnetic flux.

□ **Applying dry magnetic particles** - Dust on a light layer of magnetic particles.

□ **Blowing off excess powder** - With the magnetizing force still applied, remove the excess powder from the surface with a few gentle puffs of dry air. The force of the air needs to be strong enough to remove the excess particles but not strong enough to remove particles held by a magnetic flux leakage field.

□ **Terminating the magnetizing force** - If the magnetic flux is being generated with an electromagnet or an electromagnetic field, the magnetizing force should be terminated. If permanent magnets are being used, they can be left in place.

□ **Inspection for indications** - Look for areas where the magnetic particles are clustered.

Wet Suspension Inspection

Wet suspension magnetic particle inspection, more commonly known as wet magnetic particle inspection, involves applying the particles while they are suspended in a liquid carrier. Wet magnetic particle inspection is most commonly performed using a stationary, wet, horizontal inspection unit but suspensions are also available in spray cans for use with an electromagnetic yoke. A wet inspection has several advantages over a dry inspection. First, all of the surfaces of the component can be quickly and easily covered with a relatively uniform layer of particles. Second, the liquid carrier provides mobility to the particles for an extended period of time, which allows enough particles to float to small leakage fields to form a visible indication. Therefore, wet inspection is considered best for detecting very small discontinuities on smooth surfaces. On rough surfaces, however, the particles (*which are much smaller in wet suspensions*) can settle in the surface valleys and lose mobility, rendering them less effective than dry powders under these conditions.

Steps for performing wet particle inspection:

□ **Surface preparation** - Just as is required with dry particle inspections, the surface should be relatively clean. The surface must be free of grease, oil and other moisture that could prevent the suspension from wetting the surface and preventing the particles from moving freely. A thin layer of paint, rust or scale will reduce test sensitivity, but can sometimes be left in place with adequate results. Specifications often allow up to 0.076 mm of a nonconductive coating (*such as paint*) or 0.025 mm of a ferromagnetic coating (*such as nickel*) to be left on the surface. Any loose dirt, paint, rust or scale must be removed. Some specifications require the surface to be coated with a thin layer of white paint when inspecting using visible particles in order to improve the contrast

difference between the background and the particles (*especially when gray color particles are used*).

□ **Applying suspended magnetic particles** - The suspension is gently sprayed or flowed over the surface of the part. Usually, the stream of suspension is diverted from the part just before the magnetizing field is applied.

□ **Applying the magnetizing force** - The magnetizing force should be applied immediately after applying the suspension of magnetic particles. When using a wet horizontal inspection unit, the current is applied in two or three short bursts (*1/2 second*) which helps to improve particle mobility.

□ **Inspection for indications** - Look for areas where the magnetic particles are clustered. Surface discontinuities will produce a sharp indication. The indications from subsurface flaws will be less defined and lose definition as depth increases.

THERMOGRAPHY AND EDDY CURRENT TESTING (ET)

3.1. Introduction

Thermographic inspection refers to the nondestructive testing of parts, materials or systems through the imaging of the thermal patterns at the object's surface. Strictly speaking, the term thermography alone, refers to all thermographic inspection techniques regardless of the physical phenomena used to monitor the thermal changes. For instance, the application of a temperature sensitive coating to a surface in order to measure its temperature is a thermographic inspection contact technique based on heat conduction where there is no infrared sensor involved. Infrared thermography on the other hand, is a nondestructive, nonintrusive, noncontact mapping of thermal patterns or "thermograms", on the surface of objects through the use of some kind of infrared detector.

In addition, there are two approaches in thermographic inspection:

1. Passive, in which the features of interest are naturally at a higher or lower temperature than the background, for example: the surveillance of people on a scene; and
2. Active, in which an energy source is required to produce a thermal contrast between the feature of interest and the background, for example: an aircraft part with internal flaws.

When compared with other classical nondestructive testing techniques such as ultrasonic testing or radiographic testing, thermographic inspection is safe, nonintrusive and noncontact, allowing the detection of relatively shallow subsurface defects under large surfaces) and in a fast manner. There are many other terms widely used all referring to infrared thermography, the adoption of one or other term depends on the author's background and preferences. For instance, video thermography and thermal imaging draw attention to the fact that a sequence of images is acquired and is possible to see it as a movie. Pulse-echo thermography and thermal wave imaging are adopted to emphasize the wave nature of heat. Pulsed video thermography, transient thermography, and flash thermograph are used when the specimen is stimulated using a short energy pulse.

3.2 Technique

A wide variety of energy sources can be used to induce a thermal contrast between defective and non-defective zones that can be divided in external, if the energy is delivered to the surface and then propagated through the material until it encounters a flaw; or internal, if the energy is injected into the specimen in order to stimulate exclusively the defects. Typically, external excitation is performed with optical devices such as photographic flashes or halogen lamp, whereas internal excitation can be achieved by means of mechanical oscillations, with a sonic or ultrasonic transducer^[13] for both burst and amplitude modulated stimulation. As depicted in the figure, there are three classical active thermographic techniques based on these two excitation modes.

Thermography and pulsed thermography, which are optical techniques applied externally; and vibrothermography, which uses ultrasonic waves to excite internal features. In vibrothermography, an external mechanical energy source induces a temperature difference between the defective and non-defective areas of the object. In this case, the temperature difference is the main factor that causes the emission of a broad electromagnetic spectrum of infrared radiation, which is not visible to the human eye. The locations of the defects can then be detected by infrared cameras through the process of mapping temperature distribution on the surface of the object.

3.3 Infrared vision is the capability of biological or artificial systems to detect infrared radiation. The terms thermal vision and thermal imaging, are also commonly used in this context since infrared emissions from a body are directly related to their temperature: hotter objects emit more energy in the infrared spectrum than colder ones. The human body, as well as many moving or static objects of military or civil interest, are normally warmer than the surrounding environment. Since hotter objects emit more infrared energy than colder ones, it is relatively easy to identify them with an infrared detector, day or night. Hence, the term night vision is also used (sometimes misused) in the place of "infrared vision", since one of the original purposes in developing this kind of systems was to locate enemy targets at night. However, night vision concerns the ability to see in the dark although not necessarily in the infrared spectrum. In fact, night vision equipment can be manufactured using one of two technologies light intensifiers or infrared vision. The former technology uses a photocathode to convert light to electrons, amplify the signal and transform it back to photons. Infrared vision on the other hand, uses an infrared detector working at mid or long wavelengths (invisible to the human eye) to capture the heat emitted by an object

Infrared vision is used extensively by the military for night vision, navigation, surveillance and targeting. For years, it developed slowly due to the high cost of the equipment and the low quality of available images. Since the development of the first commercial infrared cameras in the second half of the 1960s, however, the availability of new generations of infrared cameras coupled with growing computer power is providing exciting new civilian (and military) applications, to name only a few buildings and infrastructure,^[12] works of art, aerospace components^[14] and processes, maintenance defect detection and characterization, law enforcement, surveillance and public services, medical and veterinary thermal imaging. The electronic technique that uses infrared vision to "see" thermal energy, to monitor temperatures and thermal patterns is called infrared thermography.

3.4 Infrared thermography (IRT), thermal imaging, and thermal video are examples of infrared imaging science. Thermographic cameras usually detect radiation in the long-infrared range of the electromagnetic spectrum and produce images of that radiation, called thermograms. Since infrared radiation is emitted by all objects with a temperature above absolute zero according to the black body radiation law, thermography makes it possible to see one's environment with or without visible illumination. The amount of radiation emitted by an object increases with temperature; therefore, thermography allows one to see variations in temperature. When viewed through a thermal imaging camera, warm objects stand out well against cooler backgrounds; humans and other warm-blooded animals become easily visible against the environment, day or night. As a result, thermography is particularly useful to the military and other users of surveillance cameras.

3.4.1 Thermogram of a cat

Some physiological changes in human beings and other warm-blooded animals can also be monitored with thermal imaging during clinical diagnostics. Thermography is used in allergy detection and veterinary medicine. It is also used for breast screening, though primarily by alternative practitioners as it is considerably less accurate and specific than competing techniques. Government and airport personnel used thermography to detect suspected swine flu cases during the 2009 pandemic.

3.4.2 Thermal imaging camera and screen. Thermal imaging can detect elevated body temperature, one of the signs of the virus H1N1 (Swine influenza).

3.4.3 Thermography has a long history, although its use has increased dramatically with the commercial and industrial applications of the past fifty years. Firefighters use thermography to see through smoke, to find persons, and to localize the base of a fire. Maintenance technicians use thermography to locate overheating joints and sections of power lines, which are a sign of impending failure. Building construction technicians can see thermal signatures that indicate heat leaks in faulty thermal insulation and can use the results to improve the efficiency of heating and air-conditioning units. The appearance and operation of a modern thermographic camera is often similar to a camcorder. Often the live thermogram reveals temperature variations so clearly that a photograph is not necessary for analysis. A recording module is therefore not always built-in.

Non-specialized CCD and CMOS sensors have most of their spectral sensitivity in the visible light wavelength range. However, by utilizing the "trailing" area of their spectral sensitivity, namely the part of the infrared spectrum called near-infrared (NIR), and by using off-the-shelf CCTV camera it is possible under certain circumstances to obtain true thermal images of objects with temperatures at about 280 °C and higher.

Specialized thermal imaging cameras use focal plane arrays that respond to longer wavelengths. The most common types are InSb, InGaAs, HgCdTe and QWIP FPA. The newest technologies use low-cost, uncooled microbolometers as FPA sensors. Their resolution is considerably lower than that of optical cameras, mostly 160x120 or 320x240 pixels, up to 1024x768 for the most expensive models. Thermal imaging cameras are much more expensive than their visible-spectrum counterparts, and higher-end models are often export-restricted due to the military uses for this technology. Older bolometers or more sensitive models such as InSb require cryogenic cooling, usually by a miniature Stirling cycle refrigerator or liquid nitrogen.

Advantages

- It shows a visual picture so temperatures over a large area can be compared
- It is capable of catching moving targets in real time
- It is able to find deteriorating, i.e., higher temperature components prior to their failure
- It can be used to measure or observe in areas inaccessible or hazardous for other methods
- It is a non-destructive test method
- It can be used to find defects in shafts, pipes, and other metal or plastic parts
- It can be used to detect objects in dark areas
- It has some medical application, essentially in physiotherapy

Disadvantages

- Quality cameras often have a high price range
- Many models do not provide the irradiance measurements used to construct the output image; the loss of this information without a correct calibration for emissivity, distance, and ambient temperature and relative humidity entails that the resultant images are inherently incorrect measurements of temperature

- Images can be difficult to interpret accurately when based upon certain objects, specifically objects with erratic temperatures, although this problem is reduced in active thermal imaging
- Accurate temperature measurements are hindered by differing emissivities and reflections from other surfaces
- Most cameras have $\pm 2\%$ accuracy or worse in measurement of temperature and are not as accurate as contact methods
- Only able to directly detect surface temperatures

Applications

- Condition monitoring
- Low Slope and Flat Roofing Inspections
- Building diagnostics including building envelope inspections, moisture inspections, and energy losses in buildings
- Thermal Mapping

3.5 Thermochromic Liquid Crystals (TLC)

Materials that change their reflected color as a function of temperature when illuminated by white light. Hence, reflect visible light at different wavelengths.

Liquid Crystal Thermography (LCT) in a Nutshell Simplest Implementation, household temperature indicator

Process:

A heated surface and A liquid crystal with a known color-to-temperature response Example Fish-tank thermometers, Mood rings, Color sensitive coffee cup, etc.!

Liquid Crystal Thermography

Process Specimen preparation and Light source

To ensure good measurement, the goal is to have a smooth and contaminant free calibration and the test specimen surfaces.

Results are brilliant colors and accurate measurement. Preparation Process

Clean calibration and the test specimen surfaces (if possible) with alcohol and ensure that surfaces are dry.

Apply a “thin and uniform” coat of black paint to the test specimen and the calibration surface

Dry the surfaces with a hot air gun at a mild temperature.

Spray or apply the desired TLC material to both surfaces simultaneously.

Measurement Process

A bright and stable white light source is required to obtain accurate and reliable reflected light intensity from a TLC coated surface

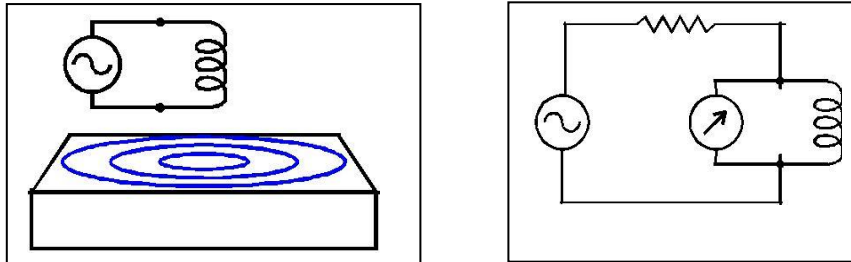
The light source must be void of infrared (IR) and ultra-violet (UV) radiation.

Any IR energy present in the incident light will cause radiant heating of the test surface.

Extended exposure to UV radiation can cause rapid deterioration of the TLC surface. This causes the surface to produce unreliable color-temperature response performance.

Consistent light source settings and lighting-viewing arrangements between calibration and actual testing are essential to minimize color-temperature interpretation errors

3.6 EDDY CUURENT TESSTING



When an AC current flows in a coil in close proximity to a conducting surface the magnetic field of the coil will induce circulating (eddy) currents in that surface. The magnitude and phase of the eddy currents will affect the loading on the coil and thus its impedance. As an example, assume that there is a deep crack in the surface immediately underneath the coil. This will interrupt or reduce the eddy current flow, thus decreasing the loading on the coil and increasing its effective impedance. This is the basis of eddy current testing, by monitoring the voltage across the coil in such an arrangement we can detect changes in the material of interest. Note that cracks must interrupt the surface eddy current flow to be detected. Cracks lying parallel to the current path will not cause any significant interruption and may not be detected. Crack parallel to eddy currents - not detected. Crack interrupts eddy currents - detected.

3.6.1 Factors affecting eddy current response

A number of factors, apart from flaws, will affect the eddy current response from a probe. Successful assessment of flaws or any of these factors relies on holding the others constant, or somehow eliminating their effect on the results. It is this elimination of undesired response that forms the basis of much of the technology of Eddy current inspection.

The main factors are:

Material conductivity

The conductivity of a material has a very direct effect on the eddy current flow: the greater the conductivity of a material the greater the flow of eddy currents on the surface. Conductivity is often measured by an eddy current technique, and inferences can then be drawn about the different factors affecting conductivity, such as material composition, heat treatment, work hardening etc.

Permeability

This may be described as the ease with which a material can be magnetised. For non-ferrous metals such as copper, brass, aluminium etc., and for austenitic stainless steels the permeability is the same as that of 'free space', i.e. the relative permeability (μ_r) is one. For ferrous metals however the value of μ_r may be several hundred, and this has a very significant influence on the eddy current response, in addition it is not uncommon for the permeability to vary greatly within a metal part due to localised stresses, heating effects etc.

Frequency

As we will discuss, eddy current response is greatly affected by the test frequency chosen, fortunately this is one property we can control.

Geometry

In a real part, for example one which is not flat or of infinite size, Geometrical features such as curvature, edges, grooves etc. will exist and will effect the eddy current response. Test techniques must recognise this, for example in testing an edge for cracks the probe will normally be moved along parallel to the edge so that small changes may be easily seen. Where the material thickness is less than the effective depth of penetration (see below) this will also effect the eddy current response.

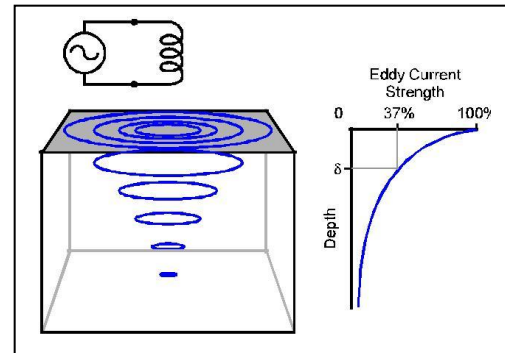
Proximity / Lift-off

The closer a probe coil is to the surface the greater will be the effect on that coil. This has two main effects:

- The “lift-off” signal as the probe is moved on and off the surface.
- A reduction in sensitivity as the coil to product spacing increases.

Depth of Penetration

The eddy current density, and thus the strength of the response from a flaw, is greatest on the surface of the metal being tested and declines with depth. It is mathematically convenient to define the “standard depth of penetration” where the eddy current is $1/e$ (37%) of its surface value.



Coil Configurations

Appropriate coil selection is the most important part of solving an eddy current application, no instrument can achieve much if it doesn't get the right signals from the probe.

Coil designs can be split into three main groups:

1. Surface probes used mostly with the probe axis normal to the surface, in addition to the basic ‘pancake’ coil this includes pencil probes and special-purpose surface probes such as those used inside a fastener hole.

2. Encircling coils are normally used for in-line inspection of round products, The product to be tested is inserted through

3. a circular coil. ID probes are normally used for in-service inspection of heat exchangers. The probe is inserted into the tube. Normally ID probes are wound with the coil axis along the centre of the tube. These categories are not exhaustive and there are obviously overlaps, for example between non-circumferential wound ID probes and internal surface probes. To this point we have only discussed eddy current probes consisting of a single coil, These are commonly used in many applications and are commonly known as absolute probes because they give an ‘absolute’ value of the condition at the test point. Absolute probes are very good for metal sorting and detection of cracks in many situations, however they are sensitive also to material variations, temperature changes etc.

4. Differential probe: Another commonly used probe type is the ‘differential’ probe this has two sensing elements looking at different areas of the material being tested. The instrument responds to the difference between the eddy current conditions at the two points. Differential probes are particularly good for detection of small defects, and are relatively unaffected by lift-off (although the sensitivity is reduced in just the same way), temperature changes and (assuming the instrument circuitry operates in a “balanced” configuration) external interference

3.7 Practical Testing

Any practical Eddy current test will require the following:

- A suitable probe
- An instrument with the necessary capabilities.
- A good idea of size, location and type of the flaws it is desired to find
- A suitable test standard to set up the equipment and verify correct operation
- A procedure or accept/reject criteria based on the above.
- The necessary operator expertise to understand and interpret the results.

Typical Instrumentation

There are a number of basic groups of eddy current instrumentation

Special purpose equipment:

Coating thickness meters, conductivity meters (e.g. Hocking AutoSigma) Generally designed to give a digital readout without the operator needing to understand much about the internal technology, except as needed to give reliable test conditions)

“Crack detectors”

fairly simple equipment, generally operates at a restricted number of frequencies typically several hundred kHz, Meter or Bar-graph display. Suitable for surface crack detection and simple sorting applications only. e.g. Hocking Locator and QuickCheck Normally have some means of compensating for lift-off (e.g. phase rotation and/or fine frequency adjustment) so that only crack-like indications give a reading on the meter or bargraph. An alarm threshold is usually included.

Portable impedance plane

Eddy current Flaw detectors Give a real impedance plane display on a CRT or other electronic display (LCD, plasma etc.) Generally have fairly extensive capabilities: Wide frequency ranges from around a hundred hertz to several megahertz, extensive alarm facilities, general purpose units may have rate filtering (see below) some instruments may be capable of multifrequency operation, allowing combination of results at two or more test frequencies in order to reduce or eliminate specific interfering effects. “Systems” eddy current units. Intended for factory operation, often in Automatic or Semi-Automatic inspection machines. Generally similar operation to impedance plane portables but usually have extensive input and output facilities such as relays and photocell inputs. May be custom built for a specific purpose, in which case features not needed for the intended application are often omitted.

Meter/CRT Instruments

Typical examples (simplified). Hocking NDT Locator UH

Typical Application:

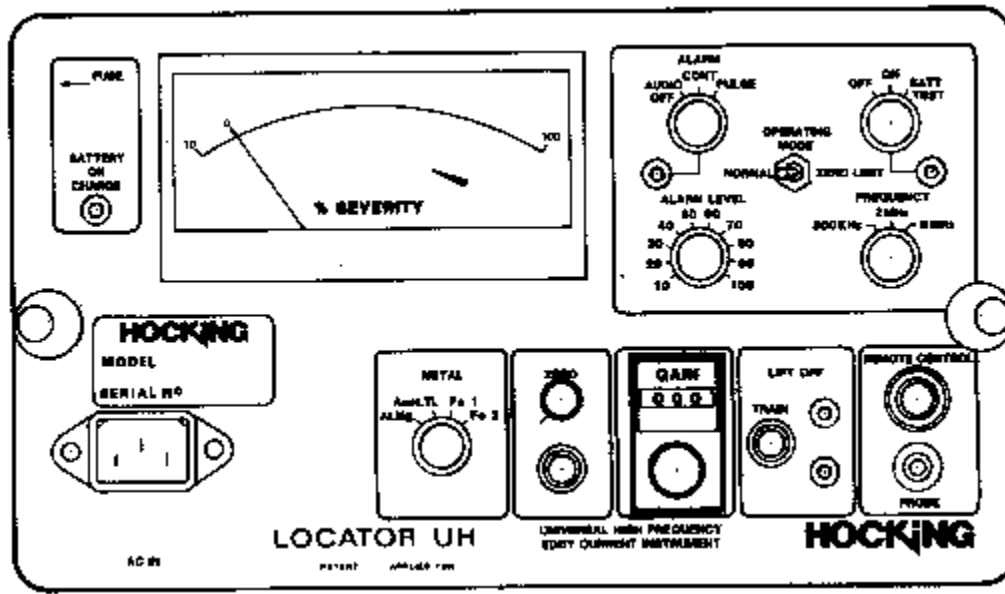
Surface crack detection in aircraft parts using absolute probe.

Controls

Meter display indicates ‘crack severity’ - imbalance from zero point.

Zero - balance internal circuitry

Zero Offset - shift zero point, useful for sorting/material



3.8 Applications

3.8.1 Surface crack detection.

Normally carried out with pencil probes or 'pancake' type probes on ferrous or non-ferrous metals. Frequencies from 100 kHz to a few MHz are commonly used. Depending on surface condition it is usually possible to find cracks as small as 0.1 mm or so deep. Differential probes are sometimes used, particularly in automated applications, care must be taken to ensure that the orientation of flaws is correct for detection.

3.8.2 Non-ferrous metal sorting

This is essentially conductivity testing and for dedicated applications a conductivity meter may be a better choice. From the impedance plane diagram it will be seen that the indication from a conductivity change is essentially the same as from a crack, and both meter and impedance plane type crack detectors can be successfully used to sort similar metals using a suitable absolute probe. It should be remembered that widely different metals may have similar conductivity and that the allowable values for similar metals may overlap, so conductivity measurement should only be used as an indication that a metal is of correct composition or heat-treatment.

3.8.3 Sub-surface crack/corrosion detection.

Primarily used in Airframe inspection. By using a low frequency and a suitable probe eddy currents can penetrate aluminium or similar structures to a depth of 10mm or so, allowing the detection of second and third layer cracking, which is invisible from the surface, or thinning of any of the different layers making up the structure.

3.8.4 Heat exchanger tube testing

Heat exchangers used for petrochemical or power generation applications may have many thousands of tubes, each up to 20m long. Using a differential ID 'bobbin' probe these tubes can be tested at high speed (up to 1 m/s or so with computerized data analysis.) and by using phase analysis defects such as pitting can be assessed to an accuracy of about 5% of tube wall thickness. This allows accurate estimation of the remaining life of the tube allowing operators to decide on appropriate action such as tube plugging, tube replacement or replacement of the complete heat exchanger. The operating frequency is determined by the tube material and wall thickness, ranging from a few kHz for thick-walled copper tube up around 600 kHz for thin-walled titanium. Tubes up to around 50mm diameter are commonly inspected with this technique. Inspection of ferrous or magnetic stainless steel tubes is not possible using standard eddy current inspection equipment. Dual or multiple frequency inspections are commonly used for tubing inspection. In particular for suppression of unwanted responses due to tube support plates.. By subtracting the result of a lower frequency test (which gives a proportionately greater response from the support) a mixed signal is produced showing little or no support plate indication, thus allowing the assessment of small defects in this area.

3.8 .5 In-Line inspection of Steel tubing

Almost all high-quality steel tubing is eddy current inspected using encircling coils . When the tube is made of a magnetic material there are two main problems:

1. Because of the high permeability there is little or no penetration of the eddy current field into the tube at practical test frequencies.
2. Variations in permeability (from many causes) cause eddy current responses which are orders of magnitude greater than those from defects.

These problems may be overcome by magnetically saturating the tube using a strong DC field. This reduces the effective permeability to a low value, allowing effective testing.

Tubes up to around 170mm diameter are commonly tested using magnetic saturation and encircling coils. When tubes are welded this is usually where the problems occur, and so welded tubes are commonly tested in-line using sector coils which only test the weld zone.

Ferrous weld inspection

The geometry and heat-induced material variations around welds in steel would normally prevent inspection with a conventional eddy current probe, however a special purpose "WeldScan" probe has been developed which allows inspection of welded steel structures for fatigue-induced cracking, the technique is particularly useful as it may be used in adverse conditions, or even underwater, and will operate through paint and other corrosion-prevention coatings. Cracks around 1mm deep and 6mm long can be found in typical welds.

Instrument set-up

While the precise details of setting up an instrument will vary depending on the type and application the general procedure is usually the same, obviously once the application has been tried the required values for many test parameters will be known, at least approximately, Connect up the appropriate probe and set any instrument configuration parameters.(mode of operation, display type etc.)

- Set the frequency as required for the test.
- Set gain to an intermediate value,

- Move the probe on/on/over the calibration test piece and set phase rotation as desired (e.g. lift-off or wobble horizontal on a CRT)
- Move over the defects and adjust gain (and horizontal/vertical gain ratio if fitted) to obtain the desired trace size/meter indication. It may be necessary to rebalance after changing gain. Further optimise phase rotation as required. Use filters etc. to further optimise signal to noise ratio.
- Set alarms etc. as required.
- Run over the calibration test piece again and verify that all flaws are clearly detected.
- Perform the test, verifying correct operation at regular intervals using the calibration test piece.

3.9 Pulsed eddy current

Conventional ECT uses sinusoidal alternating current of a particular frequency to excite the probe. Pulsed eddy current (PEC) testing uses a step function voltage to excite the probe. The advantage of using a step function voltage is that such a voltage contains a range of frequencies. As a result, the electromagnetic response to several different frequencies can be measured with just a single step. Since depth of penetration depends on the excitation frequency, information from a range of depths can be obtained all at once. If measurements are made in the time domain (that is, by looking at the strength of the signal as a function of time), indications produced by defects and other features near the inspection coil can be seen first and more distant features will be seen later in time.

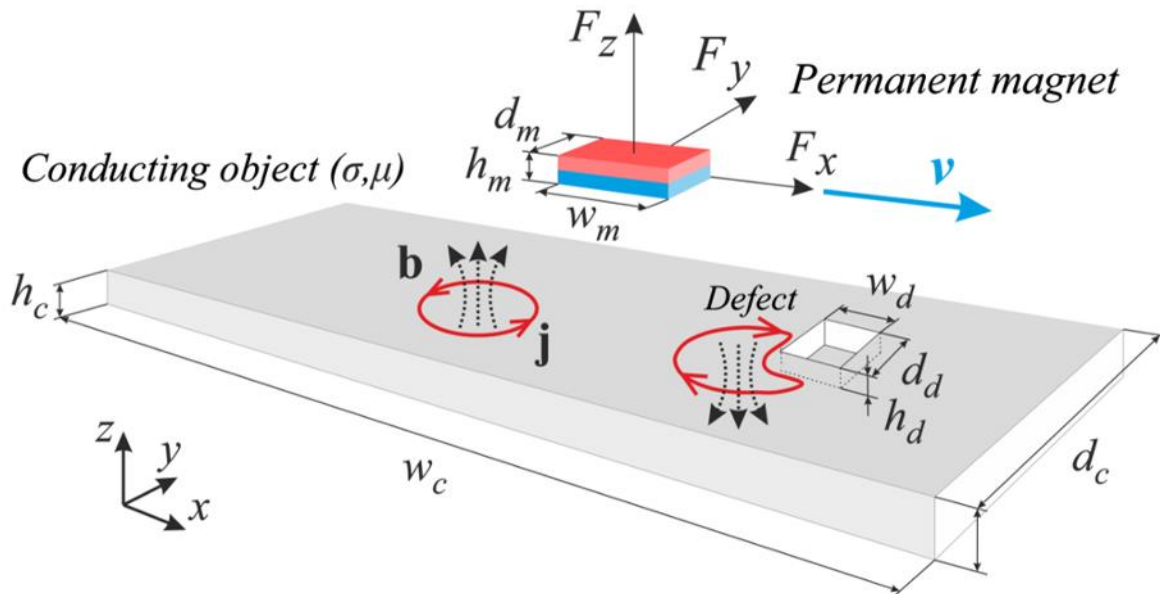
When comparing PEC testing with the conventional ECT, ECT must be regarded as a continuous-wave method where propagation takes place at a single frequency or, more precisely, over a very narrow-frequency bandwidth. With pulse methods, the frequencies are excited over a wide band, the extent of which varies inversely with the pulse length; this allows multi-frequency operation. The total amount of energy dissipated within a given period of time is considerably less for pulsed waves than for continuous waves of the same intensity, thus allowing higher input voltages to be applied to the exciting coil for PEC than conventional ECT. One of the advantage of this type of testing is that there is no need for direct contact with the tested object. Testing can be performed through coatings, sheathings, corrosion products and insulation materials. This way even high-temperature inspections are possible.

3.10 Eddy current array

Eddy current array (ECA) and conventional ECT share the same basic working principles. ECA technology provides the ability to electronically drive an array of coils (multiple coils) arranged in specific pattern called a topology that generates a sensitivity profile suited to the target defects. Data acquisition is achieved by multiplexing the coils in a special pattern to avoid mutual inductance between the individual coils. The benefits of ECA are

- Faster inspections
- Wider coverage
- Less operator dependence — array probes yield more consistent results compared to manual raster scans
- Better detection capabilities
- Easier analysis because of simpler scan patterns
- Improved positioning and sizing because of encoded data

- Array probes can easily be designed to be flexible or shaped to specifications, making hard-to-reach areas easier to inspect



3.11 Lorentz force eddy current testing

A different, albeit physically closely related challenge is the detection of deeply lying flaws and inhomogeneities in electrically conducting solid materials.

In the traditional version of eddy current testing an alternating (AC) magnetic field is used to induce eddy currents inside the material to be investigated. If the material contains a crack or flaw which make the spatial distribution of the electrical conductivity nonuniform, the path of the eddy currents is perturbed and the impedance of the coil which generates the AC magnetic field is modified. By measuring the impedance of this coil, a crack can hence be detected. Since the eddy currents are generated by an AC magnetic field, their penetration into the subsurface region of the material is limited by the skin effect. The applicability of the traditional version of eddy current testing is therefore limited to the analysis of the immediate vicinity of the surface of a material, usually of the order of one millimeter. Attempts to overcome this fundamental limitation using low frequency coils and superconducting magnetic field sensors have not led to widespread applications.

A recent technique, referred to as Lorentz force eddy current testing (LET), exploits the advantages of applying DC magnetic fields and relative motion providing deep and relatively fast testing of electrically conducting materials. In principle, LET represents a modification of the traditional eddy current testing from which it differs in two aspects, namely (i) how eddy currents are induced and (ii) how their perturbation is detected. In LET eddy currents are generated by providing the relative motion between the conductor under test and a permanent magnet (see figure). If the magnet is passing by a defect, the Lorentz force acting on it shows a distortion whose detection is the key for the LET working principle. If the object is free of defects, the resulting Lorentz force remains constant.

Advantages:

Sensitivity to surface defects. Able to detect defects of 0.5mm in length under favourable conditions.

Can detect through several layers. The ability to detect defects in multi-layer structures (up to about 14 layers), without interference from the planar interfaces.

Can detect through surface coatings. Able to detect defects through non-conductive surface coatings in excess of 5mm thickness.

Accurate conductivity measurements. Dedicated conductivity measurement instruments operate using eddy currents.

Can be automated. Relatively uniform parts can be inspected quickly and reliably using automated or semi-automated equipment, e.g. wheels, boiler tubes and aero-engine disks.

Little pre-cleaning required. Only major soils and loose or uneven surface coatings need to be removed, reducing preparation time. as small as a video cassette box and weighing less than 2kg.

Disadvantages

Very susceptible to magnetic permeability changes. Small changes in permeability have a pronounced effect on the eddy currents, especially in ferromagnetic materials. This makes testing of welds and other ferromagnetic materials difficult but, with modern digital flaw detectors and probe design, not impossible.

Only effective on conductive materials. The material must be able to support a flow of electrical current. This makes testing of fibre reinforced plastics unfeasible.

Will not detect defects parallel to surface. The flow of eddy currents is always parallel to the surface. If a planar defect does not cross or interfere with the current then the defect will not be detected.

Not suitable for large areas and/or complex geometries. Large area scanning can be accomplished, but needs the aid of some type of area scanning device, usually supported by a computer, both of which are not inexpensive. The more complex the geometry becomes, the more difficult it is to differentiate defect signals from geometry effect signals.

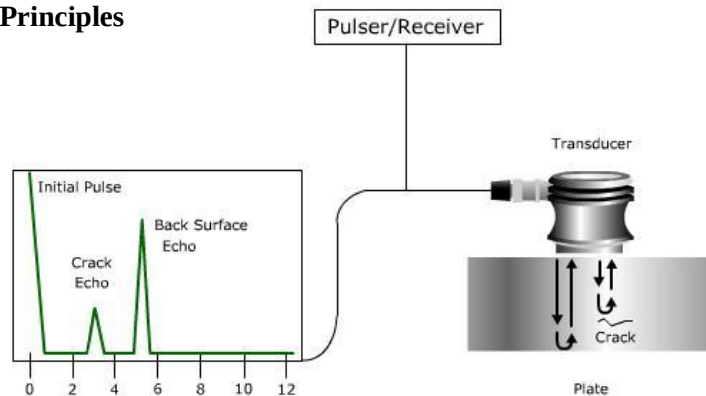
Signal interpretation required. Due to the many factors which affect eddy currents, careful interpretation of signals is needed to distinguish between relevant and non-relevant indications.

No permanent record (unless automated). Normally the only permanent record will be a paper print out or computer file when using automated systems.

4.1 Ultrasonic Testing

Ultrasonic Testing (UT) uses high frequency sound waves (typically in the range between 0.5 and 15 MHz) to conduct examinations and make measurements. Besides its wide use in engineering applications (such as *flaw detection/evaluation, dimensional measurements, material characterization, etc.*), ultrasonics are also used in the medical field (such as *sonography, therapeutic ultrasound, etc.*). In general, ultrasonic testing is based on the capture and quantification of either the reflected waves (*pulse-echo*) or the transmitted waves (*through-transmission*). Each of the two types is used in certain applications, but generally, pulse echo systems are more useful since they require one-sided access to the object being inspected.

4.1.1 Basic Principles



A typical pulse-echo UT inspection system consists of several functional units, such as the pulser/receiver, transducer, and a display device. A pulser/receiver is an electronic device that can produce high voltage electrical pulses. Driven by the pulser, the transducer generates high frequency ultrasonic energy. The sound energy is introduced and propagates through the materials in the form of waves. When there is a discontinuity (such as a crack) in the wave path, part of the energy will be reflected back from the flaw surface. The reflected wave signal is transformed into an electrical signal by the transducer and is displayed on a screen. Knowing the velocity of the waves, travel time can be directly related to the distance that the signal traveled. From the signal, information about the reflector location, size, orientation and other features can sometimes be gained.

4.1.2 Advantages and Disadvantages

Advantages

- ☐ It is sensitive to both surface and subsurface discontinuities.
- ☐ The depth of penetration for flaw detection or measurement is superior to other NDT methods
- ☐ Only single-sided access is needed when the pulse-echo technique is used.
- ☐ It is highly accurate in determining reflector position and estimating size and shape.
- ☐ Minimal part preparation is required.
- ☐ It provides instantaneous results.
- ☐ Detailed images can be produced with automated systems.
- ☐ It is nonhazardous to operators or nearby personnel and does not affect the material being tested.
- ☐ It has other uses, such as thickness measurement, in addition to flaw detection.
- ☐ Its equipment can be highly portable or highly automated.

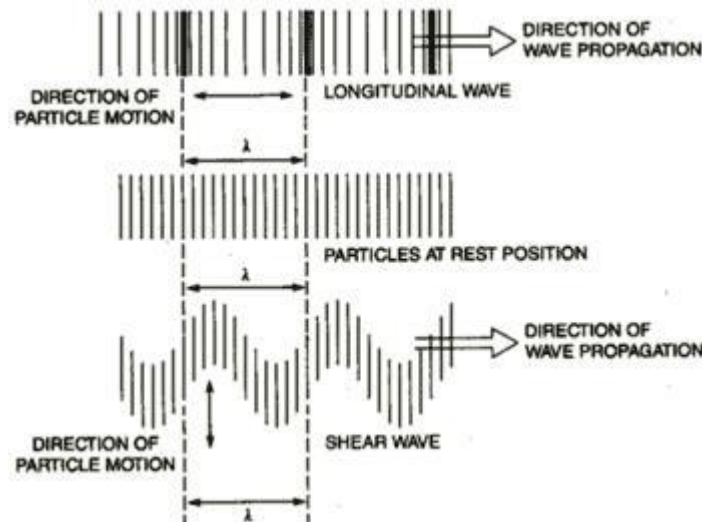
Disadvantages

- ☐ Surface must be accessible to transmit ultrasound.
- ☐ Skill and training is more extensive than with some other methods.
- ☐ It normally requires a coupling medium to promote the transfer of sound energy into test specimen.
- ☐ Materials that are rough, irregular in shape, very small, exceptionally thin or not homogeneous are difficult to inspect.
- ☐ Cast iron and other coarse grained materials are difficult to inspect due to low sound transmission and high signal noise.
- ☐ Linear defects oriented parallel to the sound beam may go undetected.
- ☐ Reference standards are required for both equipment calibration and the characterization of flaws.

4.2 PHYSICS OF ULTRASOUND

Wave Propagation

Ultrasonic testing is based on the vibration in materials which is generally referred to as acoustics. All material substances are comprised of atoms, which may be forced into vibrational motion about their equilibrium positions. Many different patterns of vibrational motion exist at the atomic level; however, most are irrelevant to acoustics and ultrasonic testing. Acoustics is focused on particles that contain many atoms that move in harmony to produce a mechanical wave. When a material is not stressed in tension or compression beyond its elastic limit, its individual particles perform elastic oscillations. When the particles of a medium are displaced from their equilibrium positions, internal restoration forces arise. These elastic restoring forces between particles, combined with inertia of the particles, lead to the oscillatory motions of the medium.



In solids, sound waves can propagate in four principal modes that are based on the way the particles oscillate. Sound can propagate as longitudinal waves, shear waves, surface waves, and in thin materials as plate waves. Longitudinal and shear waves are the two modes of propagation most widely used in ultrasonic testing. The particle movement responsible for the propagation of longitudinal and shear waves is illustrated in the figure

□ In *longitudinal waves*, the oscillations occur in the longitudinal direction or the direction of wave propagation. Since compression and expansion forces are active in these waves, they are also called pressure or compression waves. They are also sometimes called density waves because material density fluctuates as the wave moves. Compression waves can be generated in gases, liquids, as well as solids because the energy travels through the atomic structure by a series of compressions and expansion movements.

□ oscillate at a right angle or transverse to the direction of propagation. Shear waves require an acoustically solid material for effective propagation, and therefore, are not effectively propagated in materials such as liquids or gasses. Shear waves are relatively weak when compared to longitudinal waves. In fact, shear waves are usually generated in materials using some of the energy from longitudinal waves.

Modes of Sound Wave Propagation

In air, sound travels by the compression and rarefaction of air molecules in the direction of travel. However, in solids, molecules can support vibrations in other directions. Hence, a number of different types of sound waves are possible. Waves can be characterized by oscillatory patterns that are capable of maintaining their shape and propagating in a stable manner. The propagation of waves is often described in terms of what are called “*wave modes*”. As mentioned previously, longitudinal and transverse (shear) waves are most often used in ultrasonic inspection. However, at surfaces and interfaces, various types of elliptical or complex vibrations of the particles make other waves possible. Some of these wave modes such as Rayleigh and Lamb waves are also useful for ultrasonic inspection. Though there are many different modes of wave propagation, the table summarizes the four types of waves that are used in NDT

Wave Type	Particle Vibration
-----------	--------------------

Longitudinal (Compression)	Parallel to wave direction
Transverse (Shear)	Perpendicular to wave direction
Surface - Rayleigh	Elliptical orbit - symmetrical mode
Plate Wave - Lamb	Component perpendicular to surface

Since longitudinal and transverse waves were discussed previously, surface and plate waves are introduced here.

Surface (or Rayleigh) waves travel at the surface of a relatively thick solid material penetrating to a depth of one wavelength. A surface wave is a combination of both a longitudinal and transverse motion which results in an elliptical motion as shown in the image. The major axis of the ellipse is perpendicular to the surface of the solid. As the depth of an individual atom from the surface increases, the width of its elliptical motion decreases. Surface waves are generated when a longitudinal wave intersects a surface slightly larger than the second critical angle and they travel at a velocity between .87 and .95 of a shear wave. Rayleigh waves are useful because they are very sensitive to surface defects (*and other surface features*) and they follow the surface around curves. Because of this, Rayleigh waves can be used to inspect areas that other waves might have difficulty reaching.

Plate (or Lamb) waves are similar to surface waves except they can only be generated in materials a few wavelengths thick (*thin plates*). Lamb waves are complex vibrational waves that propagate parallel to the test surface throughout the thickness of the material. They are influenced a great deal by the test wave frequency and material thickness. Lamb waves are generated when a wave hits a surface at an incident angle such that the parallel component of the velocity of the wave (in the source) is equal to the velocity of the wave in the test material. Lamb waves will travel several meters in steel and so are useful to scan plate, wire, and tubes. With Lamb waves, a number of modes of particle vibration are possible, but the two most common are symmetrical and asymmetrical. The complex motion of the particles is similar to the elliptical orbits for surface waves. Symmetrical Lamb waves move in a symmetrical fashion about the median plane of the plate. This is sometimes called the “*extensional mode*” because the wave is stretching and compressing the plate in the wave motion direction. The asymmetrical Lamb wave mode is often called the “*flexural mode*” because a large portion of the motion is in a normal direction to the plate, and a little motion occurs in the direction parallel to the plate. In this mode, the body of the plate bends as the two surfaces move in the same direction.

4.3 Properties of Acoustic Waves

Among the properties of waves propagating in isotropic solid materials are wavelength, frequency, and velocity. The wavelength is directly proportional to the velocity of the wave and inversely proportional to the frequency of the wave. This relationship is shown by the following equation:

$$\lambda = V/f$$

Where;

λ : wavelength (m)

V: velocity (m/s)

f: frequency (Hz)

The velocity of sound waves in a certain medium is fixed where it is a characteristic of that medium. As can be noted from the equation, an increase in frequency will result in a decrease in wavelength. For instance, the velocity of longitudinal waves in steel is 5850 m/s and that results in a wavelength of 5.85 mm when the frequency is 1 MHz.

Wavelength and Defect Detection

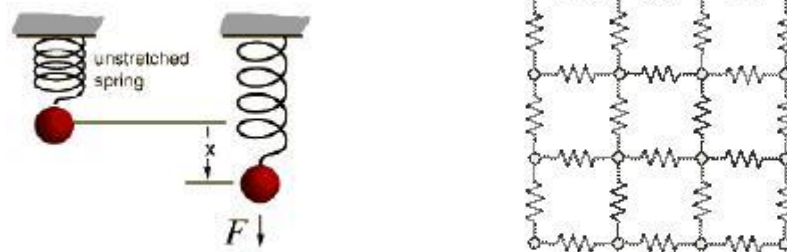
In ultrasonic testing, the inspector must make a decision about the frequency of the transducer that will be used in order to control the wavelength. The wavelength of the ultrasound used has a significant effect on the probability of detecting a discontinuity. A general rule of thumb is that a discontinuity must be larger than one-half the wavelength to stand a reasonable chance of being detected.

Sensitivity and **resolution** are two terms that are often used in ultrasonic inspection to describe a technique's ability to locate flaws. *Sensitivity* is the ability to locate small discontinuities. Sensitivity generally increases with higher frequency (*shorter wavelengths*). *Resolution* is the ability of the system

to locate discontinuities that are close together within the material or located near the part surface. Resolution also generally increases as the frequency increases. The wave frequency can also affect the capability of an inspection in adverse ways. Therefore, selecting the optimal inspection frequency often involves maintaining a balance between the favorable and unfavorable results of the selection. Before selecting an inspection frequency, the material's grain structure and thickness, and the discontinuity's type, size, and probable location should be considered. As frequency increases, sound tends to scatter from large or coarse grain structure and from small imperfections within a material. Cast materials often have coarse grains and thus require lower frequencies to be used for evaluations of these products. Wrought and forged products with directional and refined grain structure can usually be inspected with higher frequency transducers. Since more things in a material are likely to scatter a portion of the sound energy at higher frequencies, the **penetration depth** (*the maximum depth in a material that flaws can be located*) is also reduced. Frequency also has an effect on the shape of the ultrasonic beam. Beam spread, or the divergence of the beam from the center axis of the transducer, and how it is affected by frequency will be discussed later. It should be mentioned, so as not to be misleading, that a number of other variables will also affect the ability of ultrasound to locate defects. These include the pulse length, type and voltage applied to the crystal, properties of the crystal, backing material, transducer diameter, and the receiver circuitry of the instrument. These are discussed in more detail in a later section.

4.4 Sound Propagation in Elastic Materials

It was mentioned previously that sound waves propagate due to the vibrations or oscillatory motions of particles within a material. An ultrasonic wave may be visualized as an infinite number of oscillating masses or particles connected by means of elastic springs. Each individual particle is influenced by the motion of its nearest neighbor and both inertial and elastic restoring forces act upon each particle. A mass on a spring has a single resonant frequency (*natural frequency*) determined by its spring constant **k** and its mass **m**. Within the elastic limit of any material, there is a linear relationship between the displacement of a particle and the force attempting to restore the particle to its equilibrium position. This linear dependency is described by *Hooke's Law*. In terms of the spring model, the relation between force and displacement is written as $F = kx$.



The Speed of Sound

Hooke's Law, when used along with Newton's Second Law, can explain a few things about the speed of sound. The speed of sound within a material is a function of the properties of the material and is independent of the amplitude of the sound wave. Newton's Second Law says that the force applied to a particle will be balanced by the particle's mass and the acceleration of the particle. Mathematically, Newton's Second Law is written as $F = ma$. Hooke's Law then says that this force will be balanced by a force in the opposite direction that is dependent on the amount of displacement and the spring constant. Therefore, since the applied force and the restoring force are equal, $ma = kx$ can be written. Since the mass **m** and the spring constant **k** are constants for any given material, it can be seen that the acceleration **a** and the displacement **x** are the only variables. It can also be seen that they are directly proportional. For instance, if the displacement of the particle increases, so does its acceleration. It turns out that the time that it takes a particle to move and return to its equilibrium position is independent of the force applied. So, within a given material, sound always travels at the same speed no matter how much force is applied when other variables, such as temperature, are held constant.

4.5 Material Properties Affecting the Speed of Sound

Of course, sound does travel at different speeds in different materials. This is because the mass of the atomic particles and the spring constants are different for different materials. The mass of the particles is related to the density of the material, and the spring constant is related to the elastic constants of a material. The general relationship between the speed of sound in a solid and its density and elastic constants is given by the following equation:

$$V = \sqrt{\frac{C}{\rho}}$$

Where;

V: speed of sound (m/s)

C: elastic constant “in a given direction” (N/m²)

ρ : density (kg/m³)

This equation may take a number of different forms depending on the type of wave

(*longitudinal or shear*) and which of the elastic constants that are used. It must also be mentioned that the subscript “ ” attached to “ ” in the above equation is used to indicate the directionality of the elastic constants with respect to the wave type and direction of wave travel. In isotropic materials, the elastic constants are the same for all directions within the material. However, most materials are anisotropic and the elastic constants differ with each direction. For example, in a piece of rolled aluminum plate, the grains are elongated in one direction and compressed in the others and the elastic constants for the longitudinal direction differs slightly from those for the transverse or short transverse directions.

For longitudinal waves, the speed of sound in a solid material can be found as:

$$V_L = \sqrt{\frac{E(1-\nu)}{\rho(1+\nu)(1-2\nu)}}$$

where;

V_L: speed of sound for longitudinal waves (m/s)

E: Young's modulus (N/m²)

ν : Poisson's ratio

While for shear (*transverse*) waves, the speed of sound is found as:

$$V_T = \sqrt{\frac{G}{\rho}}$$

Where;

V_T: speed of sound for shear waves (m/s)

G: Shear modulus of elasticity (N/m²);

4.6 Attenuation of Sound Waves

When sound travels through a medium, its intensity diminishes with distance. In idealized materials, sound pressure (*signal amplitude*) is reduced due to the spreading of the wave. In natural materials, however, the sound amplitude is further weakened due to the scattering and absorption. Scattering is the reflection of the sound in directions other than its original direction of propagation. Absorption is the conversion of the sound energy to other forms of energy. The combined effect of scattering and absorption is called attenuation. Attenuation is generally proportional to the square of sound frequency.

The amplitude change of a decaying plane wave can be expressed as:

$$A = A_0 e^{-\alpha z}$$

Where;

A_0 : initial (*unattenuated*) amplitude

α : attenuation coefficient (Np/m)

Z : traveled distance (m)

Acoustic Impedance

Sound travels through materials under the influence of sound pressure. Because molecules or atoms of a solid are bound elastically to one another, the excess pressure results in a wave propagating through the solid.

The *acoustic impedance* () of a material is defined as the product of its density () and the velocity of sound in that material ().

$$Z = \rho V$$

Where;

Z : acoustic impedance (kg/m^2s) or ($N s/m^3$)

ρ : density (kg/m^3)

V : sound velocity (m/s)

The table gives examples of the acoustic impedances for some materials:

	Aluminum	Copper	Steel	Titanium	Water (20°C)	Air (20°C)
Acou. Imp. (kg/m^2s)	17.1 x 10 ⁶	41.6 x 10 ⁶	46.1 x 10 ⁶	28 x 10 ⁶	1.48 x 10 ⁶	413

Acoustic impedance is important in:

the determination of acoustic transmission and reflection at the boundary of two materials having different acoustic impedances. The design of ultrasonic transducers. Assessing absorption of sound in a medium

Reflection and Transmission Coefficients

Ultrasonic waves are reflected at boundaries where there is a difference in acoustic impedances () of the materials on each side of the boundary. This difference in is commonly referred to as the impedance mismatch. The greater the impedance mismatch, the greater the percentage of energy that will be reflected at the interface or boundary between one medium and another. The fraction of the incident wave intensity that is reflected can be derived based on the fact that particle velocity and local particle pressures must be continuous across the boundary. When the acoustic impedances of the materials on both sides of the boundary are known, the fraction of the incident wave intensity that is reflected (*the reflection coefficient*) can be calculated as:

$$R = \left(\frac{Z_2 - Z_1}{Z_2 + Z_1} \right)^2$$

Refraction and Snell's Law

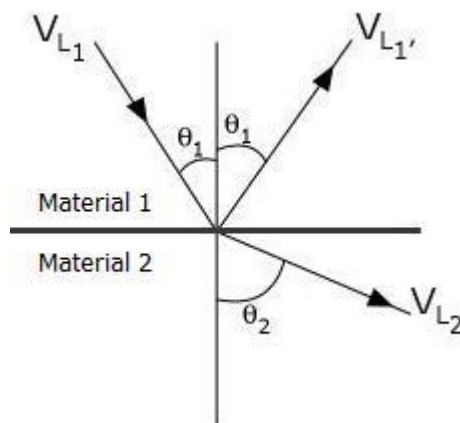
When an ultrasonic wave passes through an interface between two materials at an oblique angle, and the materials have different indices of refraction, both reflected and refracted waves are produced. This also occurs with light, which is why objects seen across an interface appear to be shifted relative to where they really are. For example, if you look straight down at an object at the bottom of a glass of water, it looks closer than it really is.

Refraction takes place at an interface of two materials due to the difference in acoustic velocities between the two materials. The figure shows the case where plane sound waves traveling in one material enters a second material that has a higher acoustic velocity. When the wave encounters the interface between these two materials, the portion of the wave in the second material is moving faster than the portion of the wave that is still in the first material. As a result, this causes the wave to bend and change its direction (*this is referred to as "refraction"*).

$$\frac{\sin \theta_1}{V_{L_1}} = \frac{\sin \theta_2}{V_{L_2}}$$

Snell's Law describes the relationship between the angles and the velocities of the waves. Snell's law equates the ratio of material velocities to the ratio of the sine's of incident and refracted angles, as shown in the following equation:

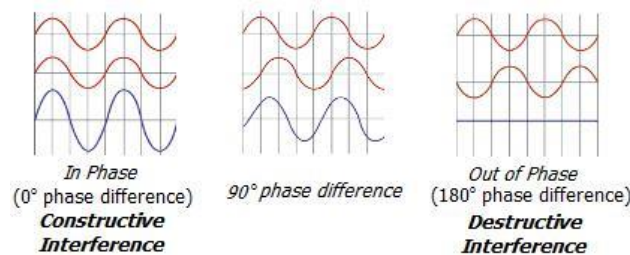
Where;



Wave Interaction or Interference

The understanding of the interaction or interference of waves is important for understanding the performance of an ultrasonic transducer. When sound emanates from an ultrasonic transducer, it does not originate from a single point, but instead originates from many points along the surface of the piezoelectric element. This results in a sound field with many waves interacting or interfering with each other. When waves interact, they superimpose on each other, and the amplitude of the sound pressure at any point of interaction is the sum of the amplitudes of the two individual waves. First, let's consider two identical waves that originate from the same point. When they are *in phase* (so that the peaks and valleys of one are exactly aligned with those of the other), they combine to double the pressure of either wave acting alone. When they are completely *out of phase* (so that the peaks of one wave are exactly aligned with the valleys of the other wave), they combine to cancel each other out. When the two waves are not completely in phase or out of phase, the resulting wave is the sum of the wave amplitudes for all points along the wave. When the origins of the two interacting waves are not the same, it is a little harder to picture the wave interaction, but the principles are the same. Up until now, we have primarily looked at

waves in the form of a 2D plot of wave amplitude versus wave position. However, anyone that has dropped something in a pool of water can picture the waves radiating out from the source with a circular wave front. If two objects are dropped a short distance apart into the pool of water, their waves will radiate out from their sources and interact with each other. At every point where the waves interact, the amplitude of the particle displacement is the combined sum of the amplitudes of the particle displacement of the individual waves. As stated previously, sound waves originate from multiple points along the face of the transducer. The image shows what the sound field would look like if the waves originated from just three points (*of course there are more than three points of origin along the face of a transducer*). It can be seen that where the waves interact near the face of the transducer and as a result there are extensive fluctuations and the sound field is very uneven. In ultrasonic testing, this is known as the “*near field*” or Fresnel zone. The sound field is more uniform away from the transducer in the “*far field*” or Fraunhofer zone. At some distance from the face of the transducer and central to the face of the transducer, a uniform and intense wave field develops.



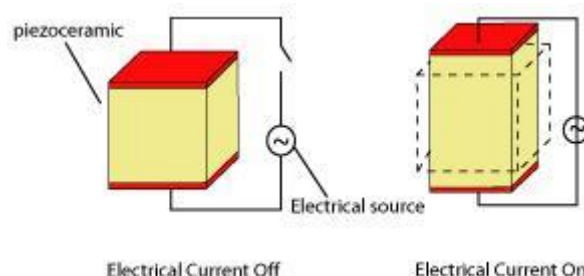
Wave Diffraction

Diffraction involves a change in direction of waves as they pass through an opening or around a barrier in their path. Diffraction of sound waves is commonly observed; we notice sound diffracting around corners or through door openings allowing us to hear others who are speaking to us from adjacent rooms. In ultrasonic testing of solids, diffraction patterns are usually generated at the edges of sharp reflectors (or discontinuities) such as cracks. Usually the tip of a crack behaves as point source spreading waves in all directions due to the diffraction of the incident wave.

4.7 EQUIPMENT & TRANSDUCERS

Piezoelectric Transducers

The conversion of electrical pulses to mechanical vibrations and the conversion of returned mechanical vibrations back into electrical energy is the basis for ultrasonic testing. This conversion is done by the transducer using a piece of piezoelectric material (*a polarized material having some parts of the molecule positively charged, while other parts of the molecule are negatively charged*) with electrodes attached to two of its opposite faces. When an electric field is applied across the material, the polarized molecules will align themselves with the electric field causing the material to change dimensions. In addition, a permanently-polarized material such as quartz (SiO_2) or barium titanate (BaTiO_3) will produce an electric field when the material changes dimensions as a result of an imposed mechanical force. This phenomenon is known as the piezoelectric effect.



The active element of most acoustic transducers used today is a piezoelectric ceramic, which can be cut in various ways to produce different wave modes. A large piezoelectric ceramic element can be seen in the image of a sectioned low frequency transducer. The most commonly employed ceramic for making transducers is lead zirconate titanate. The thickness of the active element is determined by the desired frequency of the transducer. A thin wafer element vibrates with a wavelength that is twice its thickness.

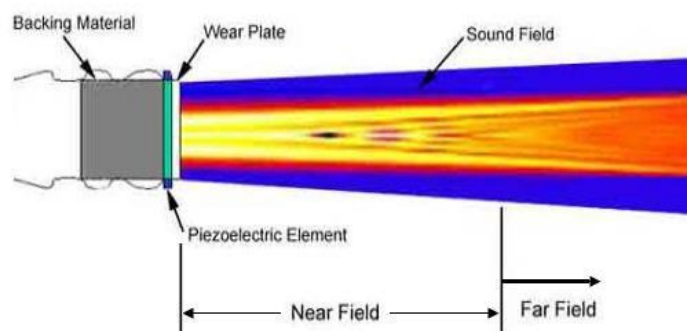
Therefore, piezoelectric crystals are cut to a thickness that is $1/2$ the desired radiated wavelength. The higher the frequency of the transducer, the thinner the active element.

4.8 Characteristics of Piezoelectric Transducers

The function of the transducer is to convert electrical signals into mechanical vibrations (*transmit mode*) and mechanical vibrations into electrical signals (*receive mode*). Many factors, including material, mechanical and electrical construction, and the external mechanical and electrical load conditions, influence the behavior of a transducer. A cut away of a typical contact transducer is shown in the figure. To get as much energy out of the transducer as possible, an impedance matching layer is placed between the active element and the face of the transducer. Optimal impedance matching is achieved by sizing the matching layer so that its thickness is $1/4$ of the desired wavelength. This keeps waves that are reflected within the matching layer in phase when they exit the layer. For contact transducers, the matching layer is made from a material that has an acoustical impedance between the active element and steel. Immersion transducers have a matching layer with an acoustical impedance between the active element and water. Contact transducers also incorporate a wear plate to protect the matching layer and active element from scratching. The backing material supporting the crystal has a great influence on the damping characteristics of a transducer. Using a backing material with an impedance similar to that of the active element will produce the most effective damping. Such a transducer will have a wider bandwidth resulting in higher sensitivity and higher resolution (*i.e., the ability to locate defects near the surface or in close proximity in the material*). As the mismatch in impedance between the active element and the backing material increases, material penetration increases but transducer sensitivity is reduced. The bandwidth refers to the range of frequencies associated with a transducer. The frequency noted on a transducer is the central frequency and depends primarily on the backing material. Highly damped transducers will respond to frequencies above and below the central frequency. The broad frequency range provides a transducer with high resolving power. Less damped transducers will exhibit a narrower frequency range and poorer resolving power, but greater penetration. The central frequency will also define the capabilities of a transducer. Lower frequencies (0.5MHz-2.25MHz) provide greater energy and penetration in a material, while high frequency crystals (15.0MHz-25.0MHz) provide reduced penetration but greater sensitivity to small discontinuities.

Radiated Fields of Ultrasonic Transducers

The sound that emanates from a piezoelectric transducer does not originate from a point, but instead originates from most of the surface of the piezoelectric element. The sound field from a typical piezoelectric transducer is shown in the figure where lighter colors indicating higher intensity. Since the ultrasound originates from a number of points along the transducer face, the ultrasound intensity along the beam is affected by constructive and destructive wave interference as discussed previously. This wave interference leads to extensive fluctuations in the sound intensity near the source and is known as the “*near field*”. Because of acoustic variations within a near field, it can be extremely difficult to accurately evaluate flaws in materials when they are positioned within this area.

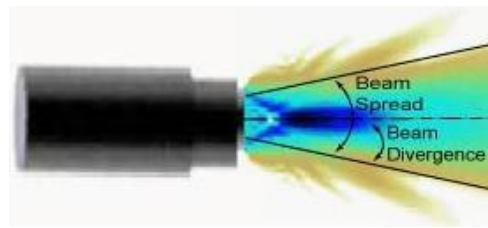


The pressure waves combine to form a relatively uniform front at the end of the near field. The area beyond the near field where the ultrasonic beam is more uniform is called the “*far field*”. The transition between the near field and the far field occurs at a distance, z , and is sometimes referred to as the “*natural focus*” of a flat (or unfocused) transducer. Spherical or cylindrical focusing changes the structure of a transducer field by “pulling” the point nearer the transducer. The area just beyond the near field is where

the sound wave is well behaved and at its maximum strength. Therefore, optimal detection results will be obtained when flaws occur in this area.

4.9 Transducer Beam Spread

As the sound waves exit the near field and propagate through the material, the sound beam continuously spreads out. This phenomenon is usually referred to as beam spread but sometimes it is also referred to as beam divergence or ultrasonic diffraction. It should be noted that there is actually a difference between beam spread and beam divergence. Beam spread is a measure of the whole angle from side to side of the beam in the far field. Beam divergence is a measure of the angle from one side of the sound beam to the central axis of the beam in the far field. Therefore, beam spread is twice the beam divergence. Although beam spread must be considered when performing an ultrasonic inspection, it is important to note that in the far field, or Fraunhofer zone, the maximum sound pressure is always found along the acoustic axis (*centerline*) of the transducer. Therefore, the strongest reflections are likely to come from the area directly in front of the transducer.



Beam spread occurs because the vibrating particle of the material (*through which the wave is traveling*) do not always transfer all of their energy in the direction of wave propagation. If the particles are not directly aligned in the direction of wave propagation, some of the energy will get transferred off at an angle. In the near field, constructive and destructive wave interference fill the sound field with fluctuation. At the start of the far field, however, the beam strength is always greatest at the center of the beam and diminishes as it spreads outward. The beam spread is largely influenced by the frequency and diameter of the transducer. For a flat piston source transducer, an approximation of the beam divergence angle at which the sound pressure has decreased by one half (-6 dB) as compared to its value at the centerline axis can be calculated as:

Transducer Types

Ultrasonic transducers are manufactured for a variety of applications and can be custom fabricated when necessary. Careful attention must be paid to selecting the proper transducer for the application. It is important to choose transducers that have the desired frequency, bandwidth, and focusing to optimize inspection capability. Most often the transducer is chosen either to enhance the sensitivity or resolution of the system. Transducers are classified into two major groups according to the application.

Contact transducers are used for direct contact inspections, and are generally hand manipulated. They have elements protected in a rugged casing to withstand sliding contact with a variety of materials. These transducers have an ergonomic design so that they are easy to grip and move along a surface. They often have replaceable wear plates to lengthen their useful life. Coupling materials of water, grease, oils, or commercial materials are used to remove the air gap between the transducer and the component being inspected.

Immersion transducers do not contact the component. These transducers are designed to operate in a liquid environment and all connections are watertight. Immersion transducers usually have an impedance matching layer that helps to get more sound energy into the water and, in turn, into the component being inspected. Immersion transducers can be purchased with a planar, cylindrically focused or spherically focused lens. A focused transducer can improve the sensitivity and axial resolution by concentrating the sound energy to a smaller area. Immersion transducers are typically used inside a water tank or as part of a squirter or bubbler system in scanning applications.

Other Types of Contact Transducers

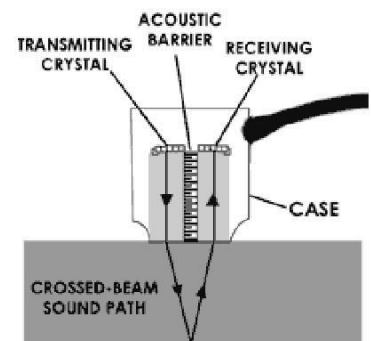
Contact transducers are available in a variety of configurations to improve their usefulness for a variety of applications. The flat contact transducer shown above is used in normal beam inspections of relatively flat surfaces, and where near surface resolution is not critical. If the surface is curved, a shoe that matches the curvature of the part may need to be added to the face of the transducer. If near surface resolution is

important or if an angle beam inspection is needed, one of the special contact transducers described below might be used.

□ *Dual element transducers* contain two independently operated elements in a single housing. One of the elements transmits and the other receives the ultrasonic signal. Dual element transducers are especially well suited for making measurements in applications where reflectors are very near the transducer since this design eliminates the ring down effect that single-element transducers experience (when single-element transducers are operating in pulse echo mode,

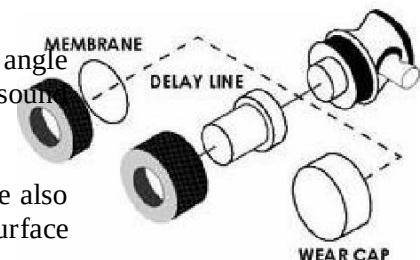
the element cannot start receiving reflected signals until the element has stopped ringing from its transmit function). Dual element transducers are very useful when making thickness measurements of thin materials and when inspecting for near surface defects. The two elements are angled towards each other to create a crossed-beam sound path in the test material.

□ *Delay line transducers* provide versatility with a variety of replaceable options. Removable delay line, surface conforming membrane, and protective wear cap options can make a single transducer effective for a wide range of applications. As the name implies, the primary function of a delay line transducer is to introduce a time delay between the generation of the sound wave and the arrival of any reflected waves. This allows the transducer to complete its "sending" function before it starts its "receiving" function so that near surface resolution is improved. They are designed for use in applications such as high precision thickness gauging of thin materials and delamination checks in composite materials. They are also useful in high-temperature measurement applications since the delay line provides some insulation to the piezoelectric element from the heat. *Angle beam transducers* and wedges are typically used to introduce a refracted shear wave into the test material. Transducers can be purchased in a variety of fixed angles or in adjustable versions where the user determines the angles of incidence and refraction. In the fixed angle versions, the angle of refraction that is marked on the transducer is only accurate for a particular material, which



is usually steel. The most commonly used refraction angles for fixed angle transducers are 45°, 60° and 70°. The angled sound path allows the sound beam to be reflected from the backwall to

improve detectability of flaws in and around welded areas. They are also used to generate surface waves for use in detecting defects on the surface of a component.



□ *Normal incidence shear wave transducers* are unique because they allow the introduction of shear waves directly into a test piece without the use of an angle beam wedge. Careful design has enabled manufacturing of transducers with minimal longitudinal wave contamination.

□ *Paint brush transducers* are used to scan wide areas. These long and narrow transducers are made up of an array of small crystals and that make it possible to scan a larger area more rapidly for discontinuities. Smaller and more sensitive transducers are often then required to further define the details of a discontinuity.



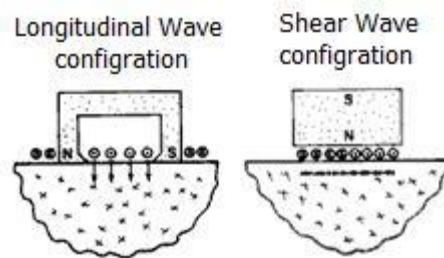
Couplant

A couplant is a material (*usually liquid*) that facilitates the transmission of ultrasonic energy from the transducer into the test specimen. Couplant is generally necessary because the acoustic impedance

mismatch between air and solids is large. Therefore, nearly all of the energy is reflected and very little is transmitted into the test material. The couplant displaces the air and makes it possible to get more sound energy into the test specimen so that a usable ultrasonic signal can be obtained. In contact ultrasonic testing a thin film of oil, glycerin or water is typically used between the transducer and the test surface. When shear waves are to be transmitted, the fluid is generally selected to have a significant viscosity. When scanning over the part, an immersion technique is often used. In immersion ultrasonic testing both the transducer and the part are immersed in the couplant, which is typically water. This method of coupling makes it easier to maintain consistent coupling while moving and manipulating the transducer and/or the part.

4.10 Electromagnetic Acoustic Transducers (EMATs)

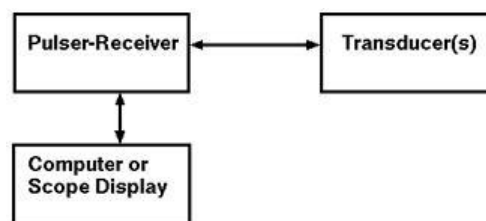
Electromagnetic-acoustic transducers (EMAT) are a modern type of ultrasonic transducers that work based on a totally different physical principle than piezoelectric transducers and, most importantly, they do not need couplant. When a wire is placed near the surface of an electrically conducting object and is driven by a current at the desired ultrasonic frequency, eddy currents will be induced in a near surface region of the object. If a static magnetic field is also present, these eddy currents will experience forces called “*Lorentz forces*” which will cause pressure waves to be generated at the surface and propagate through the material. Different types of sound waves (*longitudinal, shear, lamb*) can be generated using EMATs by varying the configuration of the transducer such that the orientation of the static magnetic field is changed.



EMATs can be used for thickness measurement, flaw detection, and material property characterization. The EMATs offer many advantages based on its non-contact couplant-free operation. These advantages include the ability to operate in remote environments at elevated speeds and temperatures.

Pulser-Receivers

Ultrasonic pulser-receivers are well suited to general purpose ultrasonic testing. Along with appropriate transducers and an oscilloscope, they can be used for flaw detection and thickness gauging in a wide variety of metals, plastics, ceramics, and composites. Ultrasonic pulser-receivers provide a unique, low-cost ultrasonic measurement capability. Specialized portable equipment that are dedicated for ultrasonic inspection merge the pulser-receiver with the scope display in one small size battery operated unit. The pulser section of the instrument generates short, large amplitude electric pulses of controlled energy, which are converted into short ultrasonic pulses when applied to an ultrasonic transducer. Control functions associated with the pulser circuit included



- ☐ **Pulse length or damping:** The amount of time the pulse is applied to the transducer.
- ☐ **Pulse energy:** The voltage applied to the transducer. Typical pulser circuits will apply from 100 volts to 800 volts to a transducer. In the receiver section the voltage signals produced by the transducer, which represent the received ultrasonic pulses, are amplified. The amplified signal is available as an output for display or capture for signal processing. Control functions associated with the receiver circuit include:

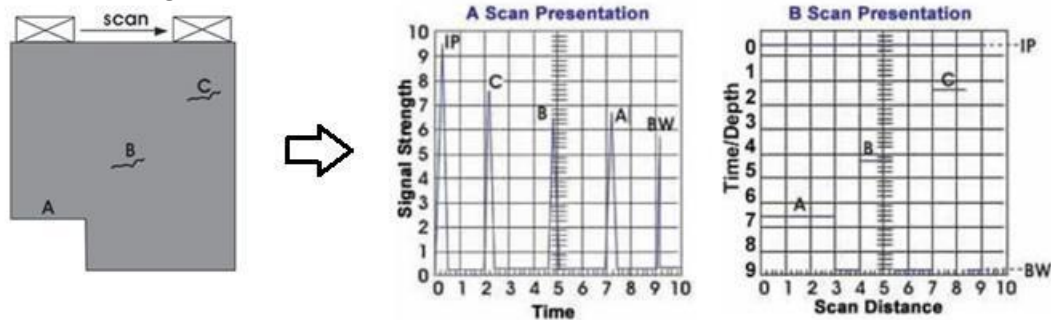
- ☐ *Signal rectification:* The signal can be viewed as positive half wave, negative half wave or full wave.
- ☐ *Filtering to shape and smoothing*
- ☐ *Gain, or signal amplification*
- ☐ *Reject control*

4.11 Data Presentation

Ultrasonic data can be collected and displayed in a number of different formats. The three most common formats are known in the NDT world as A-scan, B-scan and C-scan presentations. Each presentation mode provides a different way of looking at and evaluating the region of material being inspected. Modern computerized ultrasonic scanning systems can display data in all three presentation forms simultaneously.

A-Scan Presentation

The A-scan presentation displays the amount of received ultrasonic energy as a function of time. The relative amount of received energy is plotted along the vertical axis and the elapsed time (*which may be related to the traveled distance within the material*) is displayed along the horizontal axis. Most instruments with an A-scan display allow the signal to be displayed in its natural radio frequency form (*RF*), as a fully rectified *RF* signal, or as either the positive or negative half of the *RF* signal. In the A-scan presentation, relative discontinuity size can be estimated by comparing the signal amplitude obtained from an unknown reflector to that from a known reflector. Reflector depth can be determined by the position of the signal on the horizontal time axis.



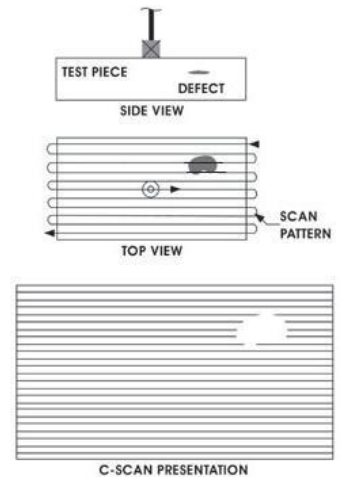
In the illustration of the A-scan presentation shown in the figure, the initial pulse generated by the transducer is represented by the signal **IP**, which is near time zero. As the transducer is scanned along the surface of the part, four other signals are likely to appear at different times on the screen. When the transducer is in its far left position, only the **IP** signal and signal **A**, the sound energy reflecting from surface **A**, will be seen on the trace. As the transducer is scanned to the right, a signal from the backwall **BW** will appear later in time, showing that the sound has traveled farther to reach this surface. When the transducer is over flaw **B**, signal **B** will appear at a point on the time scale that is approximately halfway between the **IP** signal and the **BW** signal. Since the **IP** signal corresponds to the front surface of the material, this indicates that flaw **B** is about halfway between the front and back surfaces of the sample. When the transducer is moved over flaw **C**, signal **C** will appear earlier in time since the sound travel path is shorter and signal **B** will disappear since sound will no longer be reflecting from it.

B-Scan Presentation

The B-scan presentation is a type of presentation that is possible for automated linear scanning systems where it shows a profile (*cross-sectional*) view of the test specimen. In the B-scan, the time-of-flight (*travel time*) of the sound waves is displayed along the vertical axis and the linear position of the transducer is displayed along the horizontal axis. From the B-scan, the depth of the reflector and its approximate linear dimensions in the scan direction can be determined. The B-scan is typically produced by establishing a trigger gate on the A-scan. Whenever the signal intensity is great enough to trigger the gate, a point is produced on the B-scan. The gate is triggered by the sound reflected from the backwall of the specimen and by smaller reflectors within the material. In the B-scan image shown previously, line **A** is produced as the transducer is scanned over the reduced thickness portion of the specimen. When the transducer moves to the right of this section, the backwall line **BW** is produced. When the transducer is over flaws **B** and **C**, lines that are similar to the length of the flaws and at similar depths within the material are drawn on the B-scan. It should be noted that a limitation to this display technique is that reflectors may be masked by larger reflectors near the surface

C-Scan Presentation

The C-scan presentation is a type of presentation that is possible for automated two-dimensional scanning systems that provides a plan-type view of the location and size of test specimen features. The plane of the image is parallel to the scan pattern of the transducer. C-scan presentations are typically produced with an automated data acquisition system, such as a computer controlled immersion scanning system. Typically, a data collection gate is established on the A-scan and the amplitude or the time-of-flight of the signal is recorded at regular intervals as the transducer is scanned over the test piece. The relative signal amplitude or the time-of-flight is displayed as a shade of gray or a color for each of the positions where data was recorded. The C-scan presentation provides an image of the features that reflect and scatter the sound within and on the surfaces of the test piece.

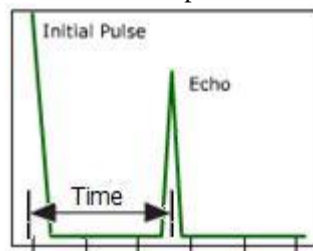


High resolution scans can produce very detailed images. The figure shows two ultrasonic C-scan images of a US quarter. Both images were produced using a pulse-echo technique with the transducer scanned over the head side in an immersion scanning system. For the C-scan image on the top, the gate was set to capture the amplitude of the sound reflecting from the front surface of the quarter. Light areas in the image indicate areas that reflected a greater amount of energy back to the transducer. In the C-scan image on the bottom, the gate was moved to record the intensity of the sound reflecting from the back surface of the coin. The details on the back surface are clearly visible but front surface features are also still visible since the sound energy is affected by these features as it travels through the front surface of the coin.

4.12 MEASUREMENT AND CALIBRATION TECHNIQUES

Normal Beam Inspection

Pulse-echo ultrasonic measurements can determine the location of a discontinuity in a part or structure by accurately measuring the time required for a short ultrasonic pulse generated by a transducer to travel through a thickness of material, reflect from the back or the surface of a discontinuity, and be returned to the transducer. In most applications, this time interval is a few microseconds or less. The two-way transit time measured is divided by two to account for the down-and-back travel path and multiplied by the velocity of sound in the test material. The result is expressed in the well-known relationship:



Where d is the distance from the surface to the discontinuity in the test piece, v is the velocity of sound waves in the material, and t is the measured round-trip transit time. Precision ultrasonic thickness gages usually operate at frequencies between 500 kHz and 100 MHz, by means of piezoelectric transducers that generate bursts of sound waves when excited by electrical pulses. Typically, lower frequencies are used to optimize penetration when measuring thick, highly attenuating or highly scattering materials, while higher frequencies will be recommended to optimize resolution in thinner, non-attenuating, non-scattering materials. It is possible to measure most engineering materials ultrasonically, including metals, plastic, ceramics, composites, epoxies, and glass as well as liquid levels and the thickness of certain biological specimens. On-line or in-process measurement of extruded plastics or rolled metal often is possible, as is measurements of single layers or coatings in multilayer materials.

Angle Beam Inspection

Angle beam transducers and wedges are typically used to introduce a refracted shear wave into the test material. An angled sound path allows the sound beam to come in from the side, thereby improving detectability of flaws in and around welded areas. Angle beam inspection is somehow different than normal beam inspection. In normal beam inspection, the backwall echo is always present on the scope display and when the transducer passes over a discontinuity a new echo will appear between the initial pulse and the backwall echo. However, when scanning a surface using an angle beam transducer there will be no reflected echo on the scope display unless a properly oriented discontinuity or reflector comes into the beam path. If a reflection occurs before the sound waves reach the backwall, the reflection is usually referred to as “*first leg reflection*”. The angular distance (*Sound Path*) to the reflector can be calculated using the same formula used for normal beam transducers

4.13 Inspection of Welded Joints

The most commonly occurring defects in welded joints are porosity, slag inclusions, lack of side-wall fusion, lack of intermediate-pass fusion, lack of root penetration, undercutting, and longitudinal or transverse cracks. With the exception of single gas pores all the listed defects are usually well detectable using ultrasonics. Ultrasonic weld inspections are typically performed using straight beam transducer in conjunction with angle beam transducers.

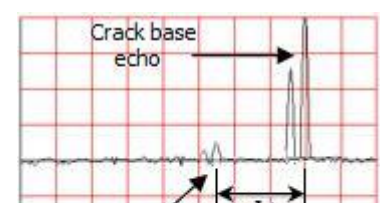
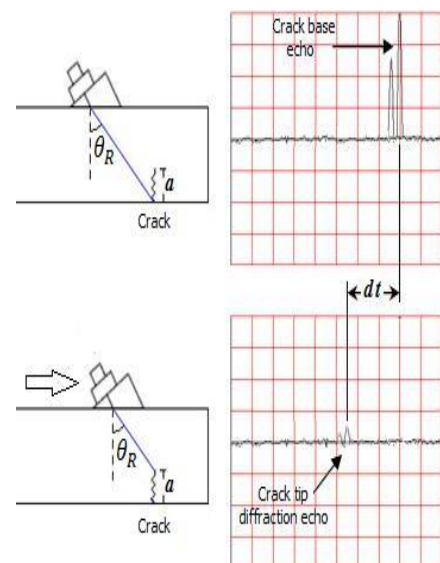
□ A straight beam transducer, producing a longitudinal wave at normal incidence into the test piece, is first used to locate any laminations in or near the heat-affected zone. This is important because an angle beam transducer may not be able to provide a return signal from a laminar flaw.

□ The second step in the inspection involves using an angle beam transducer to inspect the actual weld. This inspection may include the root, sidewall, crown, and heat-affected zones of a weld. The process involves scanning the surface of the material around the weldment with the transducer. This refracted sound wave will bounce off a reflector (*discontinuity*) in the path of the sound beam. To determine the proper scanning area for both sides of the weld, the inspector must calculate the skip distance of the sound beam using the refracted angle and material thickness as: Based on such calculations, the inspector can identify the transducer locations on the surface of the material corresponding to the face, sidewall, and root of the weld. The angle of refraction for the angle beam transducer used for inspection is usually chosen such that $(\theta = 90^\circ - \theta_R)$. Doing so, the second leg of the beam will be normal to the side wall of the weldment such that lack of fusion can be easily detected (*the first leg will also be normal to the other wall*). However, for improving the detectability of the different types of weld discontinuities, it is recommended to repeat the scanning using several transducers having different angles of refraction.

Crack Tip Diffraction

When the geometry of the part is relatively uncomplicated and the orientation of a flaw is well known, the length of a crack can be determined by a technique known as “*crack tip diffraction*”. One common application of the tip diffraction technique is to determine the length of a crack originating from on the backside of a flat plate as shown below. In this case, when an angle beam transducer is scanned over the area of the flaw, an echo appears on the scope display because of the reflection of the sound beam from the base of the crack (*top image*). As the transducer moves, a second, but much weaker, echo appears due to the diffraction of the sound waves at the tip of the crack (*bottom image*). However, since the distance traveled by the diffracted sound wave is less, the second signal appears earlier in time on the scope.

Crack height (a) is a function of the ultrasound shear velocity in the material (v_s), the incident angle (θ_R) and the difference in arrival times between the two signals (Δt). Since the beam angle and the thickness of the material is the same in both measurements, two similar right triangles are formed such that one can be overlaid on the other. A third similar right triangle is made, which is



comprised on the crack, the length and the angle . The variable is really the difference in time but can easily be converted to a distance by dividing the time in half (*to get the one-way travel time*) and multiplying this value by the velocity of the sound in the material. If the material is relatively thick or the crack is relatively short, the crack base echo and the crack tip diffraction echo could appear on the scope display simultaneously (*as seen in the figure*). This can be attributed to the divergence of the sound beam where it becomes wide enough to cover the entire crack length. In such case, though the angle of the beam striking the base of the crack is slightly different than the angle of the beam striking the tip of the crack, the previous equation still holds reasonably accurate and it can be used for estimating the crack length.

Calibration Methods

Calibration refers to the act of evaluating and adjusting the precision and accuracy of measurement equipment. In ultrasonic testing, several forms of calibrations must occur. First, the electronics of the equipment must be calibrated to ensure that they are performing as designed. This operation is usually performed by the equipment manufacturer and will not be discussed further in this material. It is also usually necessary for the operator to perform a "user calibration" of the equipment. This user calibration is necessary because most ultrasonic equipment can be reconfigured for use in a large variety of applications. The user must "calibrate" the system, which includes the equipment settings, the transducer, and the test setup, to validate that the desired level of precision and accuracy are achieved. In ultrasonic testing, reference standards are used to establish a general level of consistency in measurements and to help interpret and quantify the information contained in the received signal. The figure shows some of the most commonly used reference standards for the calibration of ultrasonic equipment. Reference standards are used to validate that the equipment and the setup provide similar results from one day to the next and that similar results are produced by different systems. Reference standards also help the inspector to estimate the size of flaws. In a pulse-echo type setup, signal strength depends on both the size of the flaw and the distance between the flaw and the transducer. The inspector can use a reference standard with an artificially induced flaw of known size and at approximately the same distance away for the transducer to produce a Signal. By comparing the signal from the reference standard to that received from the actual flaw, the inspector can estimate the flaw size.

The material of the reference standard should be the same as the material being inspected and the artificially induced flaw should closely resemble that of the actual flaw. This second requirement is a major limitation of most standard reference samples. Most use drilled holes and notches that do not closely represent real flaws. In most cases the artificially induced defects in reference standards are better reflectors of sound energy (*due to their flatter and smoother surfaces*) and produce indications that are larger than those that a similar sized flaw would produce. Producing more "realistic" defects is cost prohibitive in most cases and, therefore, the inspector can only make an estimate of the flaw size.

Reference standards are mainly used to calibrate instruments prior to performing the inspection and, in general, they are also useful for: Checking the performance of both angle-beam and normal-beam transducers (*sensitivity, resolution, beam spread, etc.*)

- ☐ Determining the sound beam exit point of angle-beam transducers
- ☐ Determining the refracted angle produced
- ☐ Calibrating sound path distance
- ☐ Evaluating instrument performance (*time base, linearity, etc*)

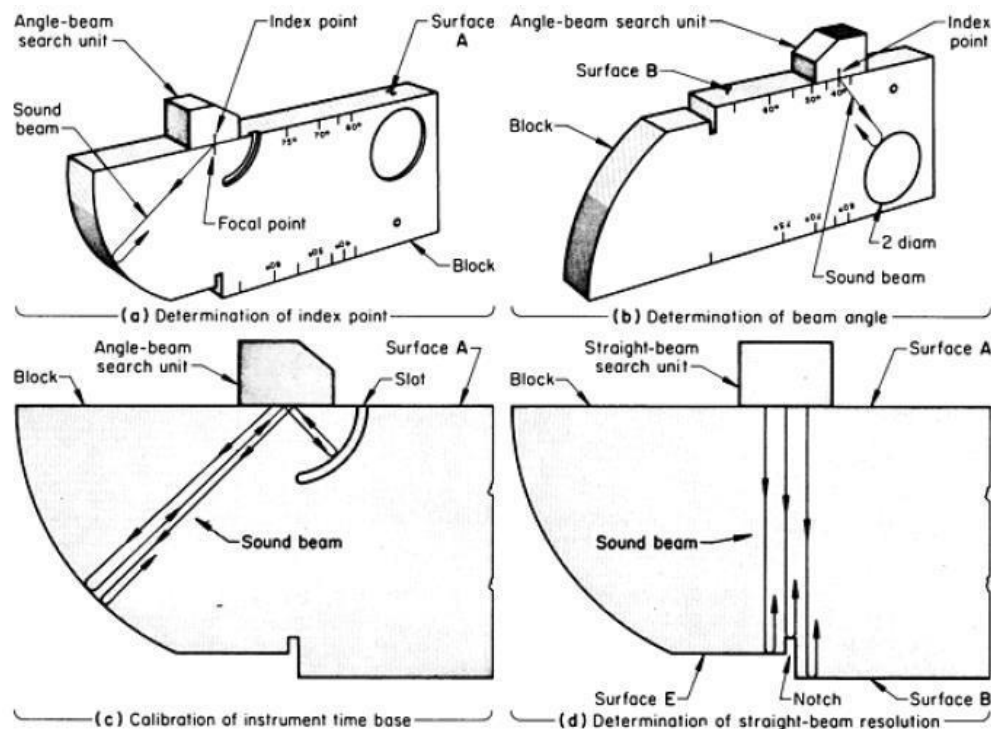
4.14 Introduction to Some of the Common Standards

A wide variety of standard calibration blocks of different designs, sizes and systems of units (*mm or inch*) are available. The type of standard calibration block used is dependent on the NDT application and the form and shape of the object being evaluated. The most commonly used standard calibration blocks are those of the; International Institute of Welding (IIW), American Welding Society (AWS) and

American Society of Testing and Materials (ASTM). Only two of the most commonly used standard calibration blocks are introduced here.

IIW Type US-1 Calibration Block

This block is a general purpose calibration block that can be used for calibrating angle-beam transducers as well as normal beam transducers. The material from which IIW blocks are prepared is specified as killed, open hearth or electric furnace, low-carbon steel in the normalized condition and with a grain size of McQuaid-Ehn No. 8 (*fine grain*). Official IIW blocks are dimensioned in the metric system of units.



ASTM - Miniature Angle-Beam Calibration Block (V2)

The miniature angle-beam block is used in a somewhat similar manner as the as the IIW block, but is smaller and lighter. The miniature angle-beam block is primarily used in the field for checking the characteristics of angle-beam transducers. With the miniature block, beam angle and exit point can be checked for an angle-beam transducer. Both the 25 and 50 mm radius surfaces provide ways for checking the location of the exit point of the transducer and for calibrating the time base of the instrument in terms of metal distance. The small hole provides a reflector for checking beam angle and for setting the instrument gain.

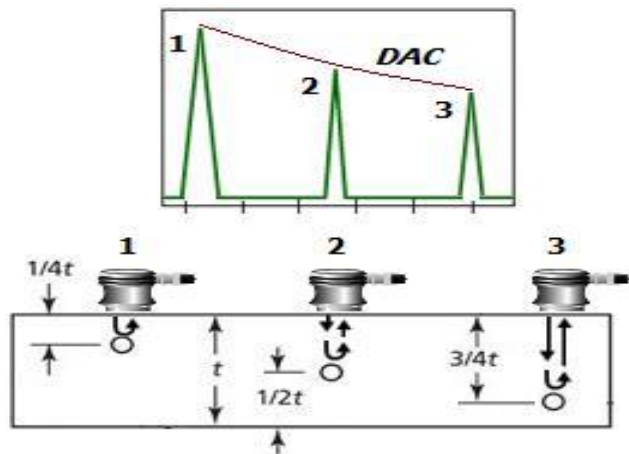
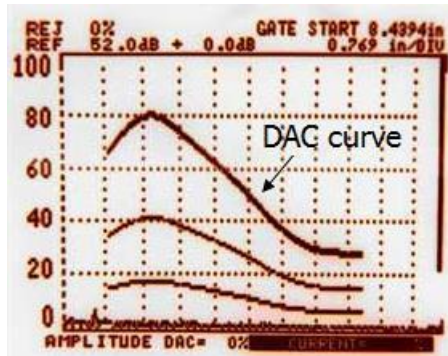
Distance Amplitude Correction (DAC)

Acoustic signals from the same reflecting surface will have different amplitudes at different distances from the transducer. A distance amplitude correction (DAC) curve provides a means of establishing a graphic "reference level sensitivity" as a function of the distance to the reflector (*i.e., time on the A-scan display*). The use of DAC allows signals reflected from similar discontinuities to be evaluated where signal attenuation as a function of depth has been correlated.

DAC will allow for loss in amplitude over material depth (time) to be represented graphically on the A-scan display. Because near field length and beam spread vary according to transducer size and frequency, and materials vary in attenuation and velocity, a DAC curve must be established for each different situation. DAC may be employed in both longitudinal and shear modes of operation as well as either contact or immersion inspection techniques.

A DAC curve is constructed from the peak amplitude responses from reflectors of equal area at different distances in the same material. Reference standards which incorporate side drilled holes (SDH), flat

bottom holes (FBH), or notches whereby the reflectors are located at varying depths are commonly used. A-scan echoes are displayed at their non-electronically compensated height and the peak amplitude of each signal is marked to construct the DAC curve as shown in the figure. It is important to recognize that regardless of the type of reflector used, the size and shape of the reflector must be constant. The same method is used for constructing DAC curves for angle beam transducers, however in that case both the first and second leg reflections can be used for constructing the DAC curve.

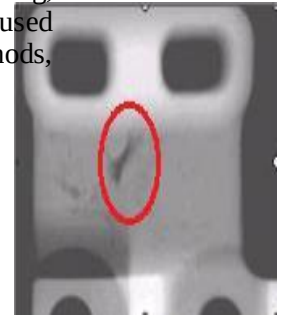


UNIT 5 Radiographic Testing

5.1 Introduction

Radiography is used in a very wide range of applications including medicine, engineering, forensics, security, etc. In NDT, radiography is one of the most important and widely used methods. Radiographic testing (RT) offers a number of advantages over other NDT methods, however, one of its major disadvantages is the health risk associated with the radiation.

In general, RT is method of inspecting materials for hidden flaws by using the ability of short wavelength electromagnetic radiation (high energy photons) to penetrate various materials. The intensity of the radiation that penetrates and passes through the material is either captured by a radiation sensitive film (*Film Radiography*) or by a planer array of radiation sensitive sensors (*Real-time Radiography*). Film radiography is the oldest approach, yet it is still the most widely used in NDT.



5.1.1 Basic Principles

In radiographic testing, the part to be inspected is placed between the radiation source and a piece of radiation sensitive film. The radiation source can either be an X-ray machine or a radioactive source (*Ir-192, Co-60, or in rare cases Cs-137*). The part will stop some of the radiation where thicker and more dense areas will stop more of the radiation. The radiation that passes through the part will expose the film and forms a shadowgraph of the part. The film darkness (*density*) will vary with the amount of radiation reaching the film through the test object where darker areas indicate more exposure (*higher radiation intensity*) and lighter areas indicate less exposure (*lower radiation intensity*).

This variation in the image darkness can be used to determine thickness or composition of material and would also reveal the presence of any flaws or discontinuities inside the material.

5.1.2 Advantages and Disadvantages

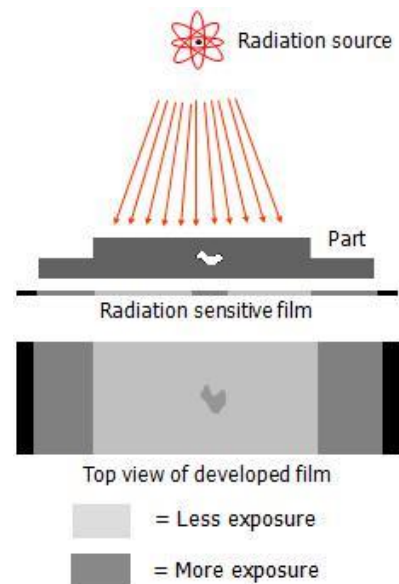
The primary advantages and disadvantages in comparison to other NDT methods are: *Advantages*

- ☐ Both surface and internal discontinuities can be detected.
- ☐ Significant variations in composition can be detected.
- ☐ It has a very few material limitations.
- ☐ Can be used for inspecting hidden areas (*direct access to surface is not required*)
- ☐ Very minimal or no part preparation is required.
- ☐ Permanent test record is obtained.
- ☐ Good portability especially for gamma-ray sources.

Disadvantages

Hazardous to operators and other nearby personnel.

- ☐ High degree of skill and experience is required for exposure and interpretation.
- ☐ The equipment is relatively expensive (*especially for x-ray sources*).
- ☐ The process is generally slow.
- ☐ Highly directional (*sensitive to flaw orientation*).
- ☐ Depth of discontinuity is not indicated.

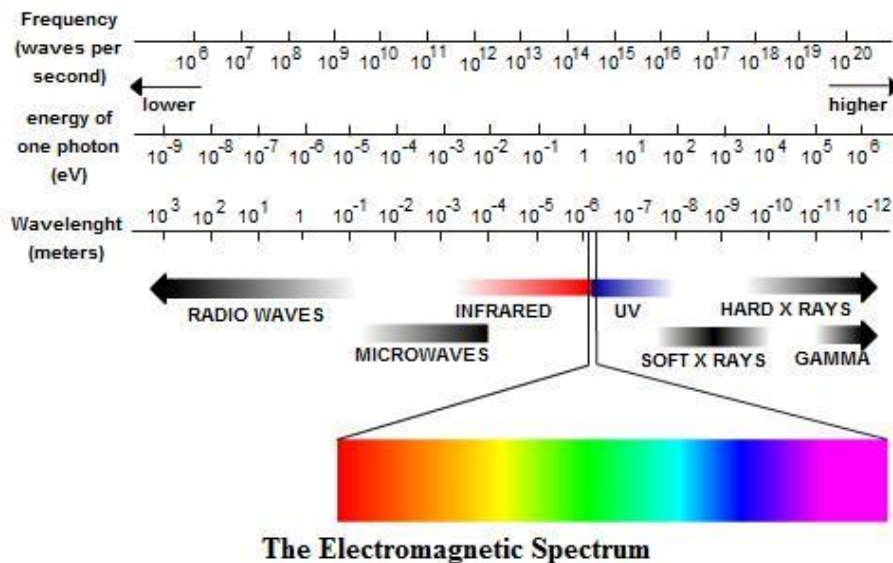


- ☐ It requires a two-sided access to the component.

PHYSICS OF RADIATION

5.1.3 Nature of Penetrating Radiation

Both X-rays and gamma rays are electromagnetic waves and on the electromagnetic spectrum they occupy frequency ranges that are higher than ultraviolet radiation. In terms of frequency, gamma rays generally have higher frequencies than X-rays as seen in the figure. The major distinction between X-rays and gamma rays is the origin where X-rays are usually artificially produced using an X-ray generator and gamma radiation is the product of radioactive materials. Both X-rays and gamma rays are waveforms, as are light rays, microwaves, and radio waves. X-rays and gamma rays cannot be seen, felt, or heard. They possess no charge and no mass and, therefore are not influenced by electrical and magnetic fields and will generally travel in straight lines. However, they can be diffracted (bent) in a manner similar to light.



Electromagnetic radiation act somewhat like a particle at times in that they occur as small “packets” of energy and are referred to as “*photons*”. Each photon contains a certain amount (*or bundle*) of energy, and all electromagnetic radiation consists of these photons. The only difference between the various types of electromagnetic radiation is the amount of energy found in the photons. Due to the short wavelength of X-rays and gamma rays, they have more energy to pass through matter than do the other forms of energy in the electromagnetic spectrum. As they pass through matter, they are scattered and absorbed and the degree of penetration depends on the kind of matter and the energy of the rays.

Properties of X-Rays and Gamma Rays

- ☐ They are not detected by human senses (cannot be seen, heard, felt, etc.).
- ☐ They travel in straight lines at the speed of light.
- ☐ Their paths cannot be changed by electrical or magnetic fields.
- ☐ They can be diffracted, refracted to a small degree at interfaces between two different materials, and in some cases be reflected.
- ☐ They pass through matter until they have a chance to encounter with an atomic particle.
- ☐ Their degree of penetration depends on their energy and the matter they are traveling through.
- ☐ They have enough energy to ionize matter and can damage or destroy living cells.

X Radiation

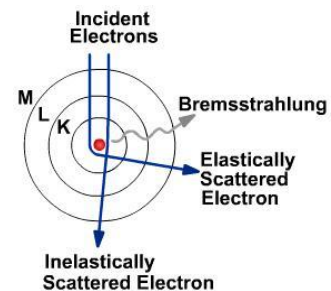
X-rays are just like any other kind of electromagnetic radiation. They can be produced in packets of energy called photons, just like light. There are two different atomic processes that can produce X-ray photons. One is called *Bremsstrahlung* (a German term meaning “braking radiation”) and the other is called *K-shell emission*. They can both occur in the heavy atoms of tungsten which is often the material chosen for the target or anode of the X-ray tube.

Both ways of making X-rays involve a change in the state of electrons. However, Bremsstrahlung is easier to understand using the classical idea that radiation is emitted when the velocity of the electron shot at the tungsten target changes. The negatively charged electron slows down after swinging around the nucleus of a positively charged tungsten atom and this energy loss produces X-radiation. Electrons are scattered elastically or inelastically by the positively charged nucleus. The inelastically scattered electron loses energy, and thus produces X-ray photon, while the elastically scattered electrons generally change their direction significantly but without losing much of their energy.

5.2 Bremsstrahlung Radiation

X-ray tubes produce X-ray photons by accelerating a stream of electrons to energies of several hundred kiloelectronvolts with velocities of several hundred kilometers per hour and colliding them into a heavy target material. The abrupt acceleration of the charged particles (electrons) produces Bremsstrahlung photons. X-ray radiation with a continuous spectrum of energies is produced with a range from a few keV to a maximum of the energy of the electron beam.

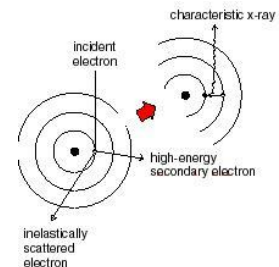
The Bremsstrahlung photons generated within the target material are attenuated as they pass through, typically, 50 microns of target material. The beam is further attenuated by the aluminum or beryllium vacuum window. The results are the elimination of the low energy photons, 1 keV through 15 keV, and a significant reduction in the portion of the spectrum from 15 keV through 50 keV. The spectrum from an X-ray tube is further modified by the filtration caused by the selection of filters used in the setup.



K-shell Emission Radiation

Remember that atoms have their electrons arranged in closed “shells” of different energies. The K-shell is the lowest energy state of an atom. An incoming electron can give a K-shell electron enough energy to knock it out of its energy state. About 0.1% of the electrons produce K-shell vacancies; most produce heat. Then, a tungsten electron of higher energy (from an outer shell) can fall into the K-shell. The energy lost by the falling electron shows up as an emitted X-ray photon. Meanwhile,

higher energy electrons fall into the vacated energy state in the outer shell, and so on. After losing an electron, an atom remains ionized for a very short time (about 10^{-14} second) and thus an atom can be repeatedly ionized by the incident electrons which arrive about every 10^{-12} second. Generally, K-shell emission produces higher-intensity X-rays than Bremsstrahlung, and the X-ray photon comes out at a single wavelength



Gamma Radiation

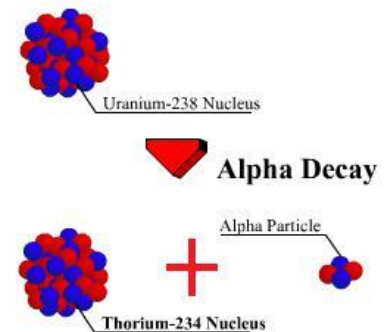
Gamma radiation is one of the three types of natural radioactivity. Gamma rays are electromagnetic radiation just like X-rays. The other two types of natural radioactivity are alpha and beta radiation, which are in the form of particles. Gamma rays are the most energetic form of electromagnetic radiation. Gamma radiation is the product of radioactive atoms. Depending upon the ratio of neutrons to protons within its nucleus, an isotope of a particular element may be stable or unstable. When the binding energy is not strong enough to hold the nucleus of an atom together, the atom is said to be unstable. Atoms with unstable nuclei are constantly changing as a result of the imbalance of energy within the nucleus. Over time, the nuclei of unstable isotopes spontaneously disintegrate, or transform, in a process known as “radioactive decay” and such material is called “radioactive material”.

5.4 Types of Radiation Produced by Radioactive Decay

When an atom undergoes radioactive decay, it emits one or more forms of high speed subatomic particles ejected from the nucleus or electromagnetic radiation (gamma-rays) emitted by either the nucleus or orbital electrons.

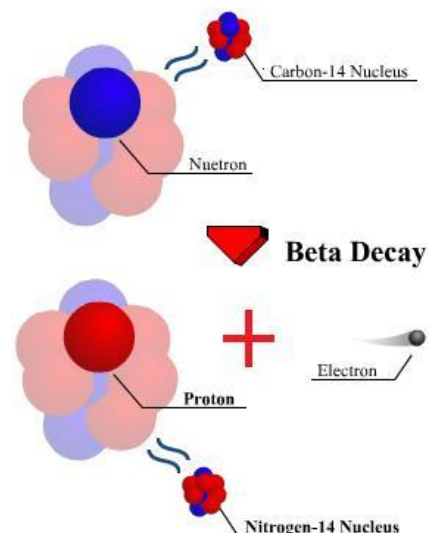
Alpha Particles

Certain radioactive materials of high atomic mass (such as *Ra-226*, *U-238*, *Pu-239*), decay by the emission of alpha particles. These alpha particles are tightly bound units of two neutrons and two protons each (*He-4 nucleus*) and have a positive charge. Emission of an alpha particle from the nucleus results in a decrease of two units of atomic number (*Z*) and four units of mass number (*A*). Alpha particles are emitted with discrete energies characteristic of the particular transformation from which they originate. All alpha particles from a particular radionuclide transformation will have identical energies.



Beta Particles

A nucleus with an unstable ratio of neutrons to protons may decay through the emission of a high speed electron called a beta particle. In beta decay a neutron will split into a positively charged proton and a negatively charged electron. This results in a net change of one unit of atomic number (*Z*) and no change in the mass number (*A*). Beta particles have a negative charge and the beta particles emitted by a specific radioactive material will range in energy from near zero up to a maximum value, which is characteristic of the particular transformation.



Gamma-rays

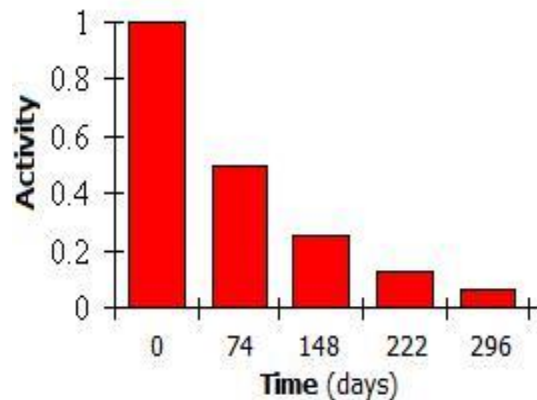
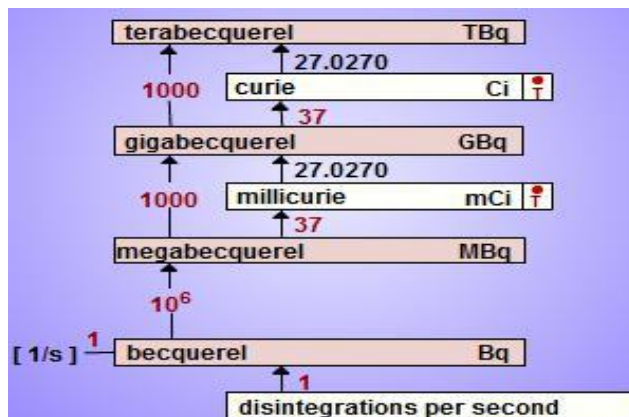
A nucleus which is in an excited state (*unstable nucleus*) may emit one or more photons of discrete energies. The emission of gamma rays does not alter the number of protons or neutrons in the nucleus but instead has the effect of moving the nucleus from a *higher to a lower energy state (unstable to stable)*. Gamma ray emission frequently follows beta decay, alpha decay, and other nuclear decay processes.

5.5 Activity (of Radioactive Materials)

The quantity which expresses the radiation producing potential of a given amount of radioactive material is called “Activity”. The *Curie (Ci)* was originally defined as that amount of any radioactive material that disintegrates at the same rate as one gram of pure radium. The International System (SI)

unit for activity is the *Becquerel (Bq)*, which is that quantity of radioactive material in which one atom is transformed per second. The radioactivity of a given amount of radioactive material does not depend upon the mass of material present. For example, two one-curie sources of the same radioactive material might have very different masses depending upon the relative proportion of non-radioactive atoms present in each source.

The concentration of radioactivity, or the relationship between the mass of radioactive material and the activity, is called “*specific activity*”. Specific activity is expressed as the number of *Curies* or *Becquerels* per unit mass or volume. Each gram of Cobalt-60 will contain approximately 50 Ci. Iridium-192 will contain 350 Ci for every gram of material. The higher specific activity of iridium results in physically smaller sources. This allows technicians to place the source in closer proximity to the film while maintaining the sharpness of the radiograph



Isotope Decay Rate (Half-Life)

Each radioactive material decays at its own unique rate which cannot be altered by any chemical or physical process. A useful measure of this rate is the “*half-life*” of the radioactivity. Half-life is defined as the time required for the activity of any particular radionuclide to decrease to one-half of its initial value. In other words one-half of the atoms have reverted to a more stable state material. Half-lives of radioactive materials range from microseconds to billions of years. Half-life of two widely used industrial isotopes are; 74 days for Iridium-192, and 5.3 years for Cobalt-60. In order to find the remaining activity of a certain material with a known half-life value after a certain period of time, the following formula may be used. The formula calculates the decay fraction (or the remaining fraction of the initial activity) as:

$$f_D = \frac{A}{A_0} = (0.5)^{\frac{t}{L_H}}$$

Where;

f_D : decay fraction (i.e., remaining fraction of the initial activity)

L_H : Half-Life value (hours, days, years, etc.)

t : Elapsed time (hours, days, years, etc.)

Radiation Energy, Intensity and Exposure

Different radioactive materials and X-ray generators produce radiation at different energy levels and at different rates. It is important to understand the terms used to describe the energy and intensity of the radiation.

Radiation Energy

The energy of the radiation is responsible for its ability to penetrate matter. Higher energy radiation can penetrate more and higher density matter than low energy radiation. The energy of ionizing radiation is measured in *electronvolts (eV)*. One electronvolt is an extremely small amount of energy so it is common to use kiloelectronvolts (*keV*) and megaelectronvolt (*MeV*). An electronvolt is a measure of energy, which is different from a volt which is a measure of the electrical potential between two positions. Specifically, an electronvolt is the kinetic energy gained by an electron passing through a potential difference of one volt. X-ray generators have a control to adjust the radiation energy, *keV* (or *kV*). The energy of a radioisotope is a characteristic of the atomic structure of the material. Consider, for example, Iridium-192 and Cobalt-60, which are two of the more common industrial Gamma ray sources. These isotopes emit radiation in two or three discrete wavelengths. Cobalt-60 will emit *1.17 and 1.33 MeV* gamma rays, and Iridium-192 will emit *0.31, 0.47, and 0.60 MeV* gamma rays. It can be seen from these values that the energy of radiation coming from Co-60 is more than twice the energy of the radiation coming from the Ir-192. From a radiation safety point of view, this difference in energy is important because the Co-60 has more material penetrating power and, therefore, is more dangerous and requires more shielding

Intensity and Exposure

Radiation intensity is the amount of energy passing through a given area that is perpendicular to the direction of radiation travel in a given unit of time. One way to measure the intensity of X-rays or gamma rays is to measure the amount of ionization they cause in air. The amount of ionization in air produced by the radiation is called the exposure. Exposure is expressed in terms of a scientific unit called a *Roentgen (R)*. The unit roentgen is equal to the amount of radiation that ionizes 1 cm^3 of dry air (at 0°C and standard atmospheric pressure) to one electrostatic unit of charge, of either sign. Most portable radiation detection safety devices used by radiographers measure exposure and present the reading in terms of *Roentgens* or *Roentgens/hour*, which is known as the “dose rate”.

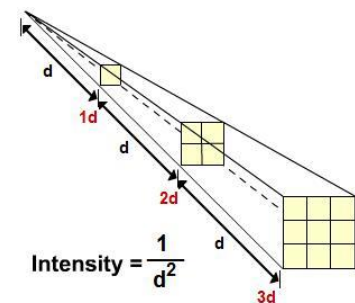
Ionization

As penetrating radiation moves from point to point in matter, it loses its energy through various interactions with the atoms it encounters. The rate at which this energy loss occurs depends upon the type and energy of the radiation and the density and atomic composition of the matter through which it is passing. The various types of penetrating radiation impart their energy to matter primarily through excitation and ionization of orbital electrons. The term “*excitation*” is used to describe an interaction where electrons acquire energy from a passing charged particle but are not removed completely from their atom. Excited electrons may subsequently emit energy in the form of X-rays during the process of returning to a lower energy state. The term “*ionization*” refers to the complete removal of an electron from an atom following the transfer of energy from a passing charged particle. Because of their double charge and relatively slow velocity, alpha particles have a relatively short range in matter (a few centimeters in air and only fractions of a millimeter in tissue). Beta particles have, generally, a greater range.

Since they have no charge, gamma-rays and X-rays proceed through matter until there is a chance of interaction with a particle. If the particle is an electron, it may receive enough energy to be ionized, whereupon it causes further ionization by direct interactions with other electrons. As a result, gamma-rays and X-rays can cause the liberation of electrons deep inside a medium. As a result, a given gamma or X-ray has a definite probability of passing through any medium of any depth.

5.6 Newton's Inverse Square Law

Any point source which spreads its influence equally in all directions without a limit to its range will obey the inverse square law. This comes from strictly geometrical considerations. The intensity of the influence at any given distance (d) is the source strength divided by the area of a sphere having a radius equal to the distance (d). Being strictly geometric in its origin, the inverse square law applies to diverse phenomena. Point sources of gravitational force, electric field, light, sound, and radiation obey the inverse square law.



As one of the fields which obey the general inverse square law, the intensity of the radiation received from a point radiation source can be characterized by the diagram above. The relation between intensity and distance according to the inverse square law can be expressed as:

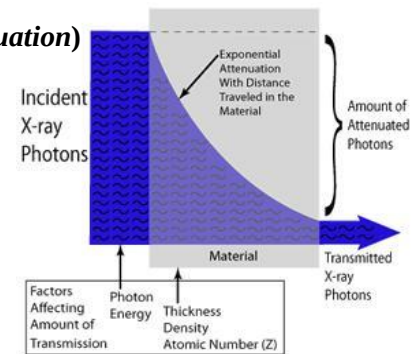
$$I_1 d_1^2 = I_2 d_2^2$$

Where I_1 & I_2 are the intensities at distances d_1 & d_2 from the source, respectively.

All measures of exposure or dose rate will drop off by the inverse square law. For example, if the received dose of radiation is 100 mR/hr at 2 cm from a source, it will be 0.01 mR/hr at 2 m .

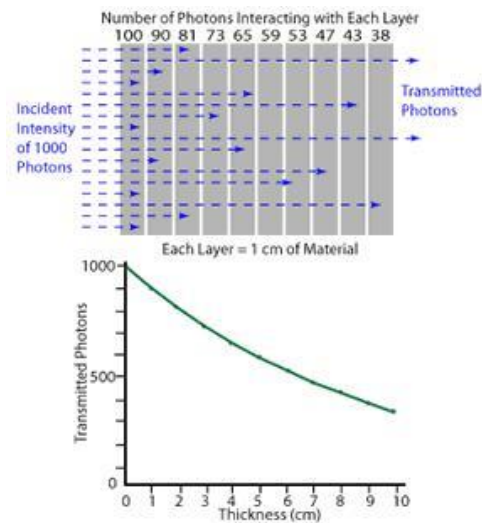
5.7 Interaction between Penetrating Radiation and Matter (Attenuation)

When X-rays or gamma rays are directed into an object, some of the photons interact with the particles of the matter and their energy can be absorbed or scattered. This absorption and scattering is called “*Attenuation*”. Other photons travel completely through the object without interacting with any of the material's particles. The number of photons transmitted through a material depends on the thickness, density and atomic number of the material, and the energy of the individual photons.



Even when they have the same energy, photons travel different distances within a material simply based on the probability of their encounter with one or more of the particles of the matter and the type of

encounter that occurs. Since the probability of an encounter increases with the distance traveled, the number of photons reaching a specific point within the matter decreases exponentially with distance traveled. As shown in the graphic to the right, if 1000 photons are aimed at ten 1 cm layers of a material and there is a 10% chance of a photon being attenuated in this layer, then there will be 100 photons attenuated. This leaves 900 photos to travel into the next layer where 10% of these photos will be attenuated. By continuing this progression, the exponential shape of the curve becomes apparent.



The formula that describes this curve is:

$$I = I_0 e^{-\mu x}$$

I

Where;

I_0 : initial (*unattenuated*) intensity

μ : linear attenuation coefficient per unit distance

x: distance traveled through the matter

Linear and Mass Attenuation Coefficients

The “*linear attenuation coefficient*” () describes the fraction of a beam of X-rays or gamma rays that is absorbed or scattered per unit thickness of the absorber (*10% per cm thickness for the previous example*). Using the transmitted intensity equation above, linear attenuation coefficients can be used to make a number of calculations. These include:

- ☐ The intensity of the energy transmitted through a material when the incident X-ray intensity, the material and the material thickness are known.
- ☐ The intensity of the incident X-ray energy when the transmitted X-ray intensity, material, and material thickness are known.
- ☐ The thickness of the material when the incident and transmitted intensity, and the material are known.
- ☐ The material can be determined from the value of when the incident and transmitted intensity, and the material thickness are known.

Linear attenuation coefficients can sometimes be found in the literature. However, it is often easier to locate attenuation data in terms of the “*mass attenuation coefficient*”

Sources of Attenuation

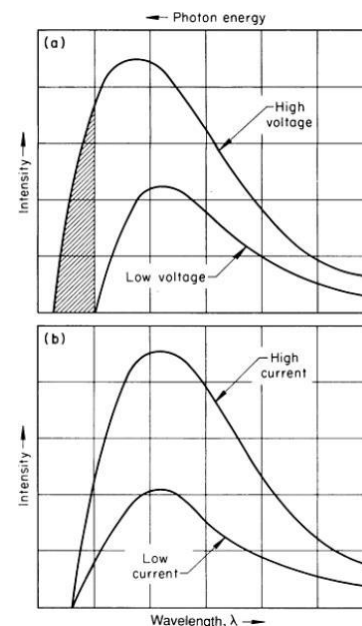
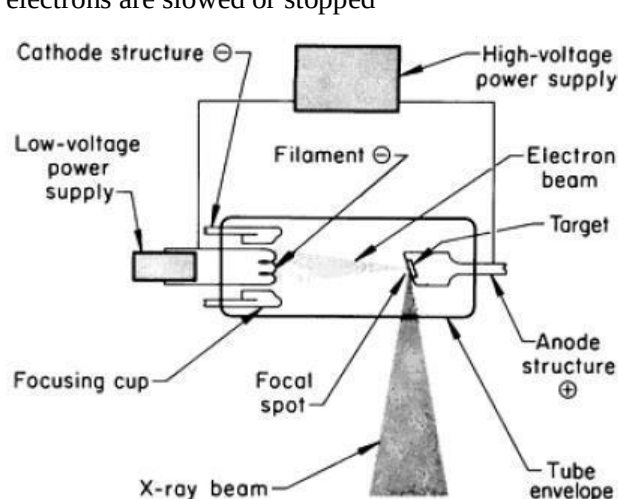
The attenuation that results due to the interaction between penetrating radiation and matter is not a simple process. A single interaction event between a primary X-ray photon and a particle of matter does not usually result in the photon changing to some other form of energy and effectively disappearing. Several interaction events are usually involved and the total attenuation is the sum of the attenuation due to different types of interactions. These interactions include the photoelectric effect, scattering, and pair production.

- **Photoelectric (PE) absorption** of X-rays occurs when the X-ray photon is absorbed, resulting in the ejection of electrons from the outer shell of the atom, and hence the ionization of the atom. Subsequently, the ionized atom returns to the neutral state with the emission of an X-ray characteristic of the atom. This subsequent emission of lower energy photons is generally absorbed and does not contribute to (or hinder) the image making process. Photoelectron absorption is the dominant process for X-ray absorption up to energies of about 500 keV. Photoelectric absorption is also dominant for atoms of high atomic numbers.
- **Compton scattering (C)** occurs when the incident X-ray photon is deflected from its original path by an interaction with an electron. The electron gains energy and is ejected from its orbital position. The X-ray photon loses energy due to the interaction but continues to travel through the material along an altered path. Since the scattered X-ray photon has less energy, it, therefore, has a longer wavelength than the incident photon.
- **Pair production (PP)** can occur when the X-ray photon energy is greater than 1.02 MeV, but really only becomes significant at energies around 10 MeV. Pair production occurs when an electron and positron are created with the annihilation of the X-ray photon. Positrons are very short lived and disappear (positron annihilation) with the formation of two photons of 0.51 MeV energy. Pair production is of particular importance when high-energy photons pass through materials of a high atomic number.

5.8 EQUIPMENT & MATERIALS

X-ray Generators

The major components of an X-ray generator are the tube, the high voltage generator, the control console, and the cooling system. As discussed earlier in this material, X-rays are generated by directing a stream of high speed electrons at a target material such as tungsten, which has a high atomic number. When the electrons are slowed or stopped



by the interaction with the atomic particles of the target, X-radiation is produced. This is accomplished in an X-ray tube such as the one shown in the figure. The tube cathode (*filament*) is heated with a low-voltage current of a few amps. The filament heats up and the electrons in the wire become loosely held. A large electrical potential is created between the cathode and the anode by the high-voltage generator. Electrons that break free of the cathode are strongly attracted to the anode target. The stream of electrons between the cathode and the anode is the tube current. The tube current is measured in

milliamps and is controlled by regulating the low-voltage heating current applied to the cathode. The higher the temperature of the filament, the larger the number of electrons that leave the cathode and travel to the anode. The milliamp or current setting on the control console regulates the filament temperature, which relates to the intensity of the X-ray output.

The high-voltage between the cathode and the anode affects the speed at which the electrons travel and strike the anode. The higher the kilovoltage, the more speed and, therefore, energy the electrons have when they strike the anode. Electrons striking with more energy result in X-rays with more penetrating power. The high-voltage potential is measured in kilovolts, and this is controlled with the voltage or kilovoltage control on the control console. An increase in the kilovoltage will also result in an increase in the intensity of the radiation. The figure shows the spectrum of the radiated X-rays associated with the voltage and current settings. The top figure shows that increasing the kV increases both the energy of X-rays and also increases the intensity of radiation (*number of photons*). Increasing the current, on the other hand, only increases the intensity without shifting the spectrum.

A focusing cup is used to concentrate the stream of electrons to a small area of the target called the “*focal spot*”. The focal spot size is an important factor in the system's ability to produce a sharp image. Much of the energy applied to the tube is transformed into heat at the focal spot of the anode. As mentioned above, the anode target is commonly made from tungsten, which has a high melting point in addition to a high atomic number. However, cooling of the anode by active or passive means is necessary. Water or oil re-circulating systems are often used to cool tubes. Some low power tubes are cooled simply with the use of thermally conductive materials and heat radiating fins.

In order to prevent the cathode from burning up and to prevent arcing between the anode and the cathode, all of the oxygen is removed from the tube by pulling a vacuum. Some systems have external vacuum pumps to remove any oxygen that may have leaked into the tube. However, most industrial X-ray tubes simply require a warm-up procedure to be followed. This warm-up procedure carefully raises the tube current and voltage to slowly burn any of the available oxygen before the tube is operated at high power.

In addition, X-ray generators usually have a filter along the beam path (*placed at or near the x-ray port*). Filters consist of a thin sheet of material (*often high atomic number materials such as lead, copper, or brass*) placed in the useful beam to modify the spatial distribution of the beam. Filtration is required to absorb the lower-energy X-ray photons emitted by the tube before they reach the target in order to produce a cleaner image (*since lower energy X-ray photons tend to scatter more*).

The other important component of an X-ray generating system is the control console. Consoles typically have a keyed lock to prevent unauthorized use of the system. They will have a button to start the generation of X-rays and a button to manually stop the generation of X-rays. The three main adjustable controls regulate the tube voltage in *kilovolts*, the tube amperage in *milliamps*, and the exposure time in *minutes and seconds*. Some systems also have a switch to change the focal spot size of the tube

Radio Isotope (Gamma-ray) Sources

Manmade radioactive sources are produced by introducing an extra neutron to atoms of the source material. As the material gets rid of the neutron, energy is released in the form of gamma rays. Two of the most

common industrial gamma-ray sources for industrial radiography are Iridium-192 and Cobalt-60. In comparison to an X-ray generator, Cobalt-60 produces energies comparable to a 1.25 MV X-ray system and Iridium-192 to a 460 kV X-ray system. These high energies make it possible to penetrate thick materials with a relatively short exposure time. This and the fact that sources are very portable are the main reasons that gamma sources are widely used for field radiography. Of course, the disadvantage of a radioactive source is that it can never be turned off and safely managing the source is a constant responsibility.

Physical size of isotope materials varies between manufacturers, but generally an isotope material is a pellet that measures 1.5 mm x 1.5 mm. Depending on the level of activity desired, a pellet or pellets are loaded into a stainless steel capsule and sealed by welding. The capsule is attached to short flexible cable called a pigtail.

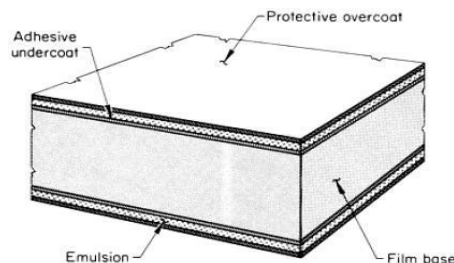
The source capsule and the pigtail are housed in a shielding device referred to as a exposure device or camera. Depleted uranium is often used as a shielding material for sources. The exposure device for Iridium-192 and Cobalt-60 sources will contain 22 kg and 225 kg of shielding materials, respectively. Cobalt cameras are often fixed to a trailer and transported to and from inspection sites. When the source is not being used to make an exposure, it is locked inside the exposure device.

To make a radiographic exposure, a crank-out mechanism and a guide tube are attached to opposite ends of the exposure device. The guide tube often has a collimator (*usually made of tungsten*) at the end to shield the radiation except in the direction necessary to make the exposure. The end of the guide tube is secured in the location where the radiation source needs to be to produce the radiograph. The crank-out cable is stretched as far as possible to put as much distance as possible between the exposure device and the radiographer. To make the exposure, the radiographer quickly cranks the source out of the exposure device and into position in the collimator at the end of the guide tube. At the end of the exposure time, the source is cranked back into the exposure device. There is a series of safety procedures, which include several radiation surveys, that must be accomplished when making an exposure with a gamma source.

5.9 Radiographic Film

X-ray films for general radiography basically consist of an emulsion-gelatin containing radiation-sensitive silver halide crystals (*such as silver bromide or silver chloride*). The emulsion is usually coated on both sides of a flexible, transparent, blue-tinted base in layers about 0.012 mm thick. An adhesive undercoat fastens the emulsion to the film base and a very thin but tough coating covers the emulsion to protect it against minor abrasion. The typical total thickness of the X-ray film is approximately 0.23 mm.

Though films are made to be sensitive for X-ray or gamma-ray, yet they are also sensitive to visible light. When X-rays, gamma-rays, or light strike the film, some of the halogen atoms are liberated from the silver halide crystal and thus leaving the silver atoms alone. This change is of such a small nature that it cannot be detected by ordinary physical methods and is called a “*latent (hidden) image*”. When the film is exposed to a chemical solution (*developer*) the reaction results in the formation of black, metallic silver.



Film Selection

Selecting the proper film and developing the optimal radiographic technique for a particular component depends on a number of different factors;

- ☐ Composition, shape, and size of the part being examined and, in some cases, its weight and location.
- ☐ Type of radiation used, whether X-rays from an X-ray generator or gamma rays from a radioactive source.
- ☐ Kilovoltage available with the X-ray equipment or the intensity of the gamma radiation.
- ☐ Relative importance of high radiographic detail or quick and economical results.

Film Packaging

Radiographic film can be purchased in a number of different packaging options and they are available in a variety of sizes. The most basic form is as individual sheets in a box. In preparation for use, each sheet must be loaded into a cassette or film holder in a darkroom to protect it from exposure to light. Industrial X-ray films are also available in a form in which each sheet is enclosed in a light-tight envelope. The film can be exposed from either side without removing it from the protective packaging. A rip strip makes it easy to remove the film in the darkroom for processing. Packaged film is also available in the form of rolls where that allows the radiographer to cut the film to any length. The ends of the packaging are sealed with electrical tape in the darkroom. In applications such as the radiography of circumferential welds and the examination of long joints on an aircraft fuselage, long lengths of film offer great economic advantage.

Film Handling

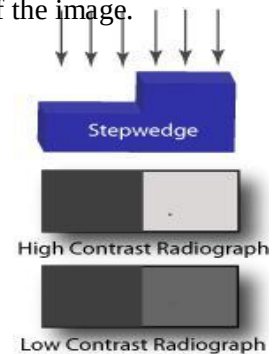
X-ray film should always be handled carefully to avoid physical strains, such as pressure, creasing, buckling, friction, etc. Whenever films are loaded in semi-flexible holders and external clamping devices are used, care should be taken to be sure pressure is uniform. Marks resulting from contact with fingers that are moist or contaminated with processing chemicals, as well as crimp marks, are avoided if large films are always grasped by the edges and allowed to hang free. Use of envelope-packed films avoids many of these problems until the envelope is opened for processing.

5.10 RADIOGRAPHY CONSIDERATIONS & TECHNIQUES

Radiographic Sensitivity

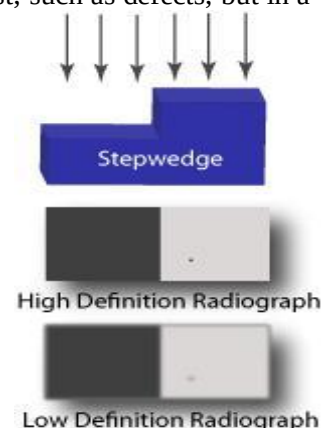
The usual objective in radiography is to produce an image showing the highest amount of detail possible. This requires careful control of a number of different variables that can affect image quality. Radiographic sensitivity is a measure of the quality of an image in terms of the smallest detail or discontinuity that may be detected. Radiographic sensitivity is dependant on the contrast and the definition of the image.

Radiographic contrast is the degree of density (darkness) difference between two areas on a radiograph. Contrast makes it easier to distinguish features of interest, such as defects, from the surrounding area. The image to the right shows two radiographs of the same stepwedge. The upper radiograph has a high level of contrast and the lower radiograph has a lower level of contrast. While they are both imaging the same change in thickness, the high contrast image uses a larger change in radiographic density to show this change. In each of the two radiographs, there is a small dot, which is of equal density in both radiographs. It is much easier to see in the high contrast radiograph.



Radiographic definition is the abruptness of change in going from one area of a given radiographic density to another. Like contrast, definition also makes it easier to see features of interest, such as defects, but in a totally different

way. In the image to the right, the upper radiograph has a high level of definition and the lower radiograph has a lower level of definition. In the high definition radiograph it can be seen that a change in the thickness of the stepwedge translates to an abrupt change in radiographic density. It can be seen that the details, particularly the small dot, are much easier to see in the high definition radiograph. It can be said that a faithful visual reproduction of the stepwedge was produced. In the lower image, the radiographic setup did not produce a faithful visual reproduction. The edge line between the steps is blurred. This is evidenced by the gradual transition between the high and low density areas on the radiograph.



Radiographic “Image” Density

After taking a radiographic image of a part and processing the film, the resulting darkness of the film will vary according to the amount of radiation that has reached the film through the test object. As mentioned earlier, the darker areas indicate more exposure and lighter areas indicate less exposure. The processed film (*or image*) is usually viewed by placing it in front of a screen providing white light illumination of uniform intensity such that the light is transmitted through the film such that the image can be clearly seen. The term “radiographic density” is a measure of the degree of film darkening (*darkness of the image*). Technically it should be called “transmitted density” when associated with transparent-base film since it is a measure of the light transmitted through the film. Radiographic density is the logarithm of two measurements: the intensity of light incident on the film () and the intensity of light transmitted through the film (). This ratio is the inverse of transmittance.

$$\text{Density} = \log (I/I_0)$$

Similar to the decibel, using the log of the ratio allows ratios of significantly different sizes to be described using easy to work with numbers. The following table shows numeric examples of the relationship between the amount of transmitted light and the calculated film density.

Transmittance (I_t/I_0)	Transmittance (%)	Inverse of Transmittance (I_0/I_t)	Density ($\log(I_0/I_t)$)
1.0	100%	1	0
0.1	10%	10	1
0.01	1%	100	2
0.001	0.1%	1000	3
0.0001	0.01%	10000	4

From the table, it can be seen that a density reading of 2.0 is the result of only one percent of the incident light making it through the film. At a density of 4.0 only 0.01% of transmitted light reaches the far side of the film. Industrial codes and standards typically require a radiograph to have a density between 2.0 and 4.0 for acceptable viewing with common film viewers. Above 4.0, extremely bright viewing lights is necessary for evaluation. Film density is measured with a densitometer which simply measures the amount of light transmitted through a piece of film using a photovoltaic sensor.

Secondary (Scatter) Radiation Control

Secondary or scatter radiation must often be taken into consideration when producing a radiograph. The scattered photons create a loss of contrast and definition. Often, secondary radiation is thought of as radiation striking the film reflected from an object in the immediate area, such as a wall, or from the table or floor where the part is resting. Control of side scatter can be achieved by moving objects in the room away from the film, moving the X-ray tube to the center of the vault, or placing a collimator at the exit port, thus reducing the diverging radiation surrounding the central beam. When scattered radiation comes from objects behind the film, it is often called “backscatter”. Industry codes and standards often require that a lead letter “B” be placed on the back of the cassette to verify the control of backscatter. If the letter “B” shows as a “ghost” image on the film, a significant amount of backscatter radiation is reaching the film. The image of the “B” is often very nondistinct as shown in the image to the right. The arrow points to the area of backscatter radiation from the lead “B” located on the back side of the film. The control of backscatter radiation is achieved by backing the film in the cassette with a sheet of lead that is at least 0.25 mm thick such that the sheet will be behind the film when it is exposed. It is a common practice in industry to place thin sheets of lead (*called “lead screens”*) in front and behind the film (0.125 mm thick in front and 0.25mm thick behind).

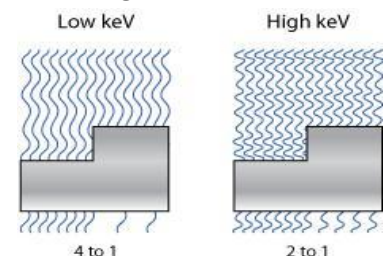
Radiographic Contrast

As mentioned previously, radiographic contrast describes the differences in photographic density in a radiograph. The contrast between different parts of the image is what forms the image and the greater the contrast, the more visible features become. Radiographic contrast has two main contributors; subject contrast and film (or detector) contrast.

Subject Contrast

Subject contrast is the ratio of radiation intensities transmitted through different areas of the component being evaluated. It is dependant on the absorption differences in the component, the wavelength of the primary radiation, and intensity and distribution of secondary radiation due to scattering. It should be no surprise that absorption differences within the subject will affect the level of contrast in a radiograph. The larger the difference in thickness or density between two areas of the subject, the larger the difference in radiographic density or contrast. However, it is also possible to radiograph a

particular subject and produce two radiographs having entirely different contrast levels. Generating X-rays using a low kilovoltage will generally result in a radiograph with high contrast. This occurs because low energy radiation is more easily attenuated. Therefore, the ratio of photons that are transmitted through a thick and thin area will be greater with low energy radiation.

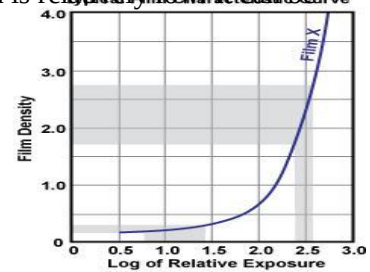


There is a tradeoff, however. Generally, as contrast sensitivity increases, the latitude of the radiograph decreases. Radiographic latitude refers to the range of material thickness that can be imaged. This means that more areas of different thicknesses will be visible in the image. Therefore, the goal is to balance radiographic contrast and latitude so that there is enough contrast to identify the features of interest but also to make sure the latitude is great enough so that all areas of interest can be inspected with one radiograph. In thick parts with a large range of thicknesses, multiple radiographs will likely be necessary to get the necessary density levels in all areas.

Film Contrast

Film contrast refers to density differences that result due to the type of film being used, how it was exposed, and how it was processed. Since there are other detectors besides film, this could be called detector contrast, but the focus here will be on film. Exposing a film to produce higher film densities will generally increase the contrast in the radiograph. A typical film characteristic curve, which shows how a film responds to different amounts of radiation exposure, is shown in the figure. From the shape of the curves, it can be seen that when the film has not seen many photon interactions (*which will result in a low film density*) the slope of the curve is low. In this region of the curve, it takes a large change in exposure to produce a small change in film density. Therefore, the sensitivity of the film is relatively low. It can be

seen that changing the log of the relative exposure from 0.75 to 1.4 only changes the film density from 0.20 to about 0.30. However, at film densities above 2.0, the slope of the characteristic curve for most films is at its maximum. In this region of the curve, a relatively small change in exposure will result in a relatively large change in film density. For example, changing the log of relative exposure from 2.4 to 2.6 would change the film density from 1.75 to 2.75. Therefore, the sensitivity of the film is high in this region of the curve. In general, the highest overall film density that can be conveniently viewed or digitized will have the highest level of contrast and contain the most useful information.



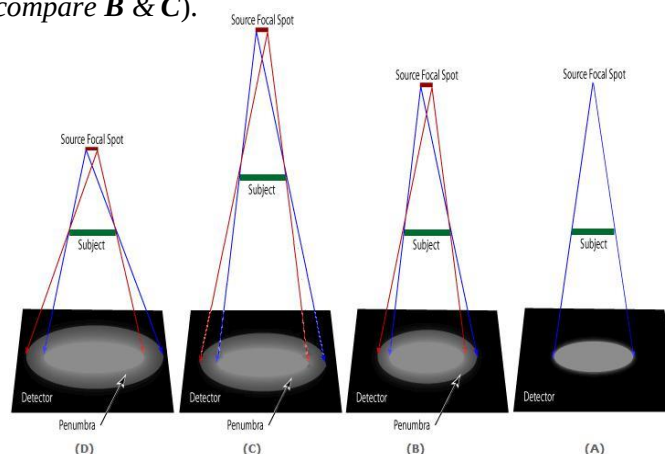
As mentioned previously, thin lead sheets (*called “lead screens”*) are typically placed on both sides of the radiographic film during the exposure (*the film is placed between the lead screens and inserted inside the cassette*). Lead screens in the thickness range of 0.1 to 0.4 mm typically reduce scatter radiation at energy levels below 150 kV. Above this energy level, they will emit electrons to provide more exposure of the film, thus increasing the density and contrast of the radiograph. Other type of screens called “fluorescent screens” can alternatively be used where they produce visible light when exposed to radiation and this light further exposes the film and increases density and contrast.

Radiographic Definition

As mentioned previously, radiographic definition is the abruptness of change from one density to another. Both geometric factors of the equipment and the radiographic setup, and film and screen factors have an effect on definition.

Geometric Factors

The loss of definition resulting from geometric factors of the radiographic equipment and setup is referred to as “*geometric unsharpness*”. It occurs because the radiation does not originate from a single point but rather over an area. The three factors controlling unsharpness are source size, source to object distance, and object to detector (film) distance. The effects of these three factors on image definition is illustrated by the images below (*source size effect; compare A & B, source to object distance; compare B & D, and object to detector distance; compare B & C*).



The source size is obtained by referencing manufacturers specifications for a given X-ray or gamma ray source. Industrial X-ray tubes often have focal spot sizes of *1.5 mm* squared but microfocus systems have spot sizes in the *30 micron* range. As the source size decreases, the geometric unsharpness also decreases. For a given size source, the unsharpness can also be decreased by increasing the source to object distance, but this comes with a reduction in radiation intensity. The object to detector distance is usually kept as small as possible to help minimize unsharpness. However, there are situations, such as when using geometric enlargement, when the object is separated from the detector, which will reduce the definition.

It should be as large as practical, and the object-to-detector distance should be as small as practical.

Codes and standards used in industrial radiography require that geometric unsharpness be limited. In general, the allowable amount is *1/100* of the material thickness up to a maximum of *1 mm*. These values refer to the width of penumbra shadow in a radiographic image.

The amount of geometric unsharpness () can be calculated using the following geometric formula

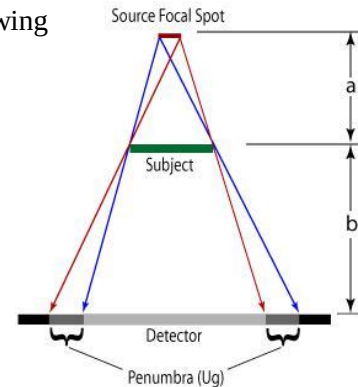
$$U_g = d_s(b/a)$$

Where;

d_s : Source focal-spot size

a: Distance from the source to the front surface of the object

b: Distance from the front surface of the object to the detector



The angle between the radiation and some features will also have an effect on definition. If the radiation is parallel to an edge or linear discontinuity, a sharp distinct boundary will be seen in the image. However, if the radiation is not parallel with the discontinuity, the feature will appear distorted, out of position and less defined in the image.

Abrupt changes in thickness and/or density will appear more defined in a radiograph than will areas of gradual change. For example, consider a circle. Its largest dimension will be a cord that passes through its centerline. As the cord is moved away from the centerline, the thickness gradually decreases. It is sometimes difficult to locate the edge of a void due to this gradual change in thickness. Lastly, any movement of the specimen, source or detector during the exposure will reduce definition. Similar to photography, any movement will result in blurring of the image. Vibration from nearby equipment may be an issue in some inspection situations.

Film and Screen Factors

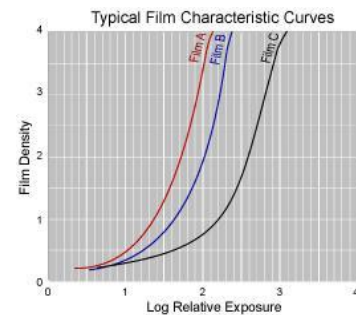
The last set of factors concern the film and the use of fluorescent screens. A fine grain film is capable of producing an image with a higher level of definition than is a coarse grain film. Wavelength of the radiation will influence apparent graininess. As the wavelength shortens and penetration increases, the apparent graininess of the film will increase. Also, increased development of the film will increase the apparent graininess of the radiograph. The use of fluorescent screens also results in lower definition. This occurs for a couple of different reasons. The reason that fluorescent screens are sometimes used is because incident radiation causes them to give off light that helps to expose the film. However, the light they produce spreads in all directions, exposing the film in adjacent areas, as well as in the areas which are in direct contact with the incident radiation. Fluorescent screens also produce screen mottle on radiographs. Screen mottle is associated with the statistical variation in the numbers of photons that interact with the screen from one area to the next

Film Characteristic Curves

In film radiography, the number of photons reaching the film determines how dense the film will become when other factors such as the developing time are held constant. The number of photons reaching the film is a function of the intensity of the radiation and the time that the film is exposed to the radiation. The term used to describe the control of the number of photons reaching the film is “exposure”.

Different types of radiographic films respond differently to a given amount of exposure. Film manufacturers commonly characterize their film to determine the relationship between the applied exposure and the resulting film density. This relationship commonly varies over a range of film densities, so the data is presented in the form of a curve such as the one for *Kodak AA400* shown to the right. This plot is usually called a film characteristic curve or density curve. A log scale is sometimes used for the x-axis or it is more common that the values are reported in log units on a linear scale as seen in the figure. Also, relative exposure values (*unitless*) are often used. Relative exposure is the ratio of two exposures. For example, if one film is exposed at 100 kV for 6 mA.min and a second film is exposed at the same energy for 3 mA.min, then the relative exposure would be 2.

The location of the characteristic curves of different films along the x-axis relates to the speed of the film. The farther to the right that a curve is on the chart, the slower the film speed (Film A has the highest speed while film C has the lowest speed). The shape of the characteristic curve is largely independent of the wavelength of the X-ray or gamma ray, but the location of the curve along the x-axis, with respect to the curve of another film, does depend on radiation quality.



Film characteristic curves can be used to adjust the exposure used to produce a radiograph with a certain density to an exposure that will produce a second radiograph of higher or lower film density. The curves can also be used to relate the exposure produced with one type of film to exposure needed to produce a radiograph of the same density with a second type of film.

5.11 RADIATION SAFETY

Radiation Health Risks

As mentioned previously, the health risks associated with the radiation is considered to be one the major disadvantages of radiography. The amount of risk depends on the amount of radiation dose received, the time over which the dose is received, and the body parts exposed. The fact that X-ray and gamma-ray radiation are not detectable by the human senses complicates matters further. However, the risks can be minimized and controlled when the radiation is handled and managed properly in accordance to the radiation safety rules. The active laws all over the world require that individuals working in the field of radiography receive training on the safe handling and use of radioactive materials and radiation producing devices.

Today, it can be said that radiation ranks among the most thoroughly investigated (and somehow understood) causes of disease. The primary risk from occupational radiation exposure is an increased risk of cancer. Although scientists assume low-level radiation exposure increases one's risk of cancer, medical studies have not demonstrated adverse health effects in individuals exposed to small chronic radiation doses.

The occurrence of particular health effects from exposure to ionizing radiation is a complicated function of numerous factors including:

- *Type of radiation involved.* All kinds of ionizing radiation can produce health effects. The main difference in the ability of alpha and beta particles and gamma and X-rays to cause health effects is the amount of energy they have. Their energy determines how far they can penetrate into tissue and how much energy they are able to transmit directly or indirectly to tissues.
- *Size of dose received.* The higher the dose of radiation received, the higher the likelihood of health effects.
- *Rate at which the dose is received.* Tissue can receive larger dosages over a period of time. If the dosage occurs over a number of days or weeks, the results are often not as serious if a similar dose was received in a matter of minutes.
- *Part of the body exposed.* Extremities such as the hands or feet are able to receive a greater amount of radiation with less resulting damage than blood forming organs housed in the upper body.
- *The age of the individual.* As a person ages, cell division slows and the body is less sensitive to the effects of ionizing radiation. Once cell division has slowed, the effects of radiation are somewhat less damaging than when cells were rapidly dividing.
- *Biological differences.* Some individuals are more sensitive to radiation than others. Studies have not been able to conclusively determine the cause of such differences.

5.12 Xeroradiography

Xeroradiography Various methods have been introduced for obtaining radiographs among which xeroradiography is a method of imaging which uses the Xeroradiographic copying process to record images produced by diagnostic X-rays. It differs from halide film technique in that it involves neither wet chemical processing nor the use of dark room.¹ Over the past 40 years, Xerox 125 system became applicable in medical sciences. A prototype Xeroradiographic imaging system specific for intraoral use was later developed² . Following these clinical trials showed that xeroradiography is superior for imaging of dental structures necessary for successful periodontal and endodontic therapy³ . Xeroradiographic radiation of 90% of more than that for silver halide radiograph has been reported, while others found that Xeroradiographic radiation is one third to half of that for halide radiograph.⁴ The imaging method was discovered by an American physicist, Chester Carlson in 1937. Later, the Xerox company followed the laboratory investigations of the technique and its potential applications in medical sciences. Others like Binnie et al Grant et al and White et al worked on phantoms and cadavers using the Xerox 125 system. Xeroradiography may be new in dentistry, but in medicine, it had long been used in the diagnosis of breast diseases, imaging of the larynx and respiratory tract for foreign bodies, Temporo-mandibular joint, skull and para-osseous soft tissues.⁵ Pogorzelska–Stronczak became the first to use xeroradiography ,to produce dental images while Xerox 125 medical system got adapted for extra oral dental use in cephalometry, Sialography and panoramic xeroradiography. Later, a prototype Xeroradiographic imaging system, specific for intraoral use, was acclaimed to be superior over halide based intraoral technique.⁴ Xeroradiography is an electrostatic process which uses an amorphous selenium photoconductor material, vacuum deposited on an aluminium substrate to form a plate .the plate enclosed in tight cassette ,may be linked to films used in halide based intraoral technique⁶ The key functional steps in the process involve the sensitization of the photoconductor plate in the charging station by depositing a uniform positive charge on its surface with a corona emitting device called scorotron. That is the uniform electrostatic charge placed on a layer of selenium is in electrical contact with a grounded conductive backing .In the absence of electromagnetic radiation ,the photoconductor remains non conductive uniform electrostatic charge when radiation is passed through an object which will vary the intensity of radiation, observed by Rawls and Owen. The photoconductor will then conduct its

electrostatic charge into its grounded base in proportion to the intensity of the exposure. After charging the cassette is inserted into a thin polyethylene bag to protect the cassette and plate from saliva.

Image Development

The generated latent image is developed through an electrophoretic development process using liquid toner. The process involves the migration to and subsequent deposition of toner particles suspended in a liquid onto an image reception under the influence of electrostatic field forces. That is by applying negatively charged powder (toner) which is attracted to the residual positive charge pattern on the photoconductor, the latent image is made visible and the image can be transferred to a transparent plastic sheet or to a paper. The toner is thereafter fixed to a receiver sheet onto which a permanent record is made. The plate is then cleaned of toner for reuse.

The Xeroradiographic Plate

The plate is made up of a 9 1/2 by 14 inch sheet of aluminium, a thin layer of vitreous or amorphous selenium photoconductor, an interface layer, and an over cutting on the thin selenium layer

The Aluminium Substrate

The substrate for the selenium photoconductor should present a clean and smooth surface. Surface defects affect the Xeroradiographic plate's sensitivity by giving rise to changes in the electrostatic charge in the photoconductor²⁹

The Interface Layer

This is a thin layer of aluminium oxide between the selenium photoconductor and aluminium substrate. The oxide is produced by heat treating the aluminium substrate. As a non conductor, the interface layer prevents charge exchange between the substrate and the photoconductor surface.

The Selenium Coating

The thickness of this layer varies from 150micron meter for powder toner development. Amorphous or vitreous selenium coating, is formed by depositing a vapour form of liquefied selenium in a high vacuum. Because of its ease of use, fabrication and durability, inherent property of electrical conduction when exposed to x-rays and ability to insulate well when shielded from all sources of light, make selenium a Xeroradiographic material of choice. On the other hand, any form of impurity adversely affects its performance. Amorphous form is used in Xeroradiography because crystalline selenium's electrical conductivity is very high which makes it unsuitable in Xeroradiography. However, amorphous selenium undergoes a dark decay of about 5% per minute. A new system of xeroradiography which uses plates with thicker selenium layer (320 micron meter) gives about 50% x-ray absorption.

Selenium Protective Coating The protective coating is a 0.1micron meter cellulose acetate overcoat. The coat bonds intimately with selenium photoconductor. It helps to prevent degradation of electrostatic lateral image through the prevention of lateral conduction of electrostatic charges. Also it impacts positively on the shelf life of the Xeroradiographic plate

Advantages Elimination of accidental film exposure, High resolution, Simultaneous evaluation of multiple tissues, Ease of reviewing, Economic benefit, Reduced exposure to radiation hazards, Wide applications- Generally Xeroradiography have interesting application in the management of neoplasm of laryngo-pharyngeal area, mammary and joint region, as well as aid in cephalometrics analysis.

Disadvantages Technical difficulties, Fragile selenium coat, Transient image retention, Slower speed, Technical limitations Uses Diagnosis of periapical pathology in endodontics, Assessment of bone height in periodontics.