

UNCLASSIFIED

Defense Technical Information Center Compilation Part Notice

ADP010762

TITLE: Various Sensors Aboard UAVs

DISTRIBUTION: Approved for public release, distribution unlimited

This paper is part of the following report:

TITLE: Development and Operation of UAVs for
Military and Civil Applications [Developpement et
utilisation des avions sans pilote [UAV] pour des
applications civiles et militaires]

To order the complete compilation report, use: ADA390637

The component part is provided here to allow users access to individually authored sections of proceedings, annals, symposia, ect. However, the component should be considered within the context of the overall compilation report and not as a stand-alone technical report.

The following component part numbers comprise the compilation report:

ADP010752 thru ADP010763

UNCLASSIFIED

VARIOUS SENSORS ABOARD UAVs

(July 1999)

E.J. Schweicher

Royal Military Academy (OMRA)

30, Avenue de la Renaissance

B-1000 Brussels-Belgium

Email: emile.schweicher @ omra.rma.ac.be

Summary : The present paper chiefly deals with imaging sensors which could be subdivided in

- * *passive ones* : visible (from 380nm to 780nm), near IR (from 0.78μ to 3μ) and thermal IR (from 3μ to 15μ) among them MWIR (from 3μ to 5μ) and LWIR (from 8μ to 12μ).
Microwave radiometers could be of interest for demining purposes;
 - * *active ones* : lidars and laser rangefinders and all kind of radar sensors among them the most important one: the SAR (Synthetic Aperture Radar) and its counterpart ISAR (Inverse Synthetic Aperture Radar).
From a communication point of view, sensor typology mainly affects the **down-link data rate**. Three categories of sensors may be defined on the base of different data rates:
 - * **Sensors featuring low information rate**
(e.g., lower than 200kbit/s), like Meteorological, NBC, ECM sensors;
 - * **Sensors featuring medium information rate**
(e.g., from 200 kbit/s to 10Mbit/s), like processed SAR/ISAR, compressed VIDEO/TV, compressed FLIR (Forward Looking Infra Red), uncompressed millimeter wave Radar image, Elint/ESM sensors;
 - * **Sensors featuring high information rate**
(e.g., over 10Mbit/s), like raw SAR/ISAR, High Definition VIDEO/TV, High definition FLIR.
- Because of a lack of time, we mainly restrict ourselves to the two most important and difficult imaging sensors, namely the thermal (IR) imager (described in part A) and the SAR (described in part B).

Introduction and perspectives

In order to deal with all possible UAV imaging sensors, we better choose the example of a recently introduced UAV: the General Atomics Predator UAV.

The Predator sensor **payload** includes an EO (Electro-Optical) suite, a Ku-band SAR sensor, Ku-band and UHF-band satellite communications (SATCOM), a C-band light-of-sight data link, and a GPS/INS navigator.

The Predator's SAR sensor is the Northrop Grumman (Westinghouse) Tactical Endurance Synthetic Aperture Radar (TESAR). TESAR provides continuous, near real time strip-map transmitted imagery over an 800 meter swath at slant ranges up to 11km. Maximum data rate is 500,000 pixels per second. The target resolution is 0.3meters. TESAR weight and power consumption are 80kg and 1200W respectively. A lighter weight, lower cost SAR is currently in development for Predator.

The Predator's EO sensor suite is the VERSATRON Skyball SA-144/18 quartet sensor. It consists of a PtSi 512×512 MWIR (Mid Wave IR) FLIR with six fields of view (to easily perform either detection or recognition or identification), a color TV camera with a 10X zoom, a color TV 900mm camera and an eyesafe pulsed Er: glass laser rangefinder (this Er: glass laser could advantageously be replaced by an eyesafe Er: YAG laser because YAG is a better heat sink than glass enabling a higher efficiency).

The diameter of the EO sensor turret is relatively small-35cm. The turret has precision pointing with a line-of-sight stabilization accuracy of 10μ rad.

It is anticipated that high performance UAV's of the year 2010 will have a broad range of missions, including surveillance, reconnaissance, communication , intelligence gathering of threat

electronic emissions, target designation for weapons attacking moving targets, and communication relay.

In view of delivering better contrasted images (especially in hazy conditions), visible cameras could be replaced by near IR cameras which are less affected by the Rayleigh scattering (proportional to $\frac{1}{\lambda^4}$) of sunlight which is so detrimental to UV and blue light.

The biggest drawback of the SAR is the heavy and sophisticated signal processing. Presently, most UAVs are using on ground SAR signal processing requiring a high information data rate for the down-link in order to send the raw SAR data to the ground station. This is, for instance, the case for the SWIFT 1 UAV SAR (weighting about 25kg) from SAGEM. But SAGEM has a SWIFT 2 UAV SAR under development which will be able to perform real time on board SAR signal processing requiring only a medium information data rate for the down-link.

To achieve this ultimate goal of **a low-weight on board real time SAR signal processing** it is worth to notice that in recent years a new transistor appeared enabling very high speed of operation (maximum frequency in excess of 100 GHz): the HEMT (High Electron Mobility Transistor) using the fundamental concept of QW (Quantum Well). The HEMT enjoyed already several improvements which are, in order of increasing speed of operation:

- the GaAs LM (lattice matched) HEMT
- the GaAs PM (pseudomorphic) HEMT
- the InP LM HEMT
- the InP PM HEMT

Very recently a last improvement appeared: the **metamorphic HEMT layers**. IMEC (Leuven) started the growth of such layers in 1998. These structures are identical to the ones on InP, but are now grown on GaAs substrates. With a 2 μ m thick buffer layer, the 3.8% lattice mismatch between GaAs and InP can be overcome. This buffer can be either a quaternary (AlGaAsSb) or a ternary (InAlAs) compound. Of both types layers have been grown successfully at IMEC and devices based on these structures yielded typically 90% of the performance of comparable InP-based devices.

With the technique of the metamorphic layers, the advantages of GaAs substrates (availability, quality, price and robustness) can be combined with those of InP based devices (superb performance).

Unmanned air vehicles (UAVs) have their nonviolent side, too, thankfully. A new type of Lucky Lindy, the Laima UAV, became the first of its kind to cross the Atlantic, and it did so on a budget-conscious 105 liters of fuel.

Equipped with temperature, pressure, humidity, and windspeed instruments, the 12-kg craft steered its way last August (1998) from St. John's, in Newfoundland, to South Uist Island, one of Scotland's Outer Hebrides. It flew the 3200-km distance in 26 hours. Sponsored initially by the Australian Bureau of Meteorology, the Laima is now being supported by weather organizations in several nations.

Micro Air Vehicles (MAVs) could provide individual soldiers with surveillance and targeting information (cfr IEEE Spectrum, January 1999, pp.73-78). Two candidate designs, initially equipped with CCD (Charge-Coupled Device) cameras, are a novel flying disk from AeroVironment Inc., in Monrovia, Calif., and a fixed-wing craft from Sanders, Nashua, N.H. The chief design hurdle, apart from the aerodynamics, is getting enough power into such small objects.

Part A.

Thermal (IR) Imager

A. 1. Thermal Imaging Devices

Like the image intensifier, the thermal imager is a *passive* night vision device being thereby *undetectable* by the opponent. But its working is independent of the ambient illuminance because a thermal imager is sensitive to the *thermal IR* radiation (composed of all wavelengths λ between 3μ and 15μ , where $1\mu = 10^{-6}\text{m}$) emanating naturally from all observed objects.

The thermal image has another aspect than the "visible" image: e.g., a white man and a black man are both seen "white" while a polar bear is seen entirely black. Consequently, we have to learn to interpret thermal images.

Because the atmosphere exhibits 2 transmission windows (cfr fig. A.1) within the thermal IR wavelengths ($3\mu < \lambda < 15\mu$), 2 kinds of thermal imagers coexist, respectively the 3μ to 5μ imager (devoted to the same window) regularly called MWIR (Mid Wave IR) and the 8μ to 12μ imager (dedicated to the 8μ to 14μ window) sometimes termed LWIR (Long Wave IR).

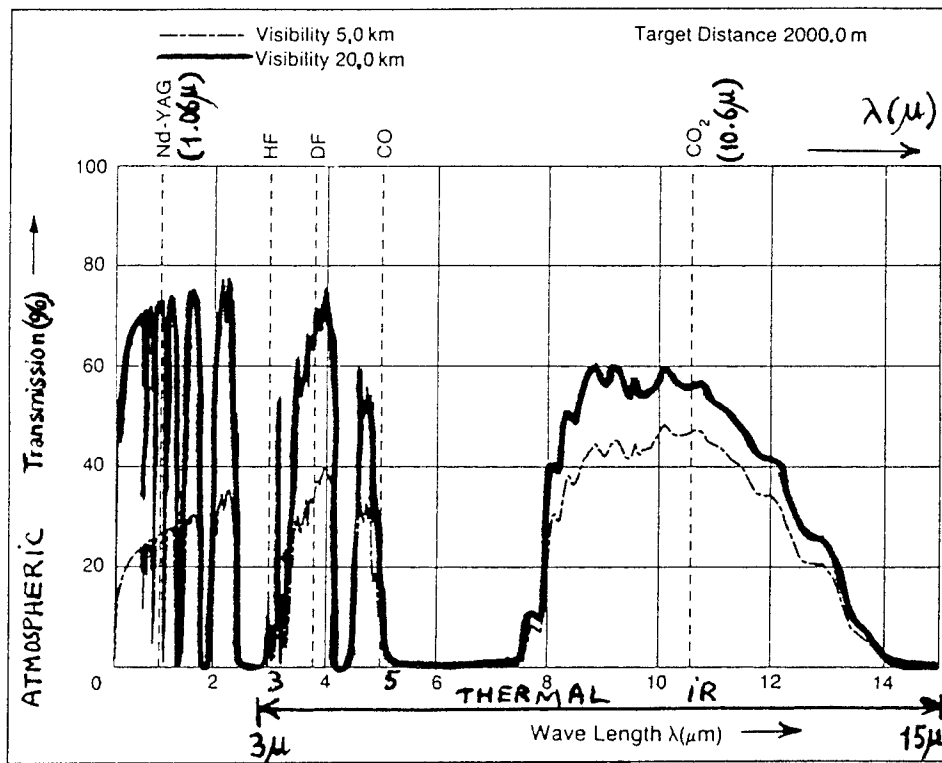
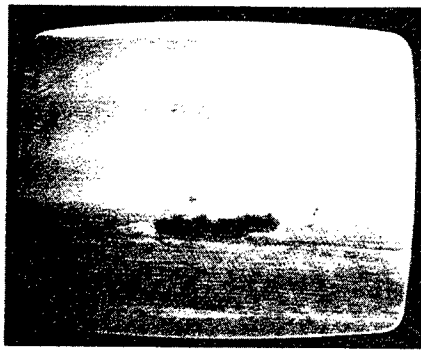
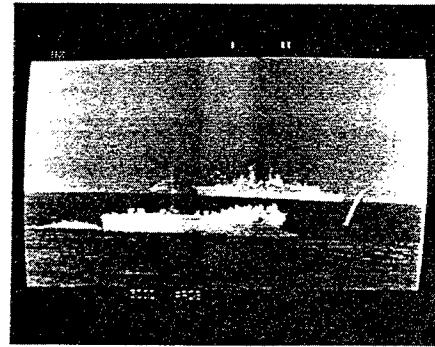


Fig.A1. Atmospheric transmission for a target distance of 2km with the wavelengths of the main lasers of military interest. The visibility distance is a distance at which a unitary contrast object is perceived with a visual contrast of 2% (which is an invisible one because the lowest contrast perceptible by the best eye is 3%). It is seen that only the CO₂ laser is compatible with thermal imagers (8 μm - 12 μm) and can penetrate mist, smoke and dust clouds.
MILTECH 5/87, p.126.



Normaal TV beeld



Thermisch TV beeld

Fig.A2. (Visible) TV image (left) and thermal image (right) taken with a Th IR camera in the ($8\mu\text{m}$ - $12\mu\text{m}$) band on a hazy day in Den helder (NL). On the lefthand "visible" image, one ship is barely distinguishable at a few hundred meters from the pier. On the (righthand) thermal image a second vessel is clearly visible at 6km from the pier.

Thermal imaging allows the user to see through external visual camouflage means, smoke screens, vegetation (even heavy brush); this is called the *decamouflaging effect* of thermal imagers.

Thermal imagers in the 8 to 12μ waveband work fairly in adverse weather conditions and can penetrate moderate natural (like drizzle, haze, fog, mist) and manmade (like dust clouds and smoke screens) obscurants. This is illustrated by fig.A2.

Furthermore the maximum range of thermal imagers is about 10 times that of image intensifiers. Some FLIR's (Forward Looking Infra Red) equipping, e.g. the F-117, can probably "see" at 40km. Spaceborne FLIR's can even "see" at several hundreds of km's.

Considering all the previous advantages of thermal imaging, one could wonder why image intensifiers are still in use. The response is that, today, thermal imagers are about 5 times bulkier, heavier and more expensive than image intensifying devices.

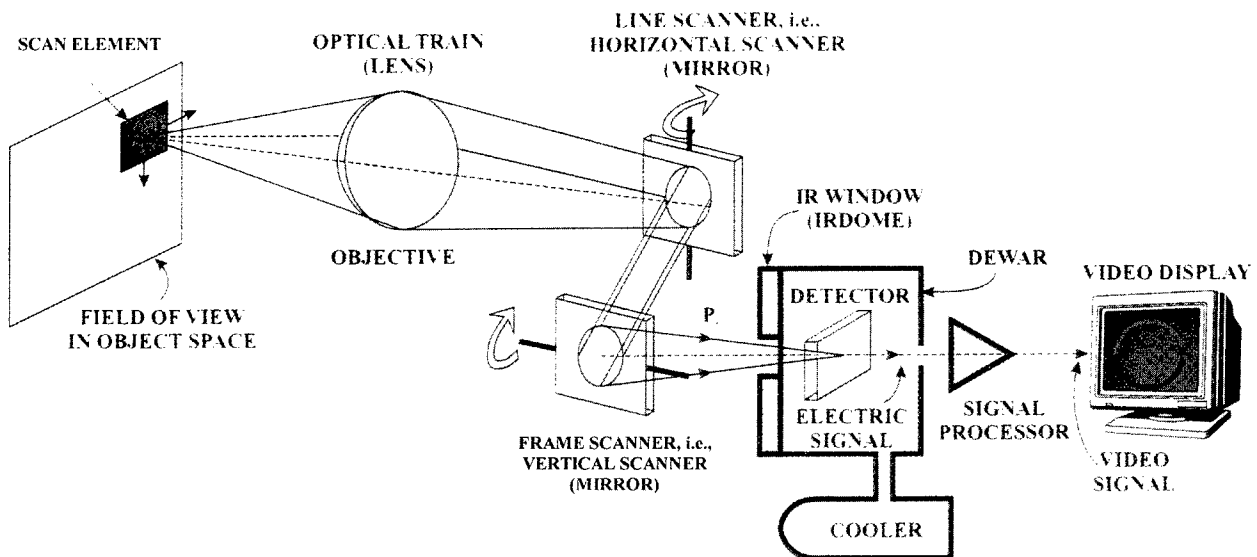


Fig.A3. Simple single-detector dual-axis thermal imager; the "scan element" is also often called "resolution element" and is nothing else but the usual "pixel", i.e., it is the detector's image in the plane of the object.

Figure A3 depicts the structure of a single-detector thermal imager using a 2D optomechanical scanning as used by FLIR's. The objective is a converging system imaging the *resolution element* (which is the detector's image in the object space and corresponds to the transverse section of the radar resolution cell) of the object onto the IR detector. This detector produces a weak electrical signal which is amplified to drive a display which is electronically scanned in *synchronism* with the detector's scanning to generate a visual counterpart of the (invisible) IR image. The opto-mechanical scanning (by both scanning mirrors) of fig.A3 is difficult to ruggedize. Therefore, there is a clear evolution towards electronic scanning which is inertionless and more rugged by nature and whose first achievement was a staring array from Mitsubishi using 512×512 PtSi detectors monolithically integrated with their Si CCD readout cells. Of course, electronic scanning requires FPA's (Focal Plane Arrays) - see fig.A4 - with no mechanical scanning at all.

Real time thermal images with high temperature resolution (enabling eventually to image objects at many tens of kilometres) require a quantum detector which is made of the following semi-conductors:

- InSb or PtSi cooled at 77K (= -196°C) for the 3μ to 5μ band
- CMT (i.e., $\text{Cd}_x\text{Hg}_{1-x}\text{Te}$ with $0 < x < 1$) cooled either at -80°C for the 3μ to 5μ band or at 77K for the 8μ to 12μ band.

The deep cooling prevents *narcissism* (explained by fig.A5 and its caption) of the quantum detector and can be performed either with Peltier elements to cool down to -80°C or with Joule-Thomson expansions or Stirling coolers to obtain 77K. Fig.A6 summarizes the main features of those three "operational" cooling mechanisms, since liquid nitrogen ($T = 77\text{K}$) cooling can only be used in laboratory equipments.

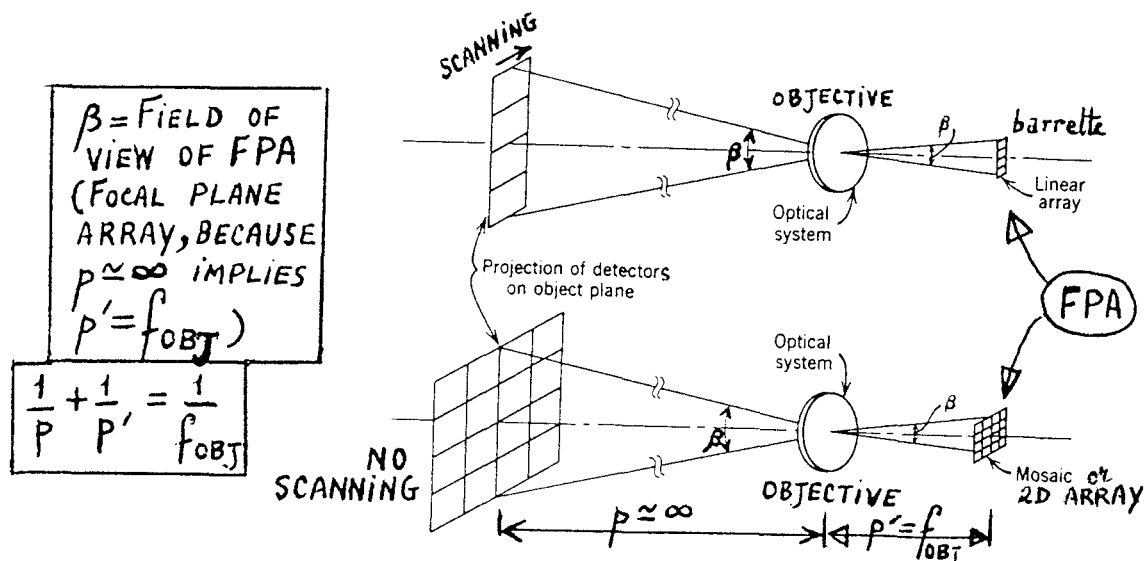


Fig.A4. FPA (Focal Plane Arrays), i.e., multielement detectors.

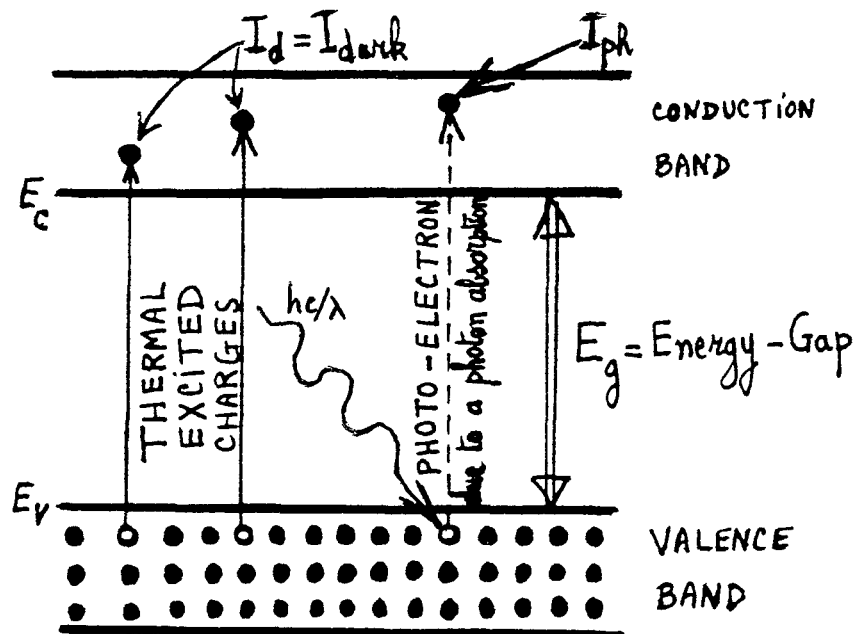


Fig.A5. Working of a semi-conductor quantum detector: an incident photon creates a hole-electron pair. Black dots represent electrons while white dots represents holes. The so-called NARCISSISM of such a detector is due to the thermal induced charges which form a dark current I_d in photo voltaic detectors while the photoelectrons constitute the photo-current I_{ph}

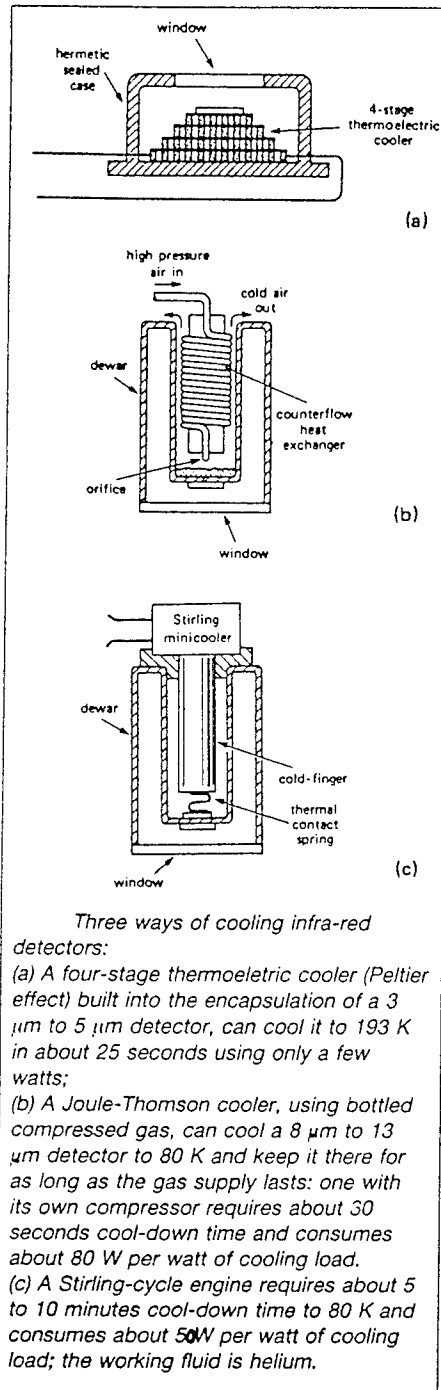


Fig.A6. Three "operational" ways of cooling IR detectors.

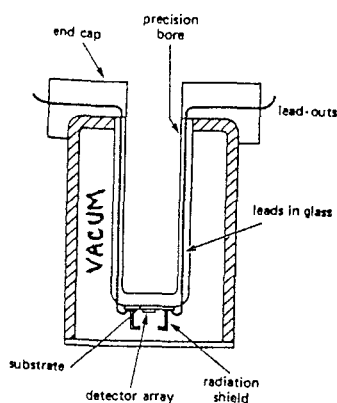


Fig.A7. Simplified cross section of DEWAR with IR detector mounted.
The cold stop (or radiation shield) prevents the detector from viewing
the "hot" parts of its direct environment.

Fig.A7. gives some details about the arrangement of the detector within the DEWAR.

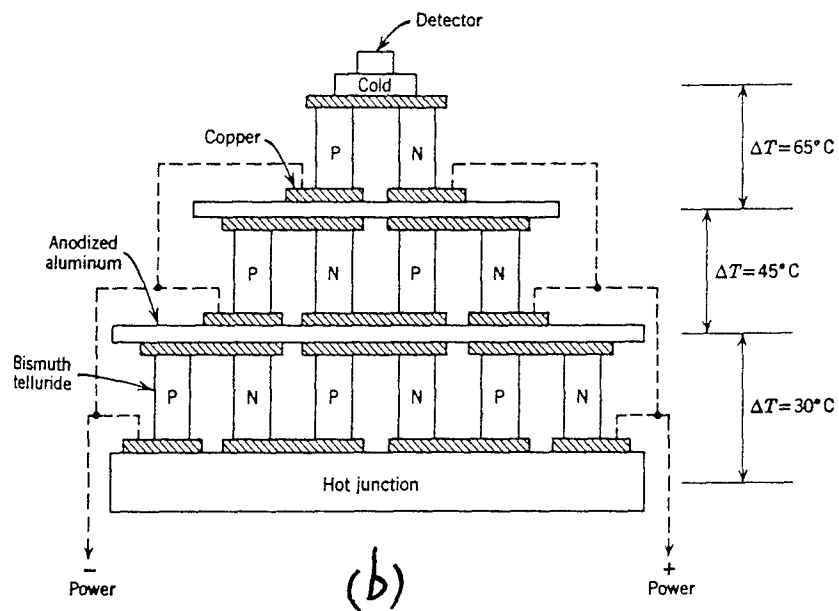
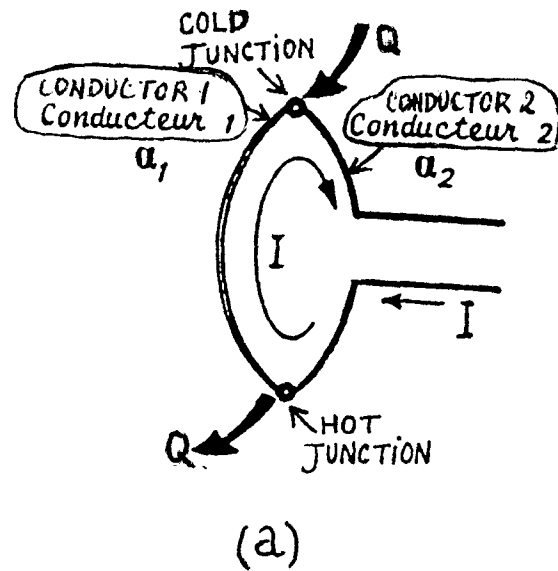
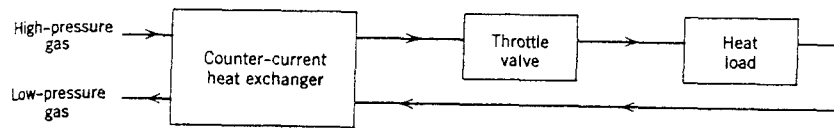
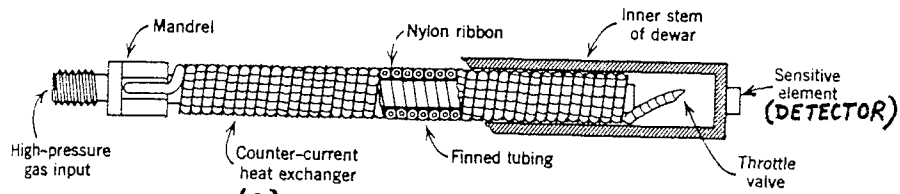


Fig.A8. Thermoelectric cooling; (a) principle of the Peltier effect;
(b) three-stage thermoelectric refrigerator.

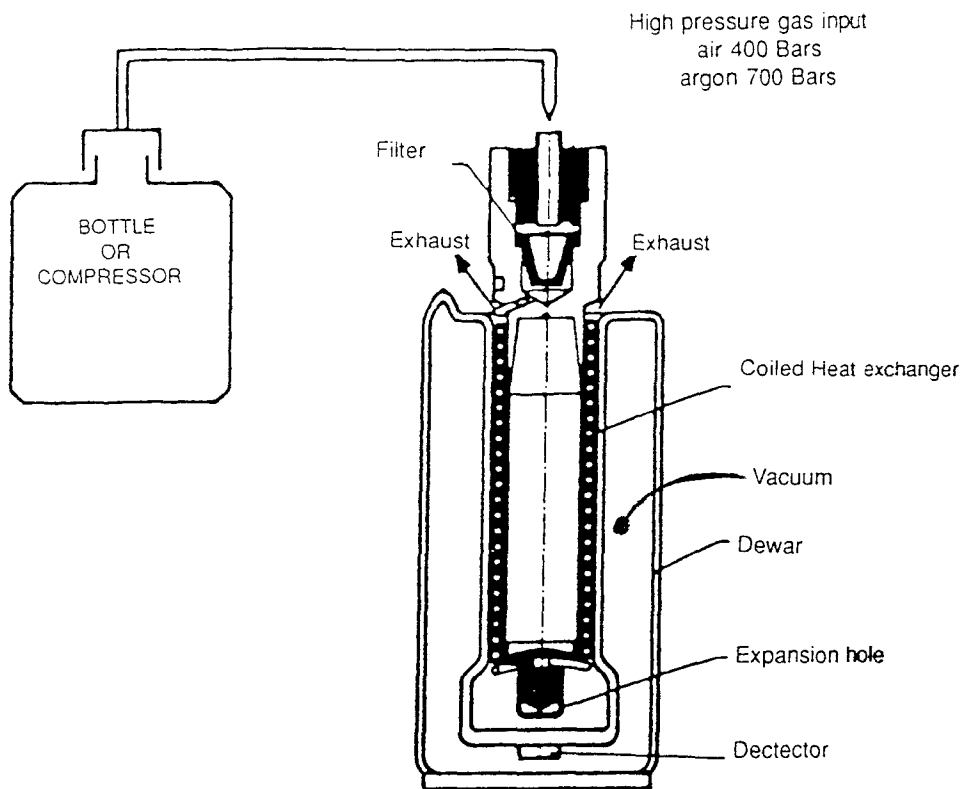
Fig.A8 (a) reminds the principle of the thermoelectric cooling by the Peltier effect while Fig.A8 (b) details a three-stage thermoelectric cooler.



(a) Schematic of Cryostat Operation



(a) Construction of a Cryostat



(b) Cooling by Joule-Thomson.

Fig.A9. (a) Principle of the Joule-Thomson cryostat; (b) more details about the Joule-Thomson cooling.

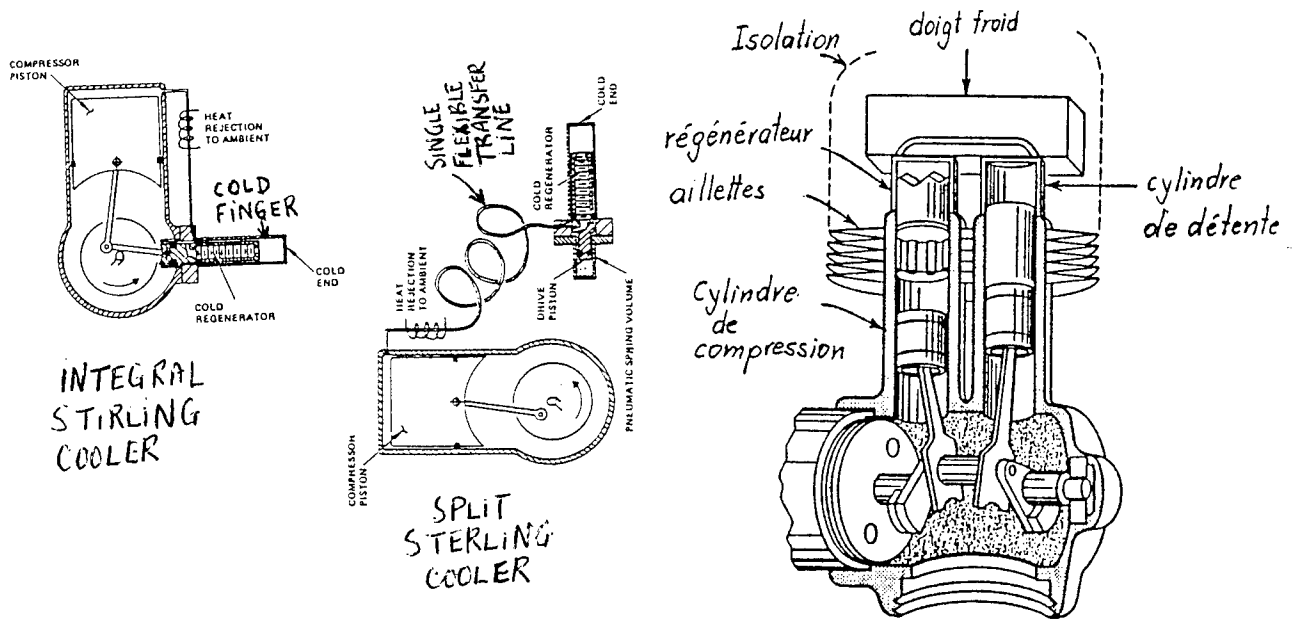


Fig.10. Views of the integral-Stirling cycle cooler (left) and the more recent split-Stirling cycle cooler (right). The latter allows the compressor to be conveniently located
[Photonics Spectra, Oct 89, p.76]

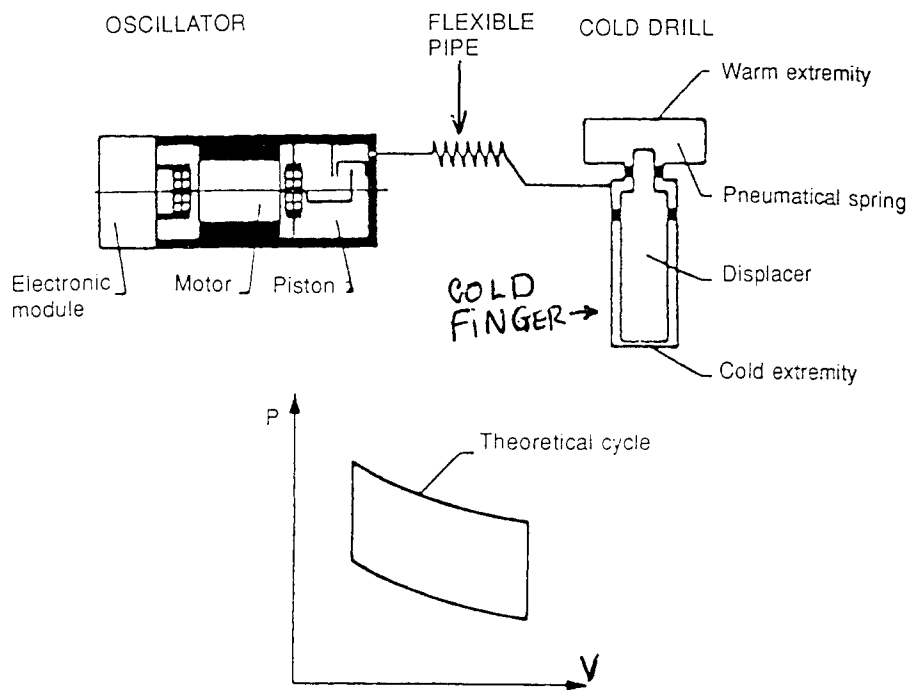


Fig.A11. Split Stirling cooler and Stirling heating cycle.

Fig.A9 explains better the cooling by a Joule-Thomson expansion of a dry high pressure gas usually contained within a bottle; for this reason this cooling method is sometimes called "bottle cooler". Figures A10 and A11 detail better the Stirling cooler. Most IR detectors exhibit also a piezoelectric effect; hence, they suffer from a microphonic noise. For that reason the previous integral Stirling coolers (cfr fig.A10) have been replaced by the split Stirling coolers in order to avoid that microphonic noise of the detector which rings like a bell when it is pinched.

B. 2. Physics and Investigation of Thermal Images

IR thermography is based on the thermal radiation laws (illustrated by fig.A12) established by Planck, Wien and Stefan-Boltzmann. Wien's law, expressing the wavelength of maximum radiation $\lambda_{\max}(\mu) = 2898 / T(K)$, explains nicely why the plume ($T = 800K$) of a jet aircraft or a missile is best seen by a 3μ to 5μ thermal imager while devices working in the 8μ to 12μ band are ideal for thermal imaging of objects around $300K$.

A thermal imager will see the power P_r radiated by the resolution element of fig.A3, i.e., it will see the product of the area of the resolution element by its radiant exitance $\epsilon\sigma T^4$ where σ is the (universal) constant of Stefan-Boltzmann while ϵ is called *emissivity*.

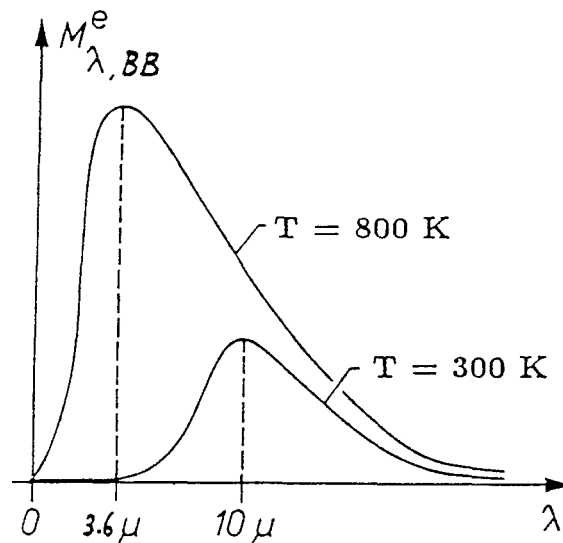


Fig.A12. Evolution of the blackbody radiation with increasing temperature

For a given imager and a given object's distance, it follows that *the thermal imager sees the product ϵT^4 from the surface of the resolution element*, i.e. the thermal imager sees heat patterns corrupted by ϵ which is a dimensionless factor whose values lie between 0 and 1. Emissivity can vary with the direction of measurement and it is a function of the type of material and, more especially, of the surface state of the observed material because the *thermal radiation is intrinsically a surface phenomenon*. Furthermore, for an opaque material, the emissivity is related to the material reflectance R by the equation $\epsilon = 1 - R$. The previous features explain why pleats in clothes are visible on thermal images and why a polished metal (behaving as a mirror) exhibits a $\epsilon = 0.1$ while the same oxidized or painted metal is characterized by a $\epsilon = 0.9$.

B. 3. Applications of IR Thermography

We are going to cite a non exhaustive list of applications (most civilian ones) of IR thermography.

(1) the search for heat leakages in buildings is obvious because a thermal imager displays heat patterns.

(2) NDI (Non Destructive Investigation) is also a straightforward application because the recording of thermal patterns (which are disturbed by hidden defects) in numerous industrial applications helps monitor potential dangers or defects, e.g., flaws in materials on a processing line can be detected.

(3) electric and electronic industry where the application is entirely based upon the heat production by Joule effect which reveals hot spots, bad connections...

(4) medical applications: it should be stressed that IR thermography can generally find breast cancer because the tumour is close to the skin.

(5) Agricultural forecast, forest surveillance, vegetation degradation and weather forecast: those applications are based on heat patterns and spatial variations of the emissivity.

(6) Archaeology: underground galleries, ruins or remains are revealed either by the decamouflaging effect or by spatial variations in the emissivity of the vegetation.

(7) Automobile: malfunctions can be found in catalytic pods, in cylinders (bad combustion), in air-filters (e.g., stops), in exhaust gases,...

(8) Pollution: monitoring of overheating or of chemical pollution (emissivity variations)

(9) Remote sensing especially in military applications like

- revealing a tank concealed behind a heavy brush or visual camouflaging means
- revealing a squad of infantrymen crossing a field under cover of darkness.
- or even the tracks left by feet of infantrymen, this last possibility having been used, during the Vietnam war, by the US Army to detect Vietcongs concealed behind vegetation.

(10) surveillance, security, alert systems:

- sensing intrusions into hazardous or restricted areas
- thermal imagers allow harbours, docks as well as airports and high security locations to be clearly viewed in both total darkness and poor visibility.

(11) fire fighting: thermal imagers enable to find the origin of fires behind thick smoke.

(12) Rescue e.g.

- locate people overcome by smoke
- detect survivors under the debris of collapsed buildings by means of their body heat
- detect concealed bodies: this possibility has been used by the British Army in Northern Ireland to find bodies (hidden behind the vegetation) of hostages killed by the IRA (a FLIR aboard of a flying helicopter was used for this purpose)

(13) Preventive maintenance: aboard a vehicle or an aircraft, items which are going to breakdown soon can be revealed either by unusual hot spots or by cold areas, enabling to replace them in advance. British Airways and the RAF use thermal imaging to locate trapped water in aircraft structures by scanning likely areas soon after the aircraft has landed from a high altitude flight; any trapped water shows as a cold spot in the structure.

(14) Law enforcement: military IR techniques are being adapted by various law enforcement agencies to battle crime and drugs. New handheld IR imagers are replacing or supplementing image intensifiers, which are less effective, particularly for seeing through camouflage or on moonless or cloudy nights because they are light-sensitive rather than thermally sensitive, and require some ambient light to be effective.

B. 4. Thermal and Quantum Detectors

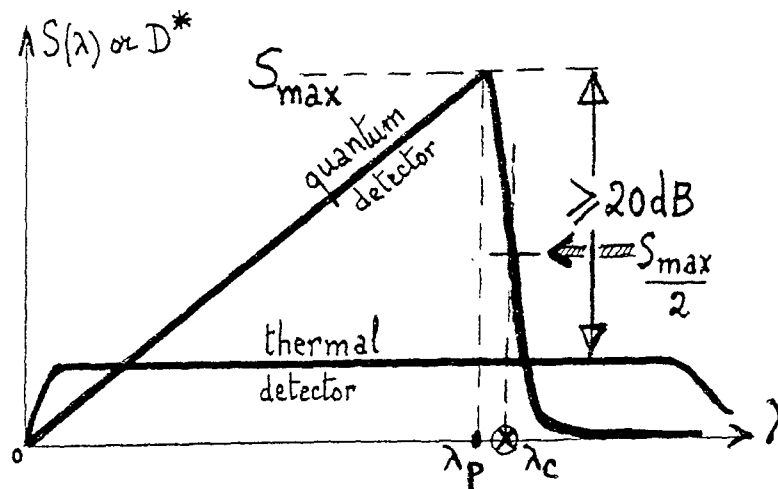


Fig.A.13. Spectral sensitivity $S(\lambda)$ and specific detectivity D^* exhibit the same shape for a detector. While the working of a thermal detector is almost wavelength independent, a quantum detector is much more sensitive but also strongly wave-

length sensitive. Notice that $D^* = \frac{(A \cdot \Delta f)^{1/2}}{\text{NETD}}$

In the case of an "uncooled" thermal imager, the detector of fig.A.3 is a **thermal detector** which may work at ambient temperature (300K) but, if compared to a quantum detector, is much slower (because of the thermal inertia) and (see fig.A.13) much less sensitive, thereby restricting the uncooled thermal imager to ranges of less than one km.

Unlike thermal detectors, **quantum detectors** are (see fig.A.13) much more sensitive (enabling high detection ranges) and faster because they use the photoelectric effect (illustrated by fig.A.5) which is a direct interaction between photon and electron.

Clearly it follows from fig.A5 that a photon

$$h\nu = \frac{hc}{\lambda}$$

can only be absorbed if its energy satisfies the condition

$$\frac{hc}{\lambda} > E_g$$

or, equivalently, if its wavelength obeys the condition

$$\lambda < \lambda_c$$

where $\lambda_c = \frac{hc}{E_g}$ is the famous cut-off wavelength of a s.c. (semi-conductor) usually computed by the practical formula

$$\lambda_c(\mu) = \frac{1.24}{E_g(\text{eV})}$$

Since a s.c. quantum detector can only detect wavelengths below the cut-off wavelength, it follows from the last formula and from fig.A.13 that a MWIR (3μ - 5μ) imager necessitates a s.c. with a gap $E_g = 0.25\text{eV}$ (corresponding to $\lambda_c = 5\mu$) while a LWIR (8μ - 12μ) imager requires a s.c. with a gap $E_g = 0.1\text{eV}$ which are **narrow gaps** compared to the typical value $E_g \approx 1\text{eV}$ of usual s.c. like Si, Ge or GaAs. Typical narrow gap s.c. quantum detectors are InSb and CMT (Cadmium Mercury Telluride) for MWIR imagers and CMT (i.e., again $\text{Cd}_x\text{Hg}_{1-x}\text{Te}$, but with another value of x because E_g and thus λ_c depend on x) for LWIR imagers.

Because of their inherent narrow gap E_g , s.c. quantum detectors feature, at room temperature (300K), a much too high value of their dark current

$$I_d = I_o \exp\left[-\frac{E_g}{(kT)}\right]$$

explaining their narcissism and, consequently, their requirement of a deep cooling to drastically reduce this narcissism.

Concerning the thermal detectors, we concentrate our study on the two fastest thermal detectors (enabling 25 images per second): the micro-bolometer (fig.A14) and the pyroelectric detector (fig.A15) which is a capacitor with temperature dependent dielectricum. Unlike the pyroelectric detector (cfr fig.A16), the (micro)bolometers --which are temperature depending resistors, i.e., thermistors-- do not require that the scene be chopped, i.e., bolometers can operate with or without a chopper.

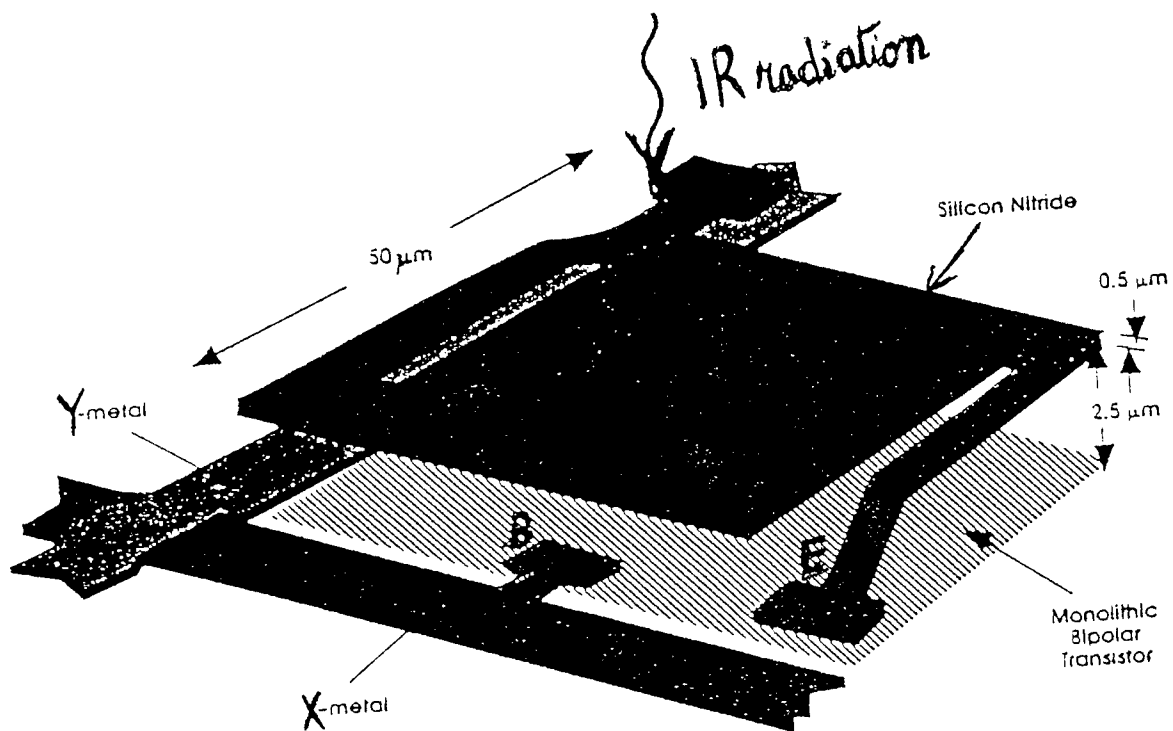


Fig.A14. Honeywell's FPA structure using microbolometers which are in fact silicon thermistors
[Photonics Spectra, May 94, p.44]

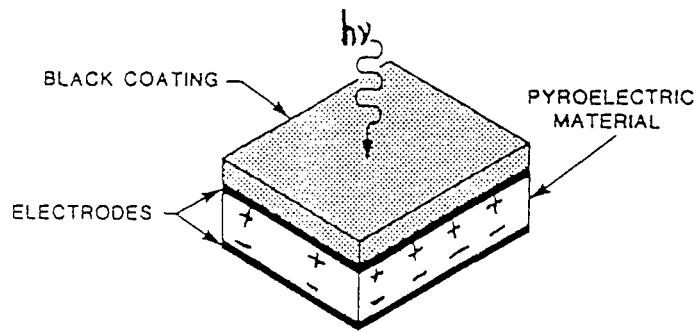


Fig.A15. A pyroelectric material has electrical polarization (permanent dipolar moment) even in the absence of an applied voltage. The materials are usually ferroelectric crystals but can be also organic polymers (like PVDF) On heating the material expands and produces a change in the polarization which builds up a charge on opposite surfaces. This causes a current to flow in the circuit which connects the electrodes. The black coating enhances the absorption of radiation and achieve (cfr fig. A13) a broad range of constant sensitivity.

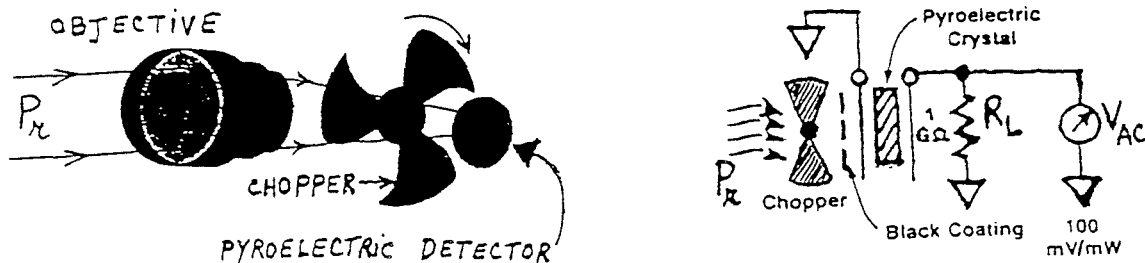


Fig.A16. Structure of a pyroelectric detecting set-up with a chopper in front of the pyroelectric detector.

We just established that a s.c. quantum detector can only "see" wavelengths obeying

$$\lambda < \lambda_c$$

This condition can be overcome with GaAs (although $\lambda_c = 0.87\mu$ for GaAs) by resorting to QWIPs (Quantum Well Infrared Photodetector) whose working is explained by fig.A17 and its caption; in the case of a QWIP, the cut-of wavelength λ_c is given by

$$\lambda_c(\mu) = 1.24/\Delta E$$

where $\Delta E - \Delta E$ is the energy separation between the single bound state E_{e1} and the continuum of the conduction band-can be as low as 0.1eV by a proper choice of the quantum well thickness t and the composition x of $\text{Al}_x\text{Ga}_{1-x}\text{As}$

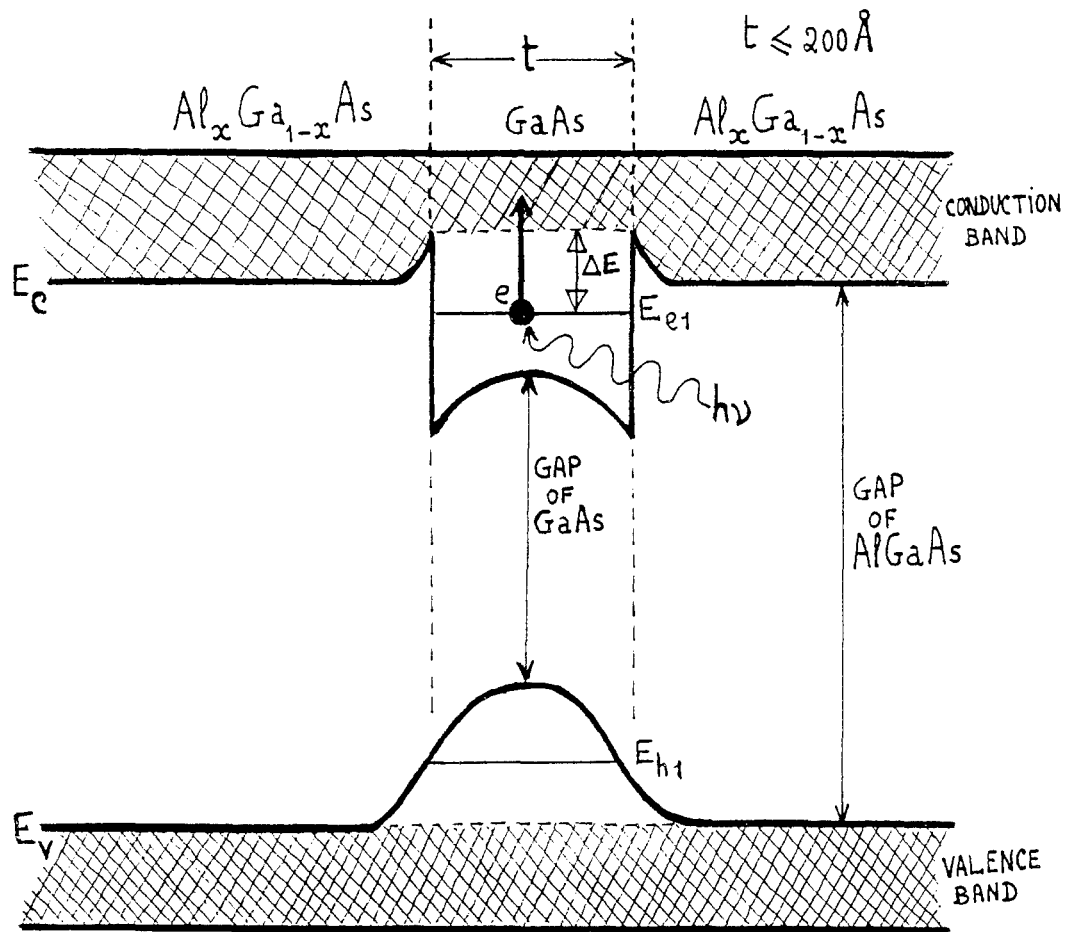


Fig.A17.Stucture and working of a QWIP. The continuum of energy levels is represented by the crosshatched areas. The QW detector uses **intraband** absorption instead of the interband absorption used by conventional detectors. The realization of the QW (Quantum Well) requires a well tickness t less than the de Broglie $\frac{h}{p}$ wavelength which is 20nm for electrons within GaAs.

B. 5. Staring Arrays for Electronic Scanning

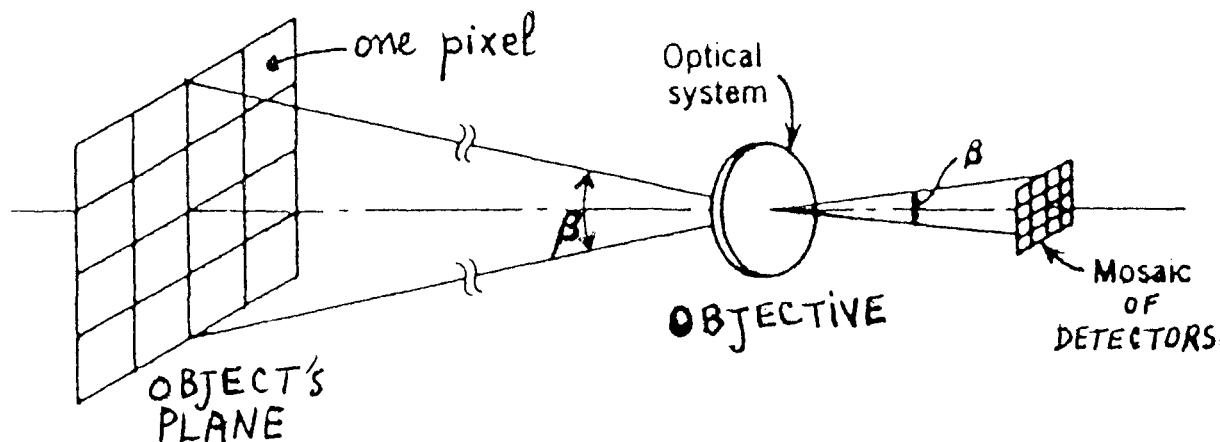


Fig.A18. A (two-dimensional) mosaic of detectors needs no mechanical scanning if its image on object's plane covers the whole FOV β of the imager: that is the definition of a staring array.

The opto-mechanical scanning of fig.A3 is difficult to ruggedize and suffers also from its mechanical inertia. Therefore, there is a clear evolution towards **electronic scanning** (comparable to phased arrays in radars) which uses special FPAs (Focal Plane Arrays) called **staring arrays**, i.e., large 2-D detector arrays which can *stare-out* into the scene giving a one-to-one correspondance (cfr fig.A18) between each pixel or scan-element (image of one single-element detector of fig.A3) and each single-element detector , which is similar in many ways to the CCD detectors used in modern lightweight TV cameras (F: camescopes). Indeed by using a mosaic, that is a 2-D array (fig.A18)

of detectors, it may be possible to cover the whole FOV (Field of View) of the imager in the object field, as depicted by fig.A18, without any mechanical or optical scanning motion provided the mosaic is big enough and comprises enough (typically, at least 512×512) single-element detectors in order not to waste the spatial resolution of the thermal imager. A staring array eliminates the expensive and tricky high-speed opto-mechanical scanning used in present thermal imagers by matching one detector to each picture element (or pixel) resulting in a compact unit which can be easily carried and used by one man.

To perform a pure electronic scanning, it "suffices" to **sequentially** read the output electrical signals of all detectors of the mosaic of fig.A18. This sequential read-out is achieved by a 2-D multiplexer which is composed either of CCD readout cells or of CMOS switches as illustrated by fig. A19.

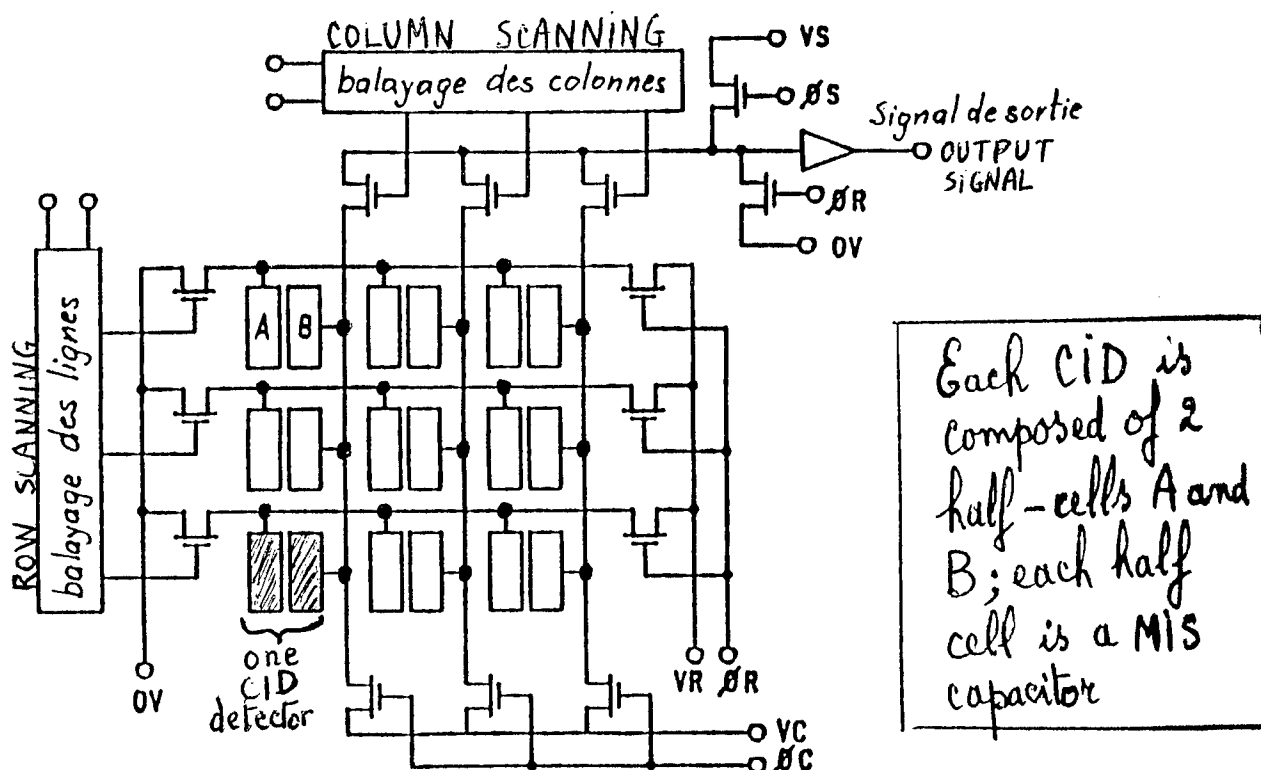


Fig.A19. Multiplexer composed of CMOS-switches for the sequential readout of 2-D array of special IR detectors called CID detectors.

Unfortunately, because of technological maturity, **most multiplexers are presently made of silicon which is lattice mismatched with most narrow gap s.c. quantum detectors**, thereby preventing monolithic integration (in one single plane) of the IR detectors and their corresponding Si readout cells. Since IR detectors like InSb and CMT are not compatible with Si, they require a **hybrid** integration (represented by fig.A20) with the additional difficulty that the whole device of fig.A20 has to be cooled inside the DEWAR of fig.A3 and that the thermal expansion coefficients of InSb and CMT are quite different from those of Si or GaAs. As illustrated by fig.A21 the same situation prevails for the microbolometer and pyroelectric detectors; fortunately the latter may work at room temperature, thereby eliminating the difficulty of the thermal expansion coefficient.

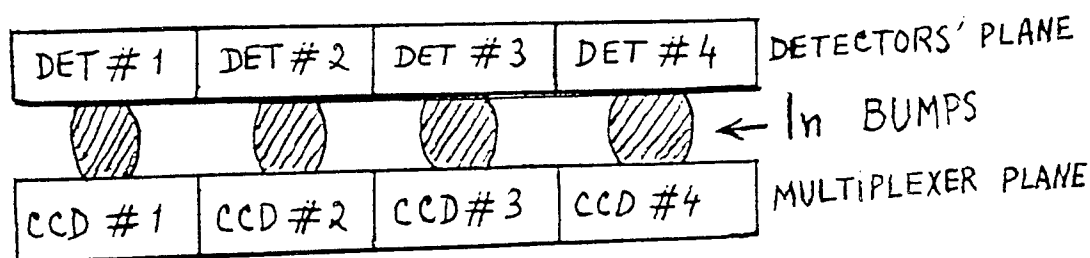


Fig.A20. Hybrid integration for InSb and CMT detectors

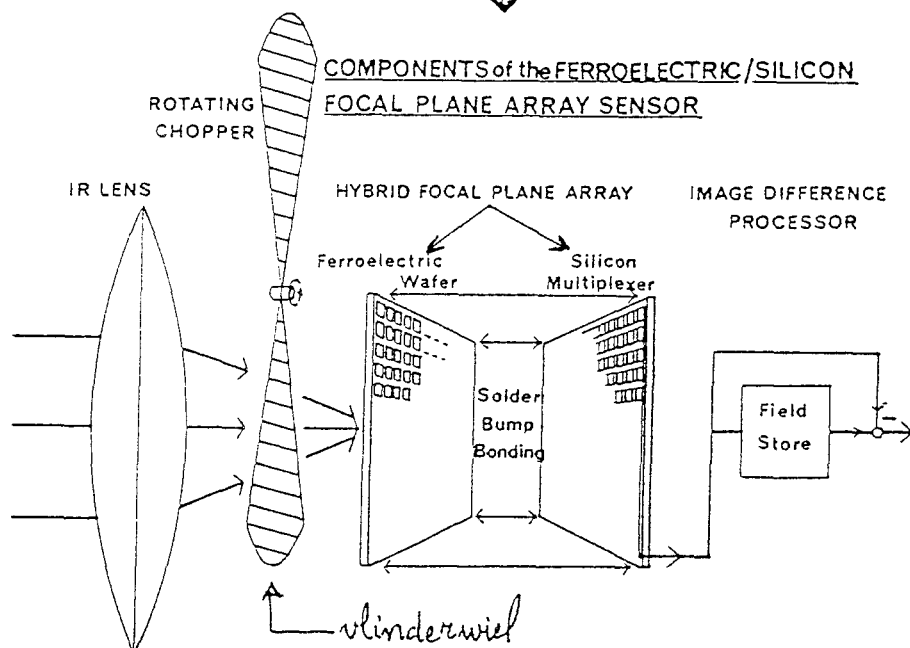
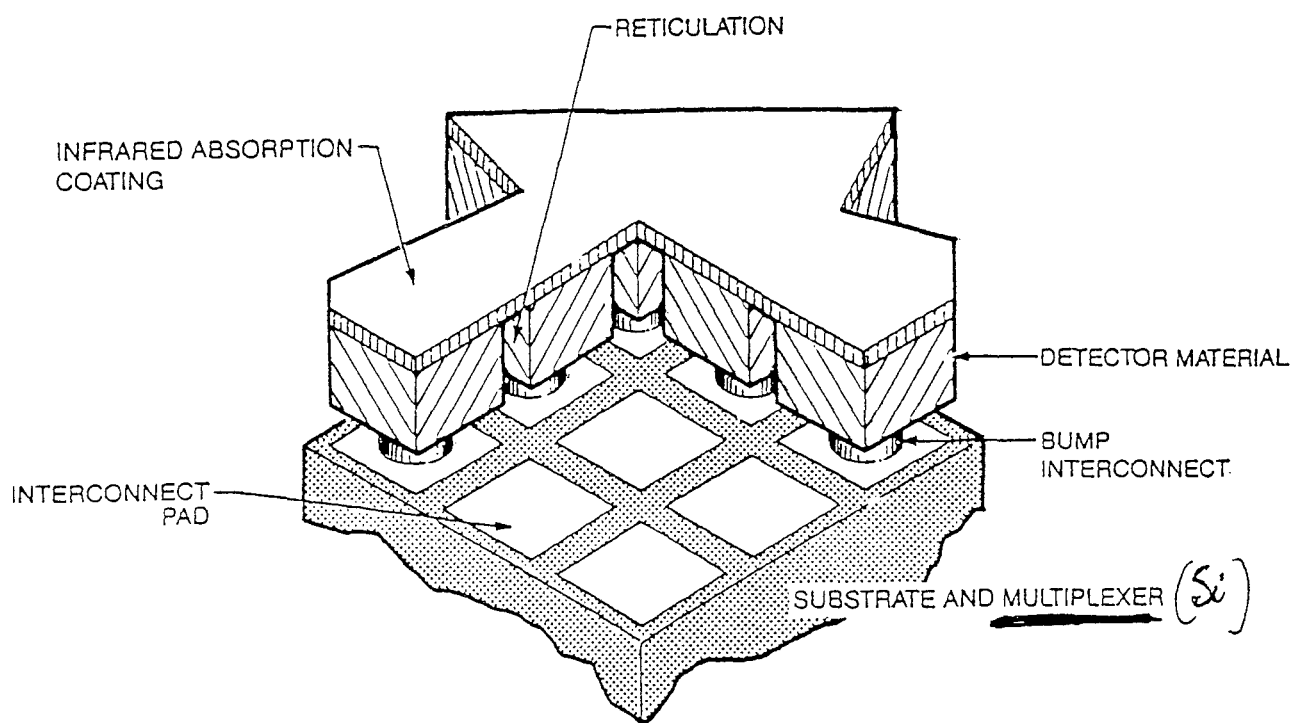


Fig.A21. Schematic of the Ferroelectric / Hybrid Focal plane Array and its use in an IR sensor head.

Recently, new thermal detectors appeared enabling monolithic integration with their readout multiplexing cells. Those detectors are arrays of micromachined thermal sensors allowing "uncooled" IR detection in the $8\mu\text{--}12\mu$ band: **the polySiGe material used as thermistor is CMOS compatible** in contrast to existing microbolometers. Imaging arrays are under development at IMEC (Leuven) as an extension of this technology.

B. 6. The Future of Thermal IR Detectors

The next considerations are coming from SPIE, OE Reports, Nr 184, April 1999.

At present, four main IR detector technologies claim the lion's share of global development: mercury-cadmium-telluride (HCT or MCT), indium antimonide InSb, quantum-well IR photodetectors (QWIP) and microbolometer / pyrometer FPAs (Focal Plane Arrays). Development within these different techniques takes many tracks, from the enlarging of array sizes to the reducing of pixel dimensions; from the movement of hybrid bump-bonding designs (cfr fig.A20) to the unity of monolithic detectors and expanded spectral sensitivity in multiple colors, e.g., **bi-color or bi-spectral FPA, combining either MWIR with LWIR or NIR (near IR) with LWIR.**

InSb, which has been around for more than 40 years, represents a mature technology for imaging MWIR ($3\mu\text{--}5\mu$). "Either InSb or MCT is highly reproducible right now", said Paul Norton of Raytheon. Later this year (1999), Raytheon will release the first 2052×2052 InSb FPA for a consortium of astronomers, giving them four times the output of the previous 1024×1024 arrays.

According to Kozlowski, Rockwell is working with Boeing on dual-color devices for two-color uncooled NIR and LWIR imaging. Using InGaAs arrays for NIR and 220×340 -pixel microbolometer arrays sensitive to 24mK for the 8μ to 14μ spectrum, these systems offer something special for military defense systems.

"Two-color is starting to become a very big area. With ships the short band (NIR) lets you see missiles at long ranges. A second, longer spectral band (LWIR) lets you discriminate between missiles and flares or other countermeasures", Kozlowski said.

According to microbolometer expert Paul Kruse, cutting-edge microbolometers are striving to reduce the current pixel size from 50μ square to 25μ . Parallel to this is the DARPA goal of achieving a noise equivalent temperature difference (NETD, sometimes called temperature resolution) of 10mK. "It is very hard to combine a smaller pixel size with a lower NETD because NETD goes inversely with area", explained Kruse. Notice that the specific detectivity D^* has therefore been defined as

$$D^* = (A \cdot \Delta f)^{1/2} / \text{NETD}$$

where A is the detector's area (always in cm^2) and Δf is the processing bandwidth (in Hz).

Researchers are taking several different tracks towards achieving more sensitive detectors, including switching from different combinations of the active materials in the vanadium oxide (VOx) group sputtered on a silicon nitride microstructure layer, to photo-enhanced chemical vapor deposition of amorphous silicon. Other exciting areas include the high temperature resistance coefficient of colossal magneto-resistance films and new structural designs using barium-strontium-titanate (a ferro-electric material) active-area pixel arrays.

Next to the interesting material development in microbolometers lies the development of QWIP technology based (see fig.A17) on GaAs / AlGaAs heterostructures. Using MBE (Molecular Beam Epitaxial) growth on large wafers, H.C.Liu of the National Research Council Canada (Ottawa) expects that QWIP technology will lead to the most easily mass-produced IR imaging devices in the mid- to longwave IR.

Furthermore QWIPs enable an easy manufacturing of bispectral FPAs.

Both the Jet Propulsion Laboratory (Pasadena, CA) and Lockheed Martin (Sanders, NH and Orlando, FL) have two-color focal plane arrays based on QWIP designs in the 3- to 5- and 8- to $12\mu\text{m}$ spectrum. Although QWIP designs promise to eventually reach out to 20 or $30\mu\text{m}$, most QWIP arrays require extensive cooling down to 70K, although recent developments at Thompson-CSF show that the devices can operate at the higher temperature of 90K.

Liu expects that exploration into combining QWIPs with GaAs multiplexing circuits will increase the utility of QWIP technology by placing both detector and readout circuit on the same chip.

Another solution, being developed at the NRC Canada, involves attaching LEDs directly to the

QWIP active area and optically reading the output with a separate charge-coupled device (CCD) chip.

Interestingly, work on IR imagers has also lead to improved visible CMOS imagers. According to Kozlowski, Rockwell's work on IR imagers at Newport Beach, CA has led to the recent spin-off of a new company that has a CMOS imager of comparable or better performance than off-the-shelf CCD imagers. According to Kozlowski, the newcomer, Conexant (Newport Beach, CA), has recently announced a visible CMOS camera on a chip with 1280×1024 pixels that has a read noise below 10 e and a sensitivity of 1.5V per lux second with a dark current of 50 pA/cm^2 .

Part B

The SAR (Synthetic Aperture Radar)

B. 1. Introduction

The SAR is chiefly a pulse SLAR (side looking airborne radar)-see fig.B1-performing ground imaging and achieving much better resolution (smaller pixel) on the ground through a sophisticated signal processing. The SAR is either space-borne (aboard a satellite) or airborne (aboard an aircraft), i.e., the SAR platform is either a satellite or an aircraft.

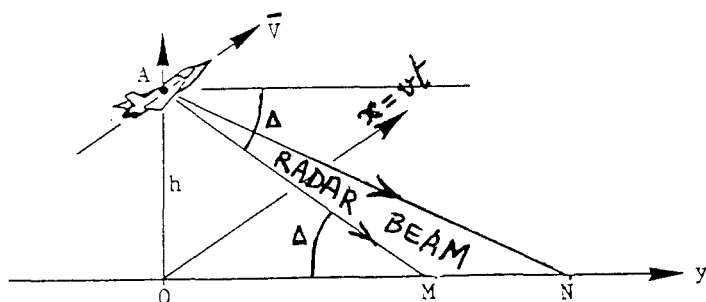


Fig. B1

SLAR with the depression angle Δ .

Motion of the SAR with respect to the ground is mandatory. The SAR requires recording of both amplitude and phase ; therefore the SAR can be claimed as being a holographic radar.

The SAR is an imaging radar achieving a pixel size comparable to that of the best EO (electro-optical) sensors which are camera's using visible and near IR radiation or thermal imagers, but the SAR is much more expensive than those EO sensors. Why then using a SAR ? Because the SAR performs much better at night, during bad weather conditions and, especially, through clouds, although the heavy signal processing does not always enable the SAR to deliver images in real time.

Indeed, Fig. B2 demonstrates clearly that EO sensors are ineffective through cloud cover, while radar sensors have good to superior performance through cloud cover and rain. EO transmission through rain is a function of the size of the raindrops, rainrate and the path length through the rain. EO passive sensors are limited from about 2 to 5 km of path length through rain. The best sensors in looking through clouds and rain are radar sensors. Radar sensors have negligible attenuation at frequencies below 10 GHz. That's the reason why most SAR's are working in the following radar bands :

L (1 - 2 GHz), S (2 - 4 GHz), C (4 - 8 GHz), X (8 - 12 GHz).

At higher frequencies, millimeter wave sensors operating in clouds and rain are limited from about 2 to 5 km length of path through the clouds and rain, with the same implications as those discussed for the EO sensors (both passive and active, like the laser rangefinder).

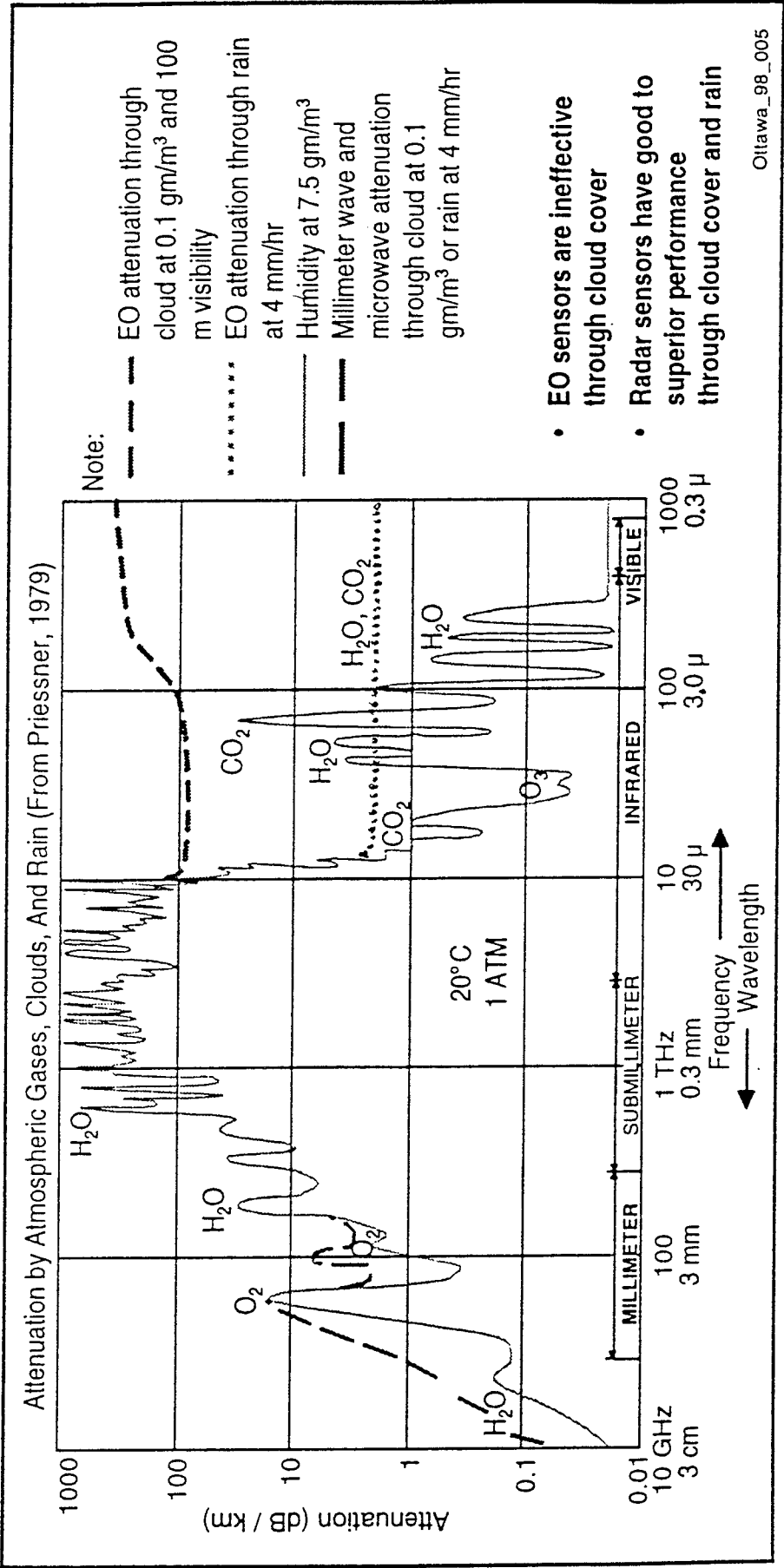


Fig. B2. Signal Attenuation Due to Weather (RTO-MP-12, AC/323 (SCI) TP/4)

The superiority of the radar (with respect to the EO sensors), which is already obvious in the presence of rain, becomes enormous in the presence of clouds. Indeed, for a 1 km length of path through a dense cloud, the 2-way (round trip) attenuation

- is 0.2 dB below 30 GHz
- is more than 200 dB for near IR and visible radiations.

Figure B3 (showing a space-borne SAR aboard the Space Shuttle) reviews the basics of SAR image formation, stressing the necessity of the platform motion for the building up of a radar image, and introduces the two image dimensions as being

- range (across-track)
- azimuth (along-track) with azimuth coordinate x . Let us now consider the image pixel, i.e., the smallest resolvable object on ground.

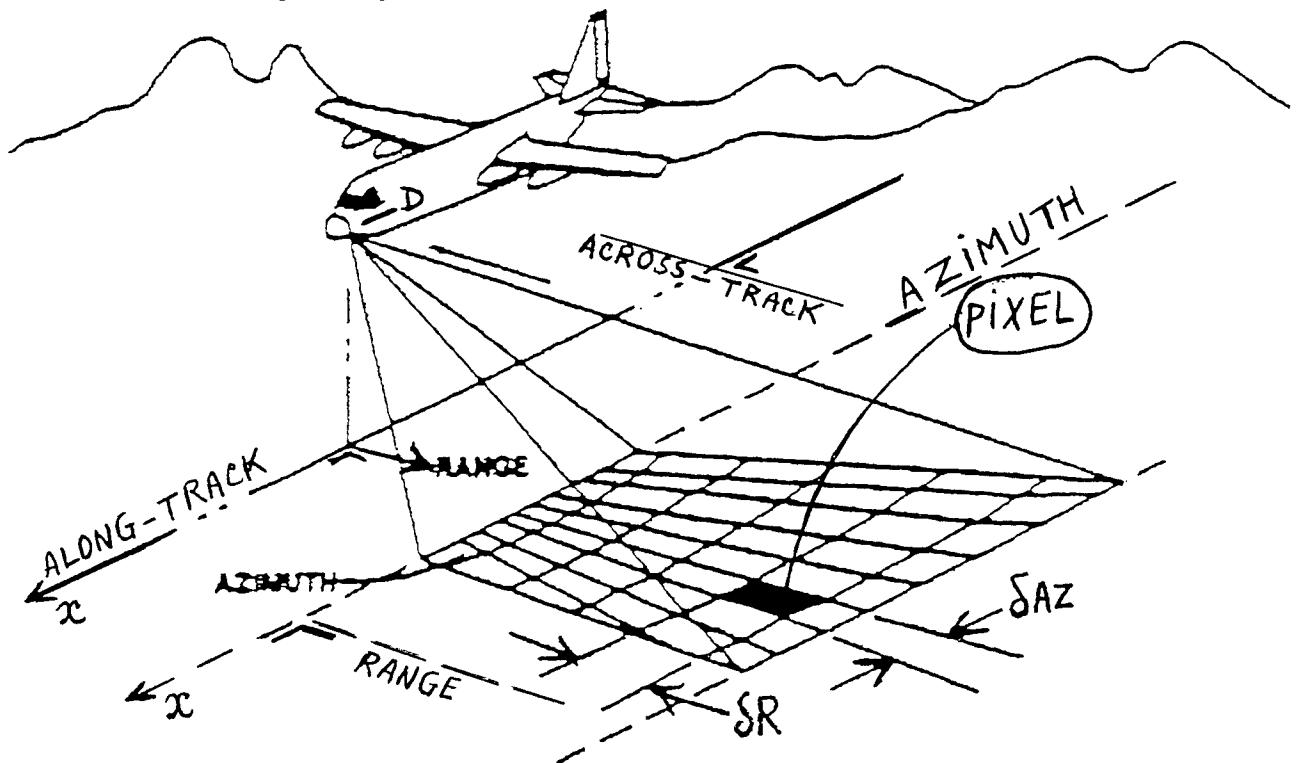


Fig. B4. The pixel of the radar image on ground with size $\delta R \cdot \delta AZ$ (pixel = smallest resolvable object)

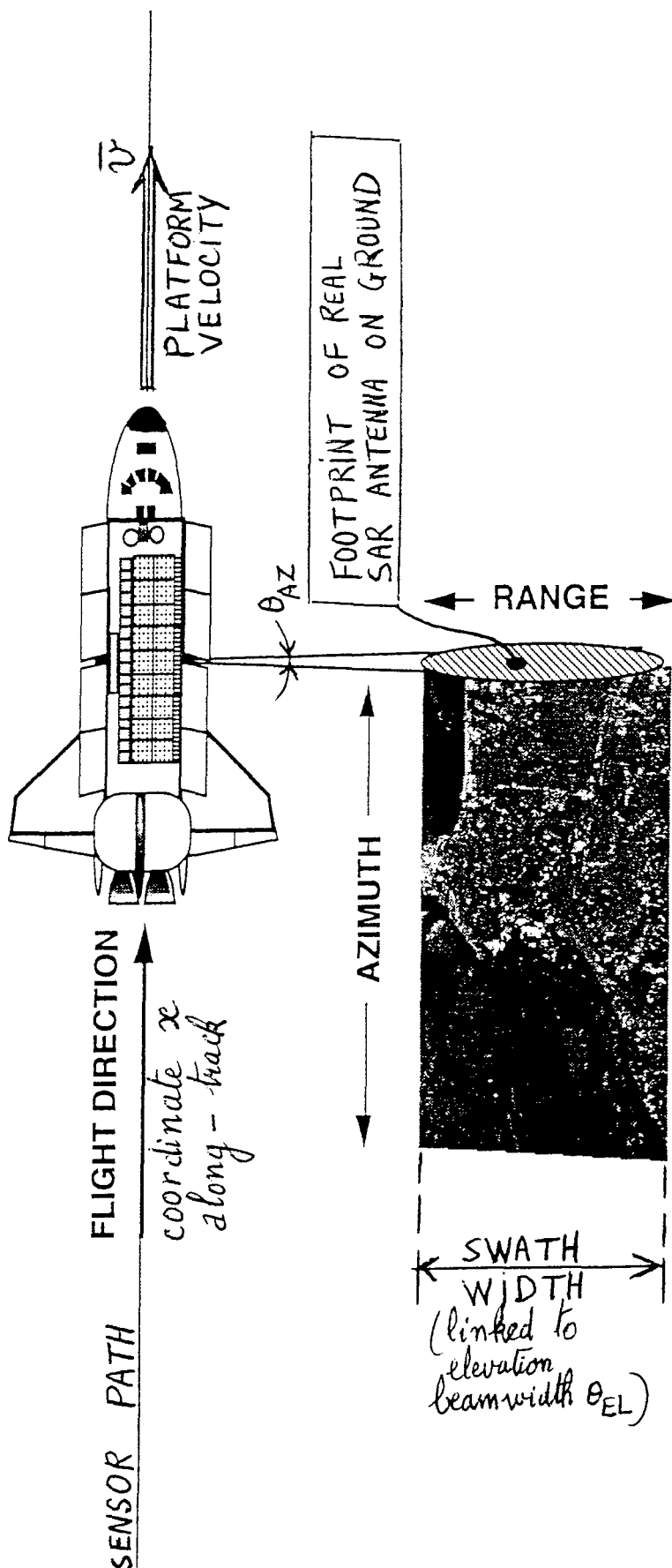
To understand the necessity of the platform motion and the subsequent signal processing of a SAR, it is of utmost importance to evaluate the size $\delta R \cdot \delta AZ$ of the radar image pixel on the ground where (see fig. B4)

- δR is the range resolution (on ground)
- δAZ is the azimuth resolution (on ground).

The radar image pixel of fig. B4 is nothing else but the footprint on the ground of the resolution cell (fig. B5) determined by the radial resolution $c \cdot T/2$ (where T is the pulse width) and the two angular resolutions θ_{AZ} and θ_{EL} (being respectively the 3-dB beamwidths in the azimuth and elevation planes) as sketched by fig. B5

FIG. B3. REVIEW OF SAR IMAGE FORMATION

- BUILDING UP A RADAR IMAGE USING PLATFORM MOTION



- RADAR BEAM ILLUMINATES A SWATH ON THE GROUND ($SWATH \equiv FAUCHEE$ in French)
- IMAGE DIMENSIONS ARE RANGE (ACROSS-TRACK) AND AZIMUTH (ALONG-TRACK)

(the volume of this resolution cell being $cT/2 \cdot R\theta_{AZ} \cdot R\theta_{EL}$ which is proportional to the square of the range R)

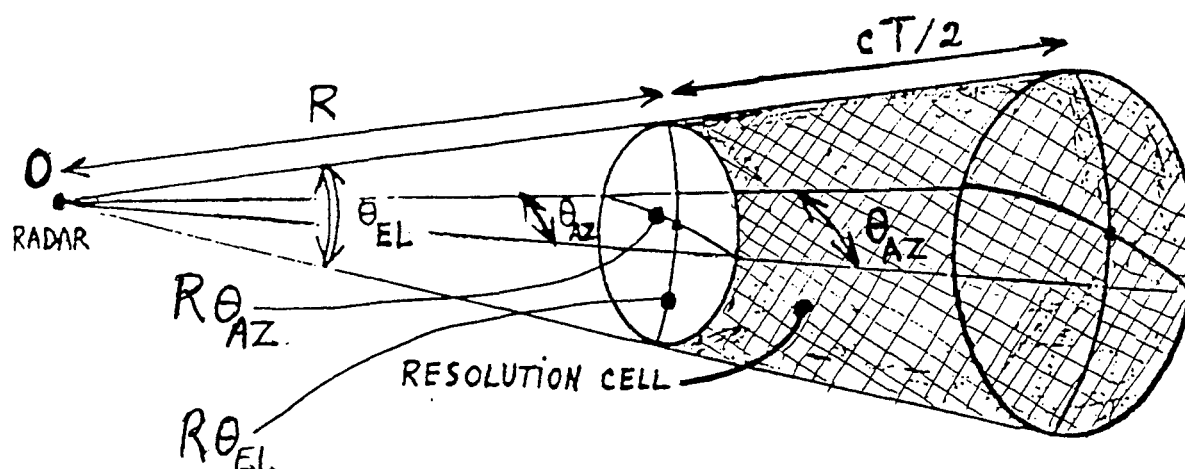


Fig B5. The resolution cell with its radial size $cT/2$ (where T is the pulse width) and its transverse sizes $R\theta_{AZ}$ and $R\theta_{EL}$

Let us first calculate the ground image pixel sizes δR and δAZ for a conventional SLAR, i.e., a radar which does not implement the sophisticated signal processing of a SAR. Therefore we introduce (see fig. B6) the coordinates used in the SAR theory, i.e., the x coordinate along track and the (slant) range R whose minimum value R_0 is the range at closed approach, i.e., when the target is purely broadside ; fig. B6 shows the azimuth angle AZ as well.

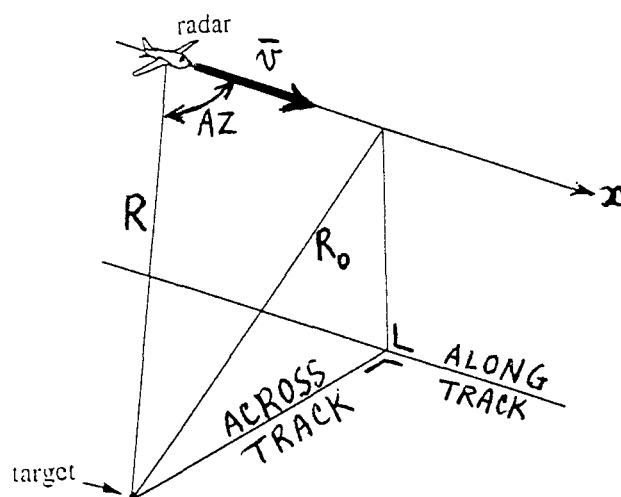


Fig. B6. SAR co-ordinate system

The range resolution δR on ground, i.e., the smallest resolvable object across track, is derived from the radial resolution $\frac{cT}{2}$ by using fig. B7 which shows immediately that $\delta R = \frac{cT}{2 \cos \Delta}$ where Δ is the depression angle ; Δ is also the grazing angle.

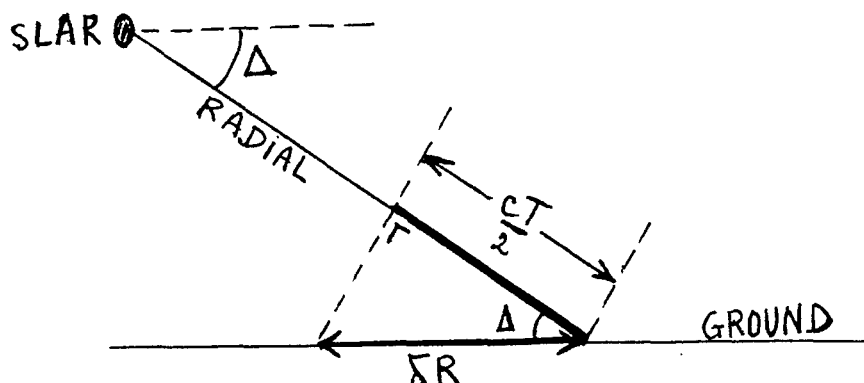


Fig. B7. Derivation of the range resolution on ground

To compute the azimuth resolution δAZ on ground, i.e., the smallest resolvable object along track, it suffices to calculate the transverse size $R_0 \theta_{AZ}$ of the resolution cell as shown by fig. B8.

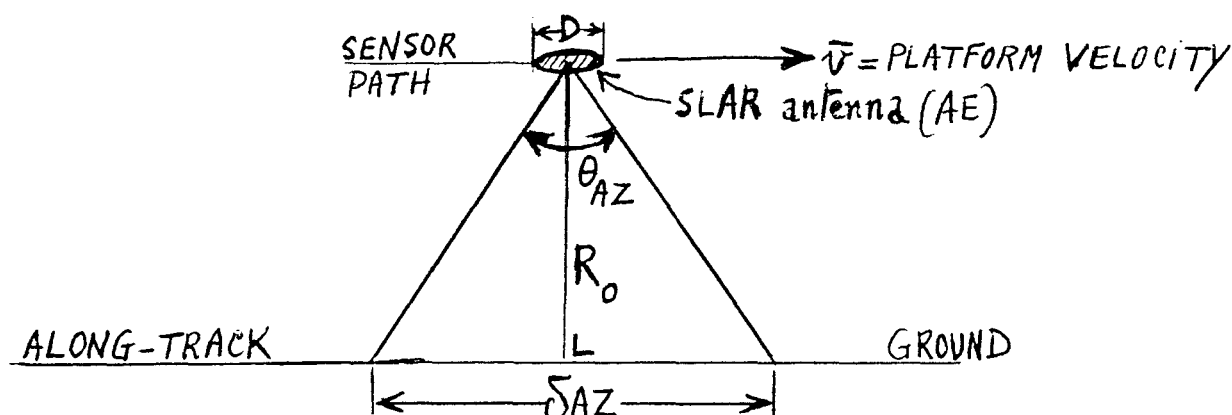


Fig. B8. Derivation of the azimuth resolution on ground ; notice that δAZ is the footprint along track of the AE radiation pattern on ground.

The real aperture length D is nothing else but the size of the real SLAR antenna along track (i.e., along the platform velocity). The corresponding real aperture (azimuth) beamwidth θ_{AZ} is then given by

$$\theta_{AZ} = \lambda/D$$

where λ is the wavelength of the SLAR.

Since θ_{AZ} is a small angle (a few degrees), δAZ is nothing else but the corresponding transverse size of the resolution cell and is given by

$$\delta AZ = R_0 \theta_{AZ} = R_0 \lambda/D$$

The latter result is a dramatical one because the azimuth resolution δAZ

- increases with the range R_0
- is unfortunately proportional to the wavelength.

Unfortunately the wavelength of a radar sensor is much bigger than the wavelength of an EO sensor. Considering fig. B2, the ratio $\lambda_{\text{radar}}/\lambda_{\text{EOsensor}}$ is of the order $3 \text{ cm}/3\mu = 10,000$ making thereby the azimuth resolution 10,000 times worse for a radar sensor.

For instance, for a 10 m antenna ($D=10\text{m}$), operating (case of a spaceborne radar sensor) at 800 km range ($R_0 = 800 \text{ km}$) at $\lambda = 25 \text{ cm}$ (P band), the formula of δAZ gives a minimum object size of 20 km, while this minimum object size would be of the order of 2 m for an EO sensor.

Conversely, the range resolution δR would be at least

- 150 m for a pulse width $T = 1\mu\text{s}$
- 3 km for a pulse width $T = 20\mu\text{s}$, because $\delta R > cT/2$

Clearly, a radar sensor necessitates a drastical reduction of both sizes δR and δAZ of the pixel on ground (the case being particularly dramatic for δAZ) in order to achieve images of the same quality as those given by the EO sensors.

The solution is quite obvious for the reduction of δR : it suffices to compress the pulse width to a value τ in the receiver as done by the PC (Pulse Compression) radar using a pulse compression ratio

$$\rho = T/\tau \approx B.T$$

which is nothing else but the well known time-bandwidth product. Values of $\rho = BT$ of the order of several thousands are quite feasible. Transmitting directly the much narrower pulse width τ (in order to bypass the range compression) is not feasible because it would require from the SAR a much higher peak power which would also increase the detectability by the enemy ELINT.

For δAZ the harmful effect comes from the large wavelength λ of all radar sensors. This very big value of λ could be compensated (in view of the formula $\delta AZ = R_0\lambda/D$) by replacing the real antenna of aperture length D by a very big linear array (see fig. B9) of aperture length L such that $L/D \approx 10,000$. Such a big linear array is evidently not feasible and totally unrealistic because such an array would be much bigger than the platform ! The solution will consist in replacing the very big linear array by a corresponding synthetic aperture with the same aperture length L as shown by fig. B9.

Assuming that the radar platform moves in a straight path, the synthetic aperture length L is (cfr fig. B10) nothing else but the path AB travelled by the radar platform during the dwel time (or integration time or illumination time)

$$T_S = t_{\text{out}} - t_{\text{in}} = L/v$$

where t_{in} is the instant where the target enters (point A) the real aperture beam and t_{out} is the instant where the target leaves (point B) the real aperture beam of beamwidth θ_{AZ} .

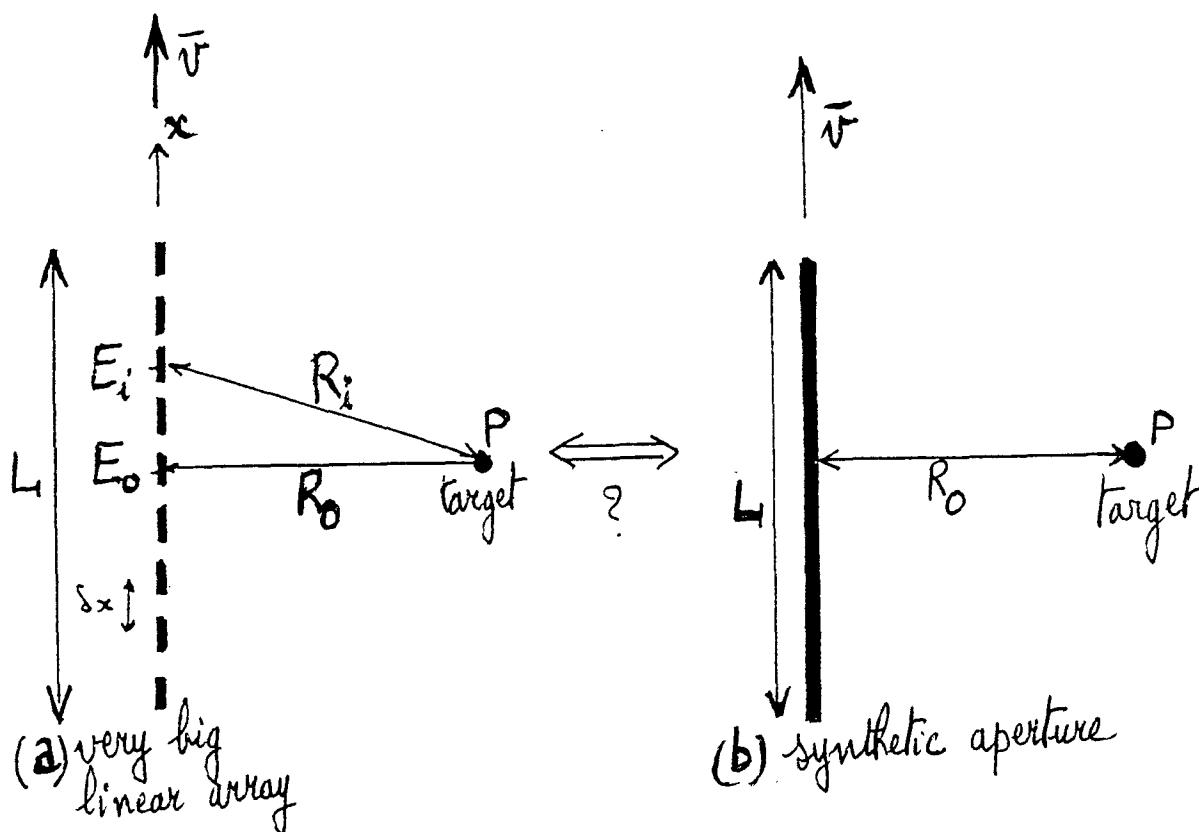


Fig. B9. Very big linear array and corresponding synthetic aperture of aperture length L

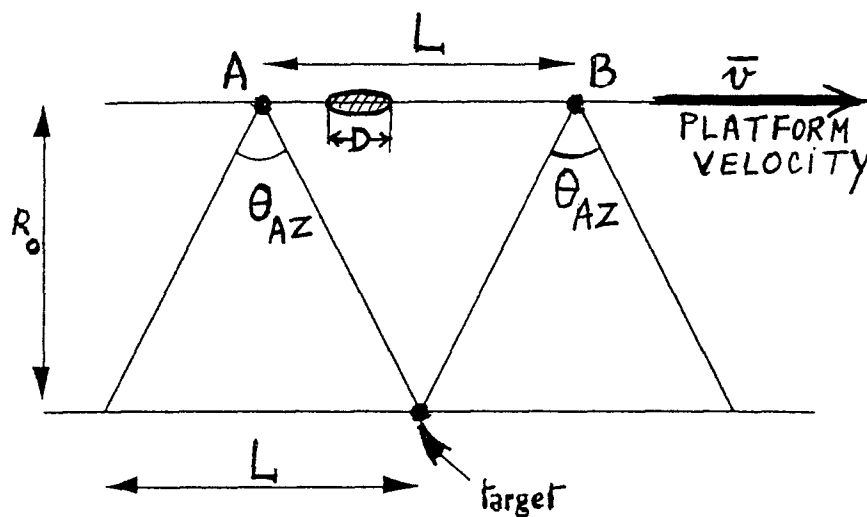


Fig. B10. Maximum synthetic aperture length L

Evidently, it is clear from fig. B10 that the (maximum) synthetic aperture length L depends on the real aperture beamwidth θ_{AZ} , which in turn depends on the real aperture length D because $\theta_{AZ} = \lambda/D$ and (since θ_{AZ} is a small angle)

$$L = R_o \theta_{AZ} = R_o \lambda / D$$

If this synthetic aperture of length L can be synthesized, the corresponding synthetic beamwidth would be

resulting in a much smaller azimuth (along-track) resolution (i.e., smallest resolvable object along-track) given by

$$\delta AZ = R_0 \theta_s = R_0 \frac{\lambda}{L} = \frac{R_0 \lambda}{R_0 \lambda} \cdot D = D$$

which is nothing else but the real aperture length along-track and is surprisingly range independent. The previous result has been obtained by a crude calculation. Accurate (exact) calculations show that the SAR performs even better and that the smallest resolvable object along-track (or azimuth resolution of the SAR) is given by

$$\delta AZ = D/2$$

the half of the real aperture length along track. Figures B11 and B12 summarize the previous results successively for a conventional SLAR (no SAR-like signal processing) and for a SAR, T being the transmitted pulse width and τ the compressed pulse width.

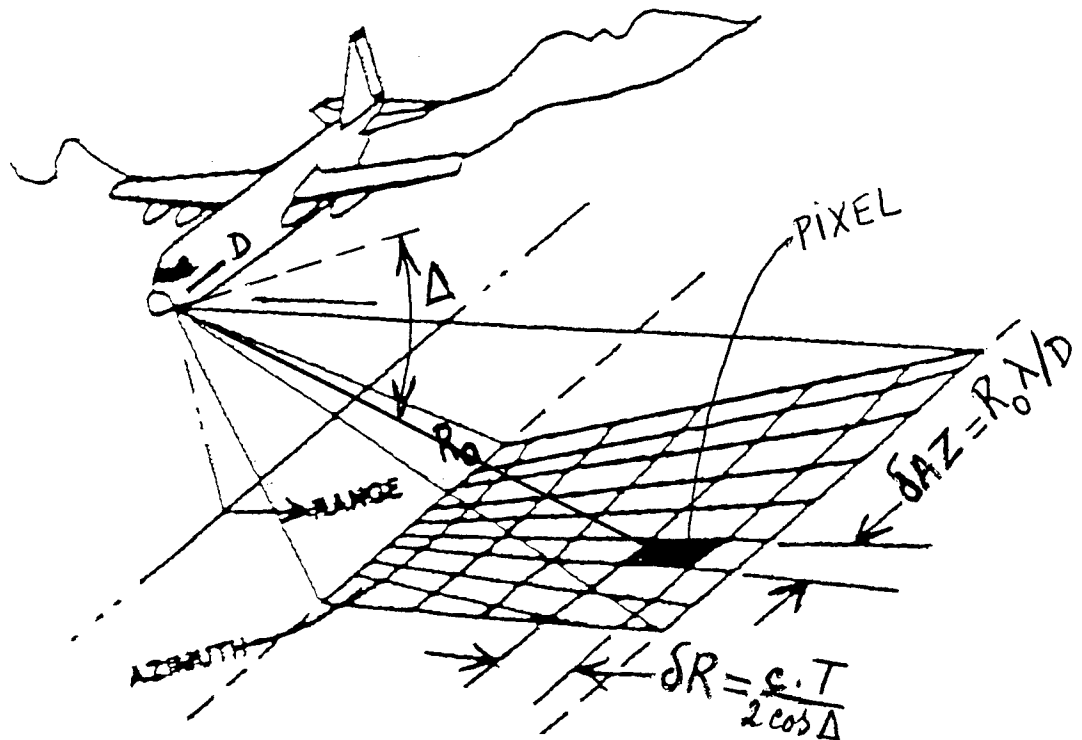
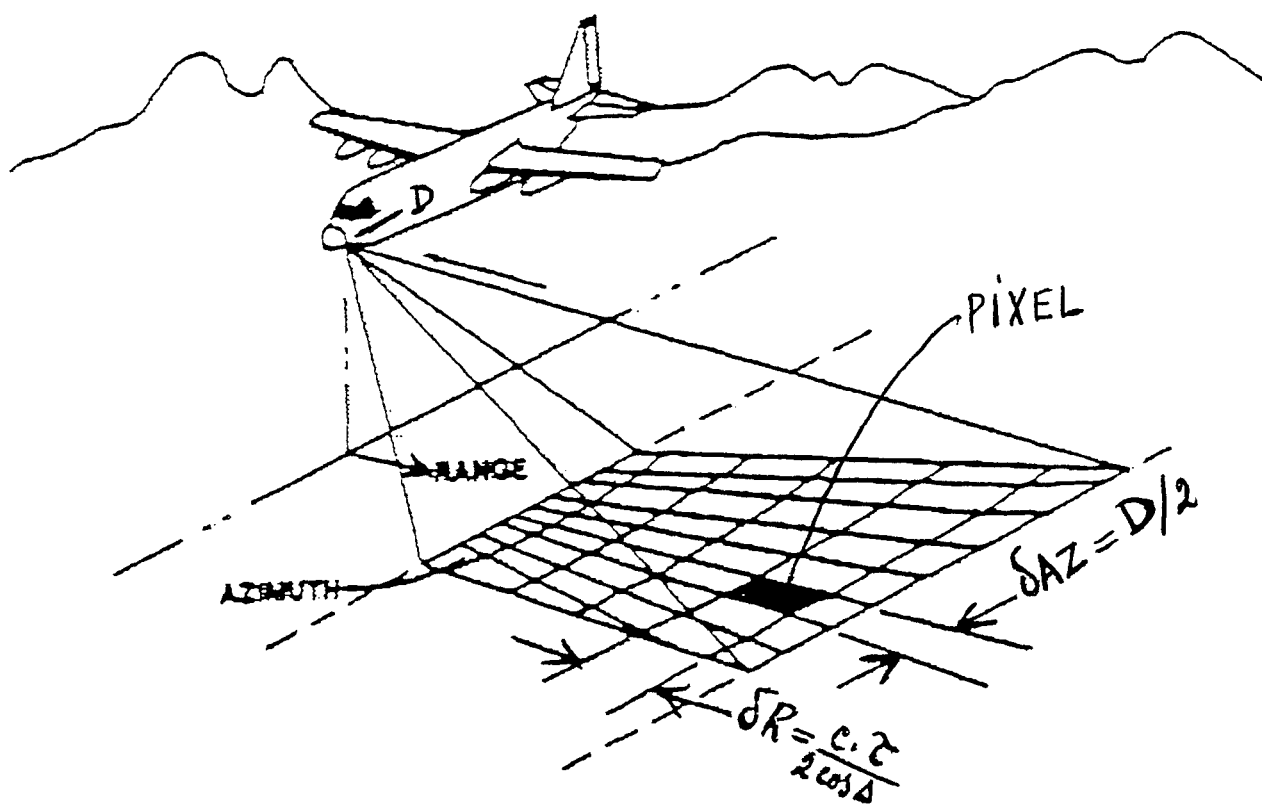


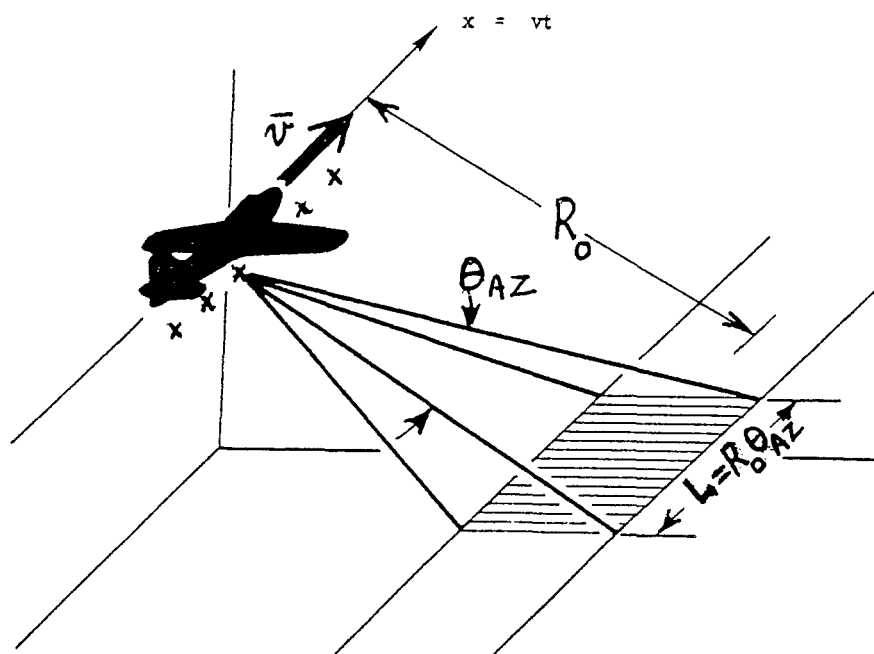
Fig. B11. Conventional SLAR pixel ; Δ is the depression angle

The SAR signal processing reducing the range resolution δR from $\frac{cT}{2\cos\Delta}$ to $\frac{c\tau}{2\cos\Delta}$ is termed **range compression** (which is exactly the pulse compression of a PC radar). The SAR signal processing reducing the azimuth resolution from $R_0 \lambda / D$ to $D/2$ is termed **azimuth compression** ; the latter achievement is entirely due to the synthetic aperture which we shall discuss now in more details.

Notice that the platform straight path assumption is not a very crucial one and that target motion corrections (with respect to a straight path) are small.

Fig. B12. SAR pixel

B. 2. The Synthetic Aperture

Fig. B13. Sidelooking radar configuration with the synthetic aperture length L

The principle of synthetic aperture processing is to realize the azimuth resolution appropriate to a large linear array (cfr fig. B9. a) by synthesizing the aperture of length L (see fig. B13) sequentially using just one radiating element (which is nothing else but the real SAR antenna of length D) moved along the whole synthetic aperture rather than instantaneously with all the signals available simultaneously in parallel. That means that the role of the various radiating elements E_i of the big linear array of fig. B9 (a) is played by the sequential positions (indicated by the crosses on fig. B13) of the single real SAR antenna at the successive instants of the transmitted pulses, provided the SAR records and stores both amplitudes and phases of the echoes at those sequential positions. This means that at each position indicated by the crosses on fig. B13, the single real antenna radiates a pulse, then receives and stores amplitude and phase of the reflected signal (called echo). These stored data are then processed in a manner analogous to the coherent weighted summation carried out in a large linear array. We remind that the focusing operation (it means the build up of the high directivity) of a large linear array represents a phase adjustment of the signal received on each radiating element E_i (of fig. B9. a) of the array so that, in the summation process, the contributions from all array elements E_i are combined in phase.

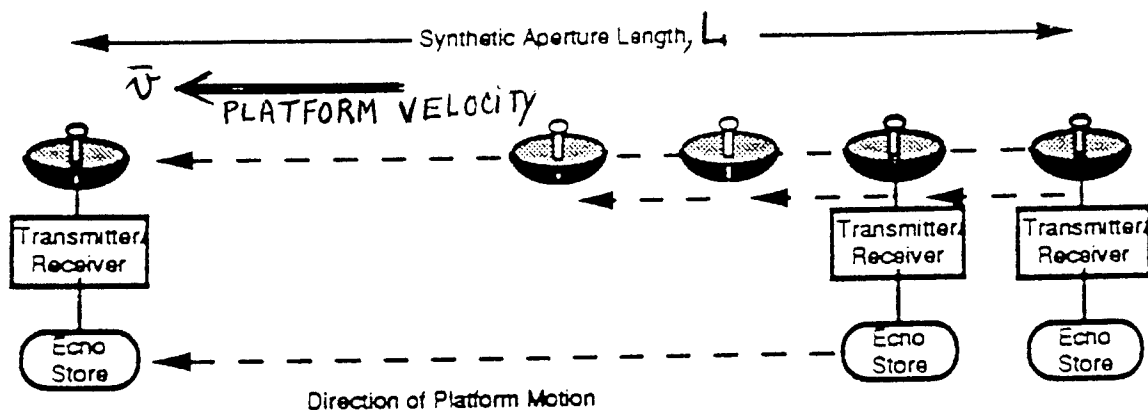


Fig. B14. The sequential positions (at the instants of the transmitted pulses) of the single real SAR antenna

Fig. B14 summarizes quite good the SAR principle :

- the larger the antenna (the larger the length L) the finer the details the radar can resolve
- in a SAR, one synthesizes a very long antenna (of length L) by combining the echoes received by the single real antenna as it moves along the targets.
- at each position of fig. 14 a pulse is transmitted, then return echoes pass through the receiver and are recorded (in amplitude and in phase) in an " echo store", for later processing into a SAR image.

Fig. B15 gives the (slant) ranges corresponding to the sequential positions of the single real antenna. At position E_i , the echo phase is proportional to $2R_i$ (because of the round trip).

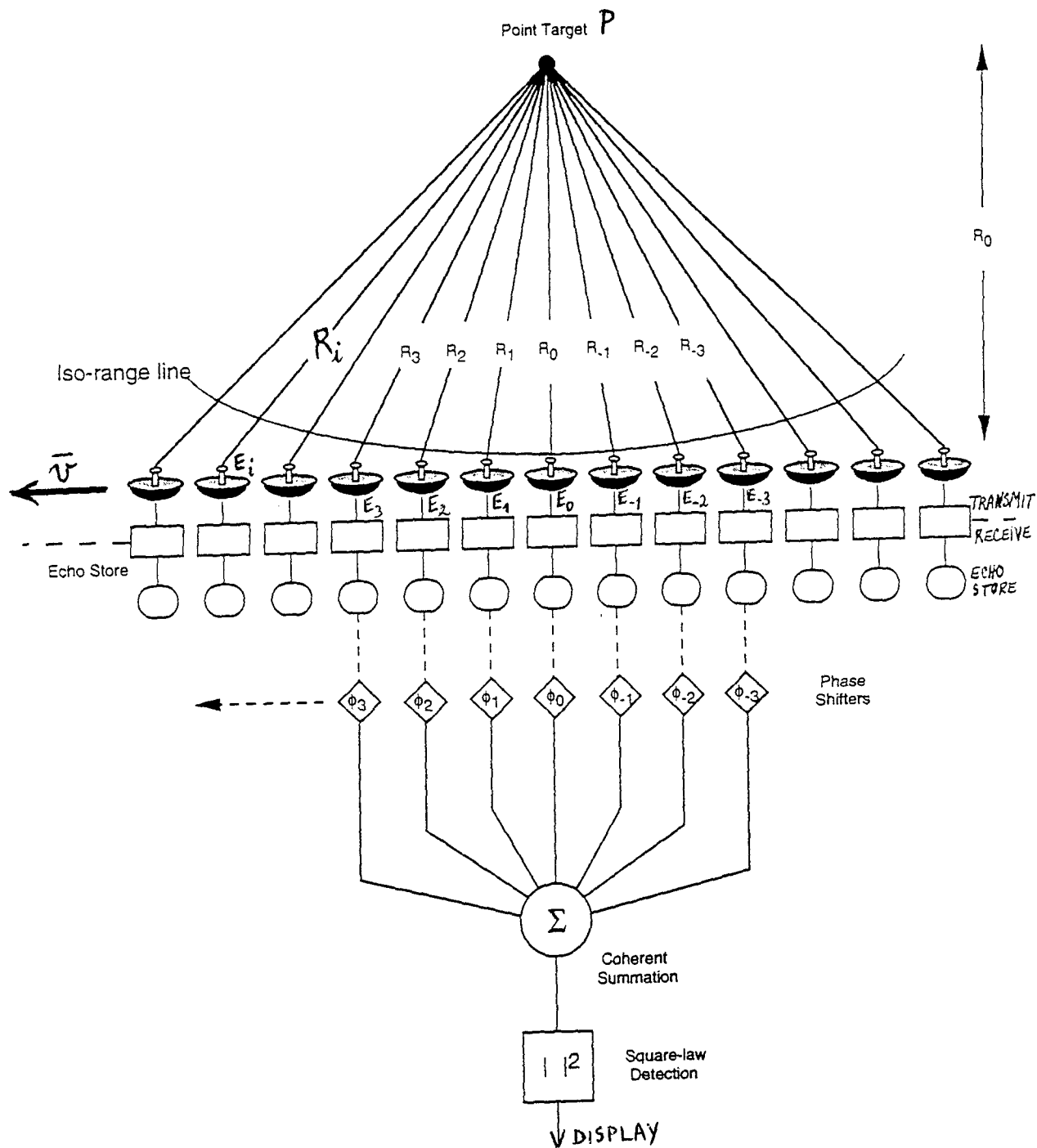


Fig. B15. Forming a SAR image of a point target by phase correction ; the iso-range line sketches in fact a diverging lens whose focal length is linked to R_0 , the radius of curvature of the spherical diopter of this diverging lens

The principle of synthetic aperture is to combine the sequence of echoes (received at the sequential positions E_i of fig. B15) coherently - coherently means adding all echoes from the target P after compensating (i.e., matching) all phases - like a big linear array to get the effect of a large along-track antenna, and hence azimuth resolution.

The process of matching the phases of the echoes can be thought of as a frequency-domain matched filtering processing. In this respect it is very similar to matched filtering of a chirped pulse in a radar pulse compression (performed in the compression filter of a PC radar), and for this reason it is often referred to as azimuth compression, or azimuth correlation.

Let us summarize the SAR signal processing as follows :

1. The range compression (which is a pulse compression) is obtained by a matched filter, or, equivalently, by a correlation with a range reference signal which is a replica of the expected return from a point target at the same range.
2. Conversely, the azimuth compression is obtained by a correlation with an azimuth reference signal which is again a replica of the expected return from a point target at the same range.
3. Since (as we see later on) the azimuth reference signal is range dependent, range compression must be performed before azimuth compression.
4. Those 2 subsequent compressions are summarized, for a point target, by fig. B16 showing at each step the obtained improvement of the resolution materialized by the corresponding reduction of the pixel.

The final range resolution shown on fig. B16 corresponds to its minimum value $c\tau/2$ obtained for a depression angle $\Delta=0$. Since, for a PC radar, the compressed pulse width τ is related to the radar bandwidth B by

$$B \approx \frac{1}{\tau}$$

the optimistic value of the final range resolution (i.e., the smallest resolvable object across track) is given by

$$\delta R = \frac{c}{2B}$$

We remind the amazing final azimuth resolution (i.e., the smallest resolvable object along track) as being

$$\delta AZ = \frac{D}{2}$$

where D is the aperture length along track of the real antenna.

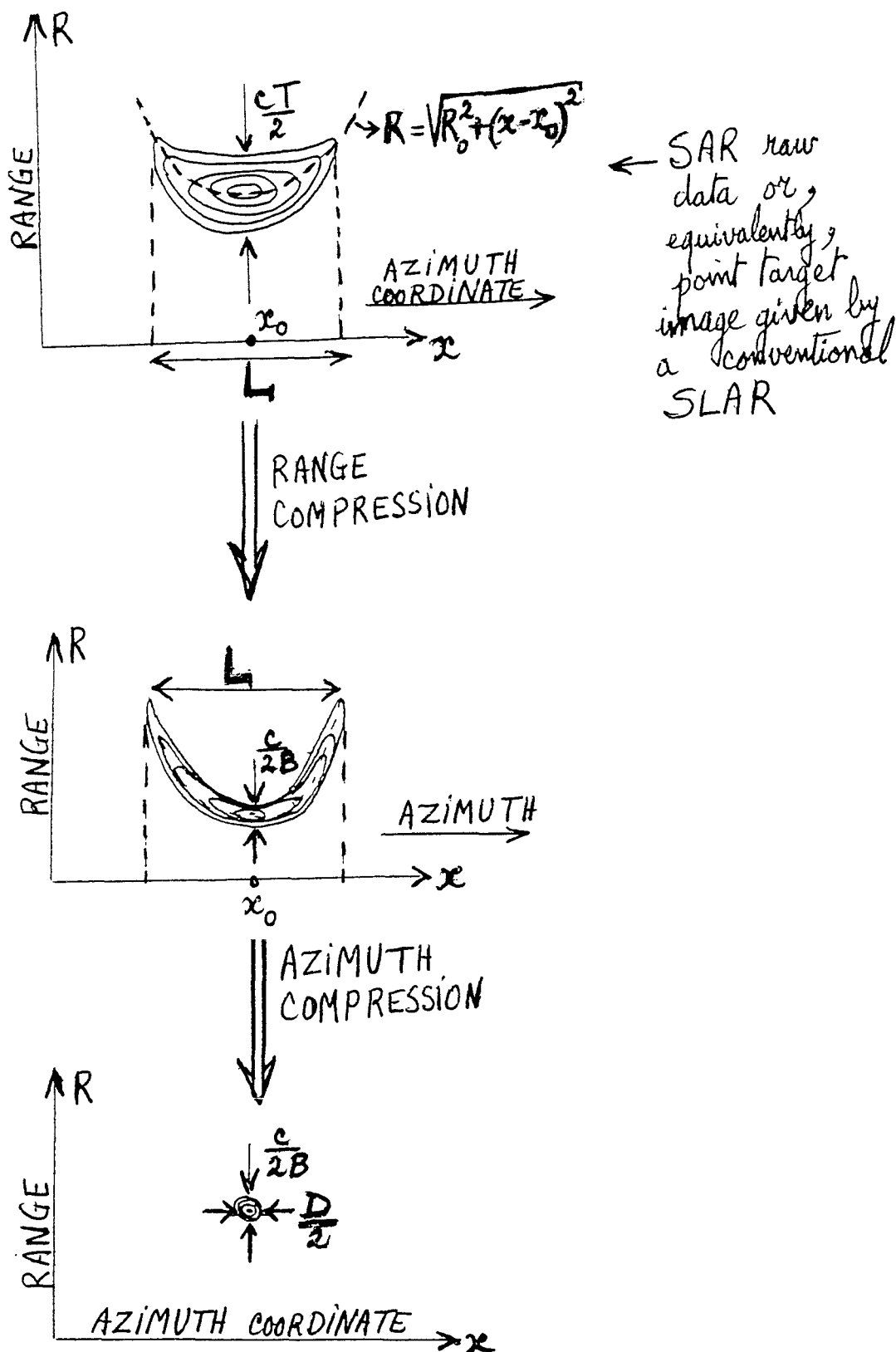


Fig. B16. The 2 successive compressions of a point target performed by the SAR signal processing. L is the synthetic aperture length but also the ground footprint along-track of the real antenna beam. The value $c/2B$ of the range resolution corresponds to the optimistic case of a depression angle $\Delta = 0$

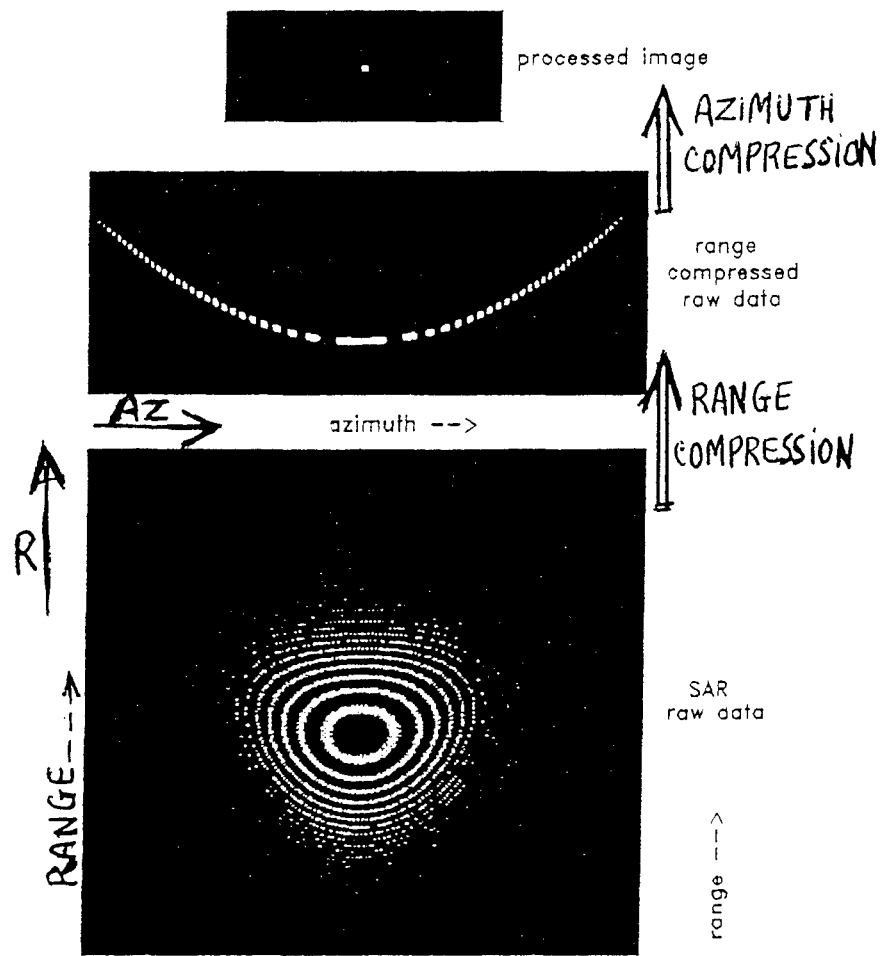


Fig. B17. The raw data received from a point target (lower image) show as well as the range compressed data the influence of range curvature

Fig. B17 illustrates the two successive compressions performed by a SAR on the raw data received from a real point target. The range curvature of fig.B17 is also called range migration and corresponds to the curve

$$R = \left[R_0^2 + (x - x_0)^2 \right]^{1/2}$$

of fig. B16. The range migration curve described by the last formula gives the variation of the (slant) range of a particular point target when the radar platform moves along that target.

As we are going to see, the azimuth compression is based on the Doppler shift (or Doppler frequency) f_D due to the relative motion of the target with respect to the SAR and given by

$$f_D = \frac{2v_R}{\lambda}$$

where v_R is the radial component of the velocity of the target with respect to the SAR.

We remind that

- v_R and f_D are positive for a closing target
- v_R and f_D are negative for an opening target

B. 3. Basic Hardware Considerations

A simplified block diagram for a SAR is shown in fig. B18 assuming a digital signal processing ; indeed the previous optical processing has almost disappeared although it could come to a new life : electro-optical processing (using light valves) is again used by a US compagny performing almost real time processing ; nevertheless, presently, the vast majority of SAR's use digital processing for both range compression and azimuth compression.

The SAR is mounted on a platform moving at a constant velocity. The PRF must be sufficiently high to avoid azimuth ambiguities which are nothing else but velocity ambiguities. This criterion requires that the radar platform displacement cannot exceed one-half the real antenna size between successive transmit pulses;

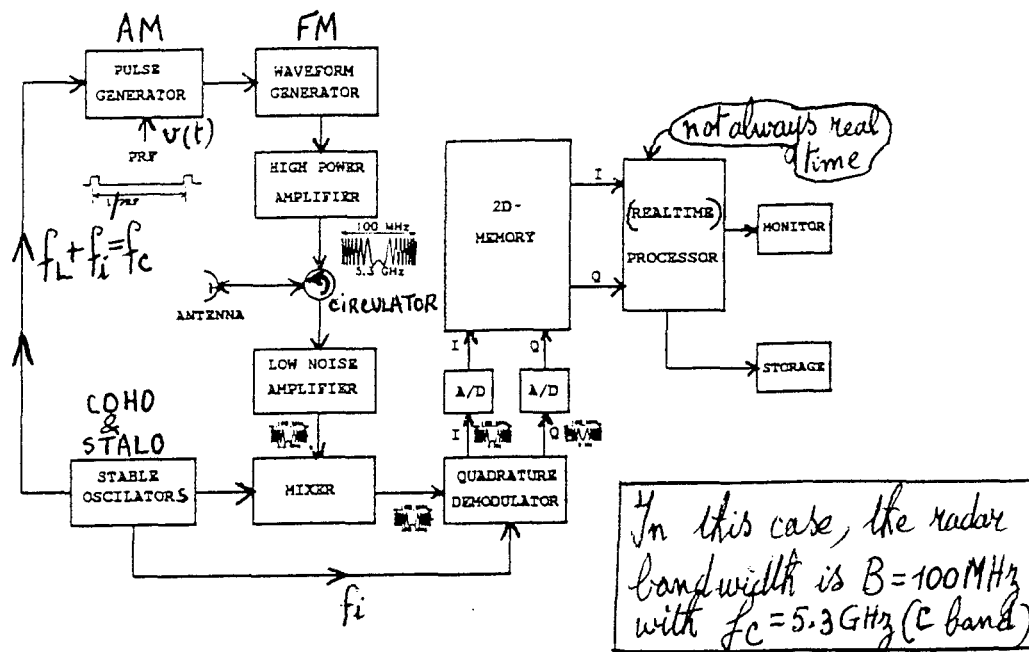


Fig. B18. Block diagram of a SAR (schematically) [W.KEYDEL]

in other words, if v is the magnitude of the SAR platform velocity,

$$v \cdot T_R = v \cdot IPP = \frac{v}{PRF} \leq \frac{D}{2} \quad \text{or}$$

$$PRF \geq \frac{2v}{D}$$

or

$$PRF \geq \frac{v}{D/2}$$

The latter condition is a constraint on the PRF which can easily be demonstrated as follows. It is very easy to show that the Doppler shifts f_D at the extremes (i.e., when the target enters or leaves the real antenna beam) are just $\pm v/D$.

$$\text{So } f_{D_{\max}} = v/D$$

Since velocity ambiguities are avoided if

$$f_{D \max} \leq \frac{\text{PRF}}{2}$$

this requires that

$$v/D \leq \text{PRF}/2$$

or

$$\text{PRF} \geq 2v/D$$

Considering the block diagram of fig. B18, it should be stressed that

- (a). In a SAR, phase stability is exceedingly important. The prime oscillators which provide the signal (STALO and COHO) for the transmitter as well as the reference (COHO) for the receiver must be very stable. The timing of the transmit pulses must be very precise with respect to the prime oscillators ; in other words the PRF must be derived from the COHO (just as it is done in a PC radar to enable full cancellation of the fixed clutter) by dividing the COHO frequency f_i by an integer number N :

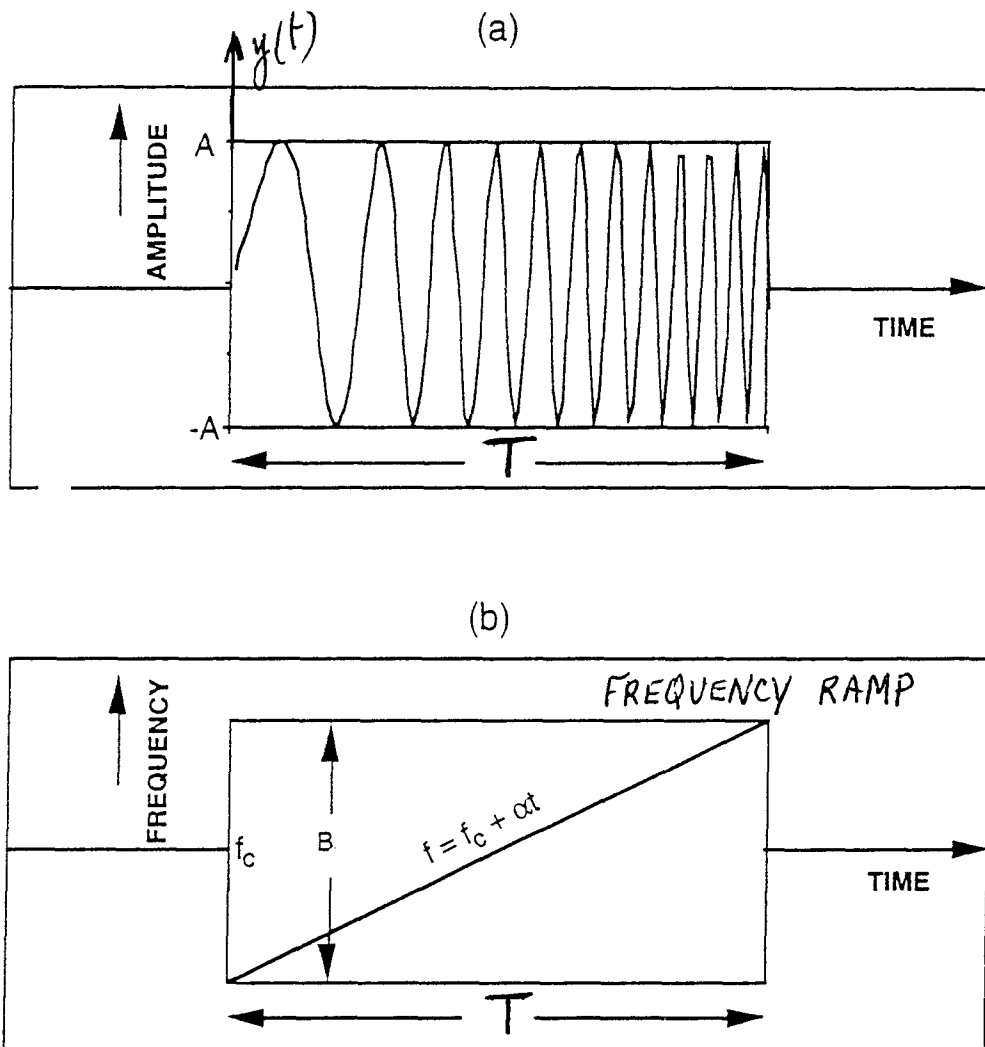
$$\text{PRF} = f_i / N$$

- (b) Electronic circuits using a VCO (voltage controlled oscillator) can provide the desired (chirped) transmit pulse. Indeed, the transmit pulse must include a chirp signal in order to be able to perform pulse compression (which is nothing else but the range compression) in the receiver.

As suggested by the block diagram of fig. B18, the whole SAR signal processing is a digital one. Therefore, both range compression and azimuth compression must be performed in the baseband, i.e., on the bipolar video(s) I (and Q if necessary). That is a noticeable difference between the SAR and the PC radar : unlike the PC radar where the pulse compression is achieved within the IF stages, the SAR performs its pulse (i.e., range) compression in the baseband, i.e., on the bipolar video(s) I (and Q eventually).

The main reason for the baseband digital processing is that a digital processing requires an A/D (Analog to Digital) conversion and thus a sampling of the echo signal ; obviously this sampling is much easier (because of a lower sampling frequency) on the bipolar video(s) than on an IF (intermediate frequency) signal. Fig. B19 displays a typical (up) chirped transmit pulse with carrier frequency f_c and transmitted pulse width T . The corresponding echo signal could be down-converted directly to the baseband by a homodyne detection performed in a PSD (Phase Sensitive Detector) subtracting f_c from the frequency ramp of fig. B19.

SAR THEORY

LINEAR FM OR 'CHIRP' WAVEFORM

$$y(t) = A(t) \exp \left\{ j2\pi \left[f_c t + \frac{1}{2} \alpha t^2 \right] \right\}$$

$\alpha = \text{slope of the frequency ramp}$

$$\text{FREQUENCY, } f = \frac{1}{2\pi} \frac{d\phi(t)}{dt} = f_c + \alpha t$$

TRANSMIT PULSE LENGTH = T

Fig. B19. LINEAR FM OR "CHIRP" WAVEFORM

But usually the down-conversion to the baseband is done from the IF echo signal (IF = intermediate frequency) in a PSD (or two PSD's in case both bipolar video's I and Q are needed) which will subtract the IF frequency

$$f_i = (f_1 + f_2) / 2$$

from the IF frequency ramp. In the previous formula, f_1 and f_2 are the extreme frequencies of the IF frequency ramp ; those frequencies are simply related to the radar bandwidth B and the compressed pulse width by

$$B = f_2 - f_1 = 1/\tau$$

Fig. B20 sketches the whole process of down-conversion from the IF signal to the baseband frequency f_{BB} history, the corresponding baseband phase

$$\phi_{BB} = \int_0^t f_{BB}(t) dt$$

and finally the baseband signal $\cos \phi_{BB}$ which is nothing else but the echo bipolar video corresponding to the chirped transmitted pulse of fig. B19.

The baseband signal $\cos \phi_{BB}$ of fig. B20 will be also the range reference function (displayed by fig. B21) for the correlation operation needed to achieve the baseband pulse compression, i.e., the range compression as a part of the SAR signal processing.

A correlation needs always a reference function. In the case of any radar, the appropriate reference function is a replica of the expected return from a point target at the same range.

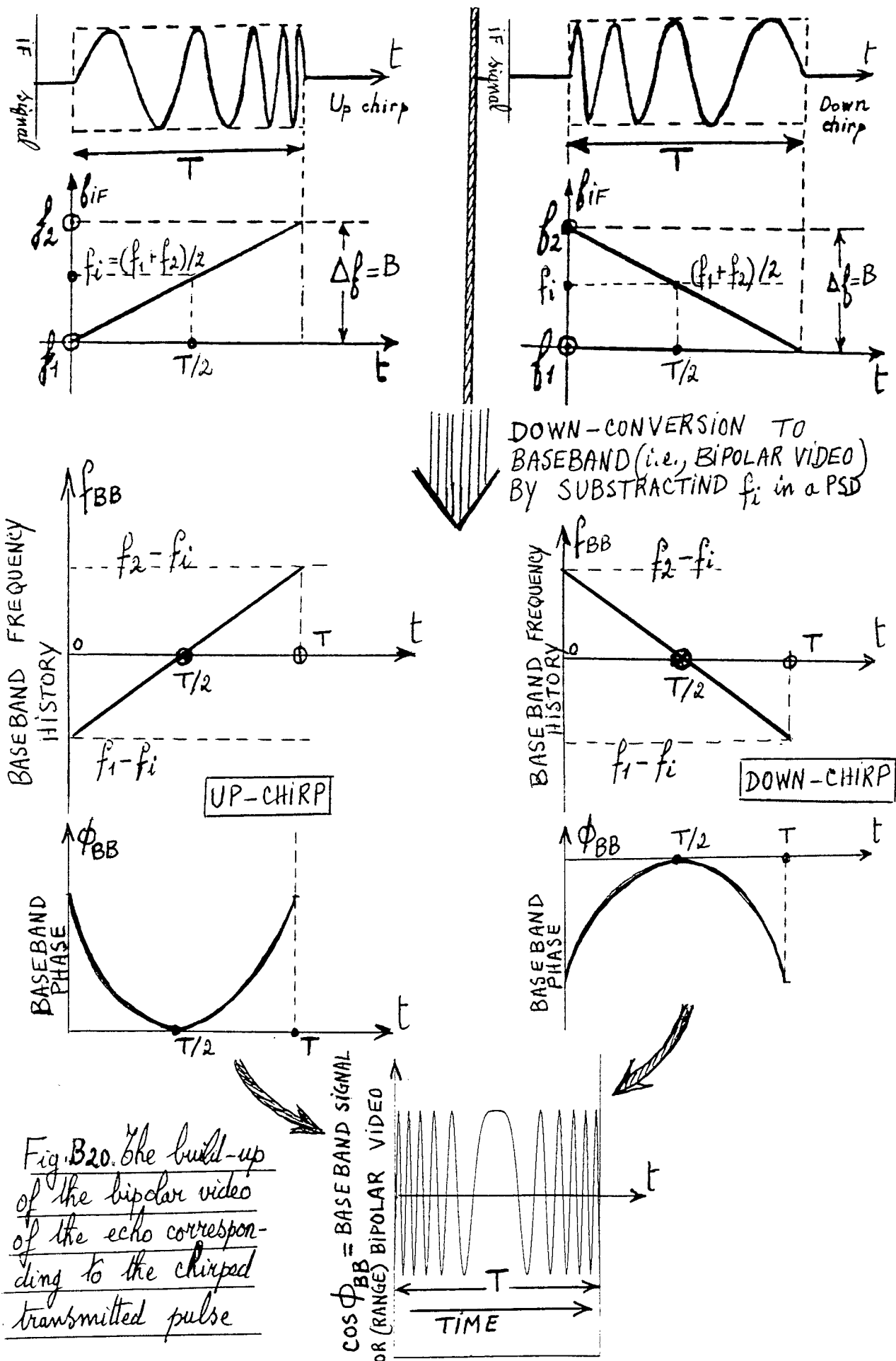


Fig. B20. The build-up of the bipolar video of the echo corresponding to the chirped transmitted pulse

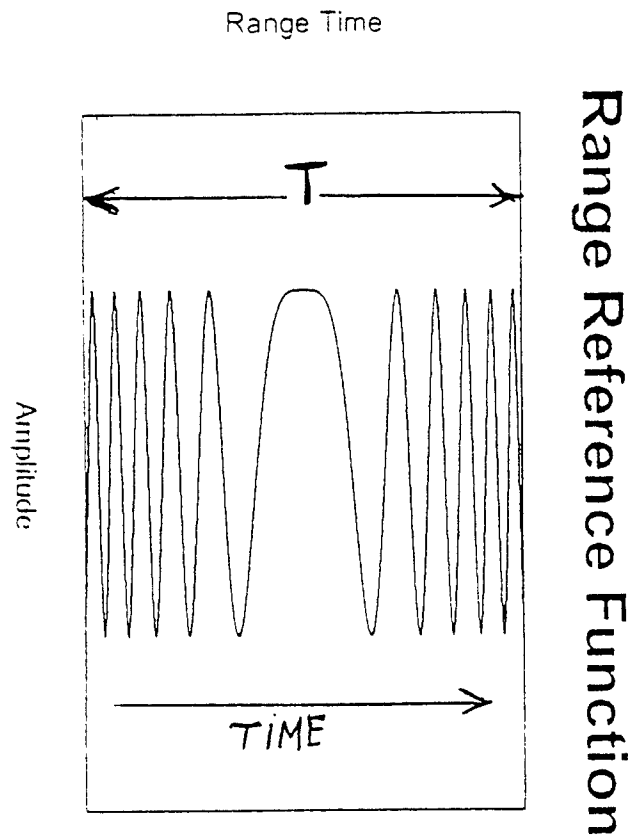


Fig. B21. The range reference function or range reference signal whose complex conjugate of the Fourier transform is the transfer function of the matched filter needed for pulse (i.e., range) compression

B. 4. The Doppler History and the Azimuth Reference Functions for Azimuth Compression

4. a. The straightforward but difficult way for the derivation of the Doppler history

We consider the case of a point target. Fig. B22 displays again the coordinate system with the coordinate along-track x corresponding to the azimuth dimension of the radar images. The origin of x as well as the origin of the time t are chosen when $R = R_0$, i.e., when the point target is purely broadside, or, equivalently, at the closest approach range R_0 . In other words, the origins of x and t correspond to the point of closest approach ($R = R_0$), so we can write

$$R^2 = R_0^2 + x^2 \quad \text{and } x = vt$$

hence

$$R = (R_0^2 + x^2)^{1/2} = R_0 \left[1 + \frac{x^2}{R_0^2} \right]^{1/2}$$

This can be expanded

$$R = R_0 \left[1 + \frac{1}{2} \frac{x^2}{R_0^2} - \frac{1}{8} \frac{x^4}{R_0^4} + \dots \right]$$

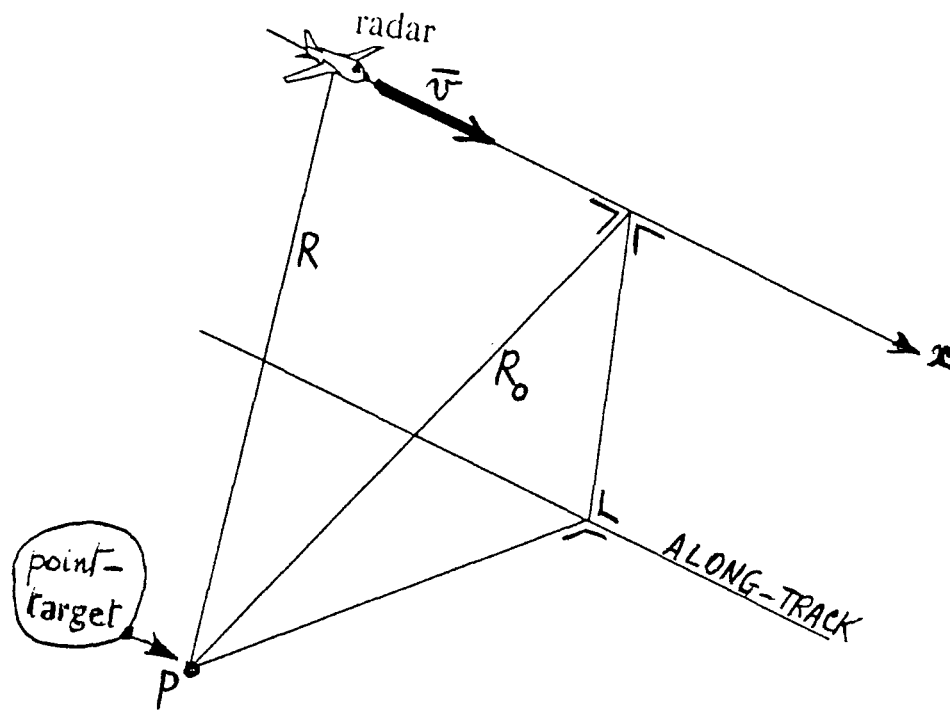


Fig. B22. SAR co-ordinate system for a point-target

If we limit this series expansion to the first two terms :

$$R = R_0 + \frac{x^2}{2R_0}$$

The two-way (because of the round trip of the radar wave) phase of the sequence of echoes (we remind that those echoes are received at a rate given by the PRF of the radar) is then given ($2R$ because of the round trip) by :

$$\phi(x) = -2R \cdot \frac{2\pi}{\lambda}$$

the negative sign occuring because an increasing path length represents a phase lag. So

$$\phi(x) = - \left[\frac{2R_0}{\lambda} \cdot 2\pi + \frac{2x^2}{2R_0\lambda} \cdot 2\pi \right]$$

or
$$\phi(x) = \phi_0 - \frac{2\pi x^2}{R_0\lambda}$$

with ϕ_0 being $\left[-\frac{2R_0}{\lambda} \cdot 2\pi \right]$, i.e., the phase at closest approach.

We see that the phase is a quadratic (or parabolic) function of x displayed by fig. B23 which is formerly identical with the base-band phase ϕ_{BB} of a down-chirp on fig. B20, provided we take into account that

$$x = vt$$

where v is the magnitude of the platform velocity. The previous calculation is obviously valid for a point-target only.

Fig. B23 displays also the extremes of the along-track footprint of the beam of the real radar antenna ; since (see fig. B8) this along-track footprint is $R_0\lambda/D$, those extremes are $\pm \frac{1}{2} \frac{R_0\lambda}{D}$.

Since $x = vt$, we can also consider the phase variation of fig. B23 as a function of time, to give the Doppler frequency f_D (due to the relative motion of the target with respect to the radar).

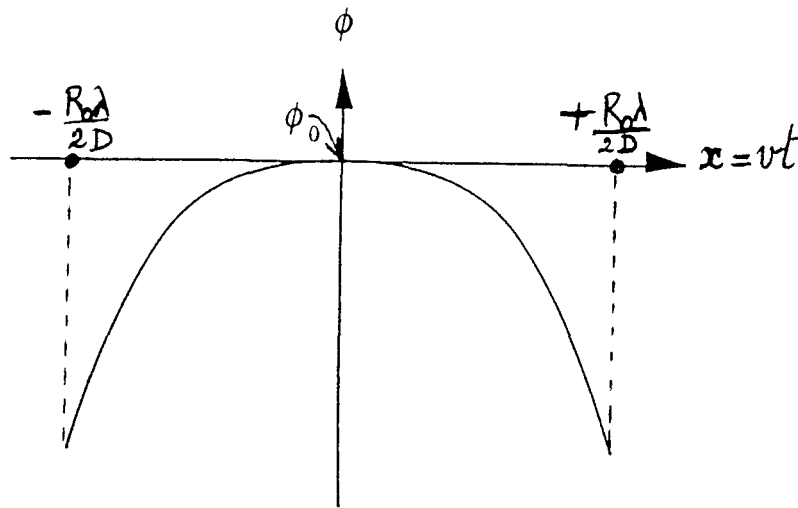


Fig. B23. Echo phase as a function of $x=vt$

So

$$\phi(t) = \phi_0 - \frac{2\pi v^2 t^2}{R_0 \lambda}$$

Thus the Doppler frequency history is expressed by

$$f_D = \frac{1}{2\pi} \frac{d\phi}{dt} = -\frac{2v^2 t}{R_0 \lambda} = \frac{-2v}{R_0 \lambda} x$$

So this quadratic variation of phase represents a linearly-varying Doppler frequency displayed by fig. B24 which is formerly identical to the baseband frequency history of a down-chirp on fig. B20.

Again, fig. B24 displays the extremes $\pm \frac{R_0 \lambda}{2D}$ of the footprint of the beam of the real antenna on ground, where D is the along-track dimension of the real antenna (the real aperture) ; the corresponding limits (as we announced when the constraint $PRF \geq 2v/D$ was computed) of the Doppler frequency are $\pm v/D$.

The expression

$$f_D = -\frac{2v^2 t}{R_0 \lambda}$$

and the corresponding fig. B24 are those of a down-chirp baseband frequency, thus a negative frequency ramp whose slope $\left[-\frac{2v^2}{R_0 \lambda} \right]$ is clearly range (the range R_0 at closest approach) dependent. The corresponding baseband signal (i.e., the corresponding bipolar video)

$\cos \phi = \cos \left[\int_0^t f_D(t) dt \right]$ is then represented by fig. B25 for various values of the range R_0 at closest approach. Obviously the azimuth reference functions of fig. B25 depend not only on the range R_0 but also on the platform velocity v .

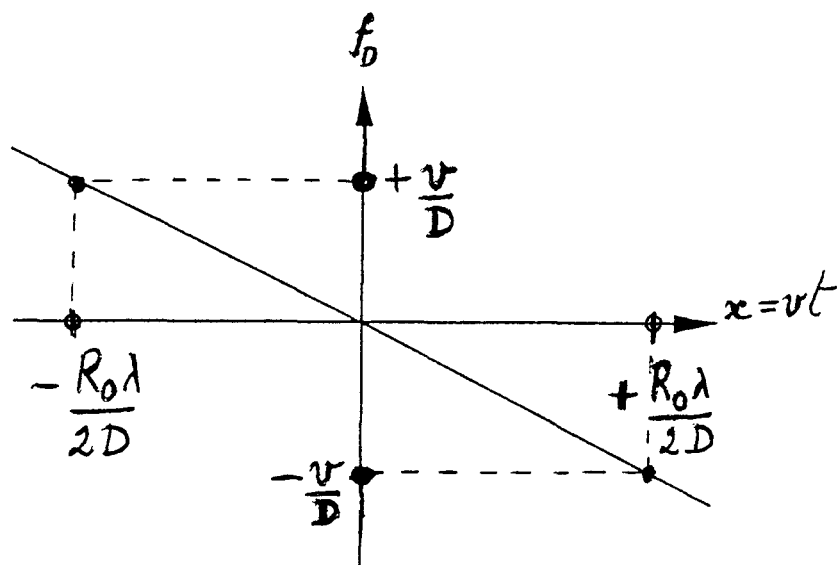


Fig. B24. Doppler frequency as a function of $x = vt$

Azimuth Reference Functions

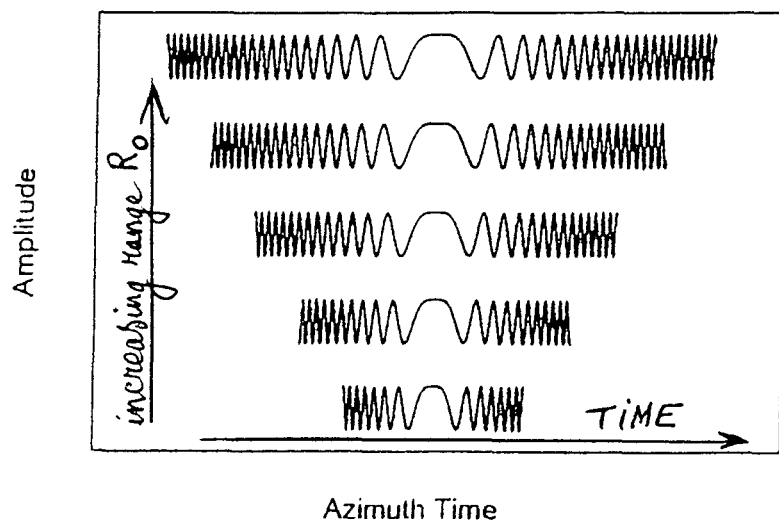


Fig. B25. The azimuth reference functions (replica of the expected returns from a point target at the same ranges)

The curves of fig. B25 represent the azimuth reference functions for the correlation operation needed to achieve the baseband azimuth compression which, after the range compression, will constitute the second part of the SAR signal processing. The curves of fig. B25 are similar to the range reference function of fig. B21.

The azimuth reference functions of fig. B25 are not the true azimuth bipolar video's (for different ranges R_0) because the azimuth bipolar video is not obtained instantaneously ; on the contrary this bipolar video is acquired after numerous pulses have been transmitted, the number of those transmitted pulses being (if T_R represents the IPP) : $\frac{T_S}{T_R} = \frac{L}{v} \cdot \text{PRF} = \frac{\lambda R_0}{D} \cdot \frac{\text{PRF}}{v}$

In order to retrieve the azimuth bipolar video, the successive echoes (from the successive transmitted pulses) have to be stored in amplitude and phase.

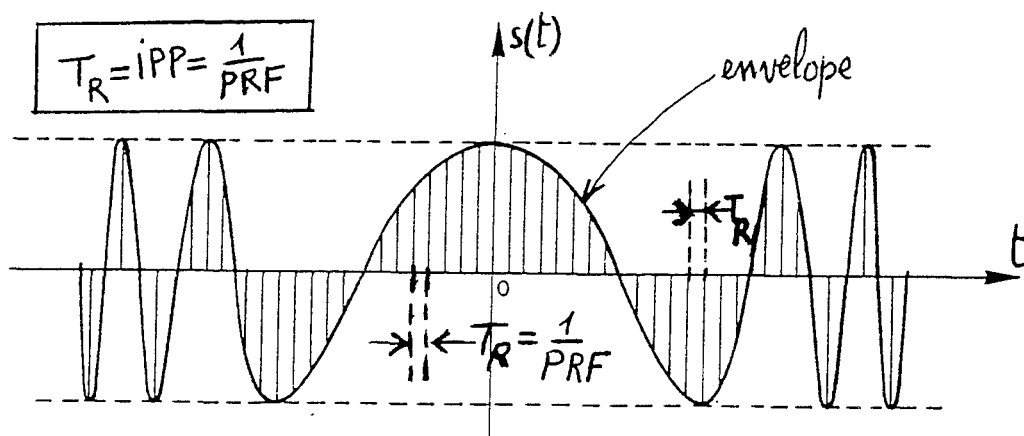


Fig. B26. True bipolar video of a point-target along track at a given R_0

As shown by fig. B26, the true azimuth bipolar video (i.e., the bipolar video along-track) of a point-target is composed of equidistant pulses separated by a time interval $T_R = \frac{1}{\text{PRF}}$; those pulses have an envelope given by the azimuth reference function of fig. B25 corresponding to the considered range R_0 .

The vertical stripes of fig. B26 are the echo pulses emanating from the same point target and originated from the successive transmitted pulses (i.e., the successive pulses transmitted by the radar).

In essence, the SAR azimuth bipolar video of fig. B26 is similar to the bipolar video's I and Q of fig. B27 except that the Doppler frequency (which is the frequency of the envelope) is linearly varying with time in the case of the SAR.

Notice that the azimuth bipolar video represented by fig. B26 does not encompass the influence of the true antenna radiation pattern. The upper curve of fig. B28 represents the envelope of the true azimuth bipolar video taking into account the radiation pattern of the true antenna. The

lower curves of fig. B28 express the fantastic effect of the azimuth compression.

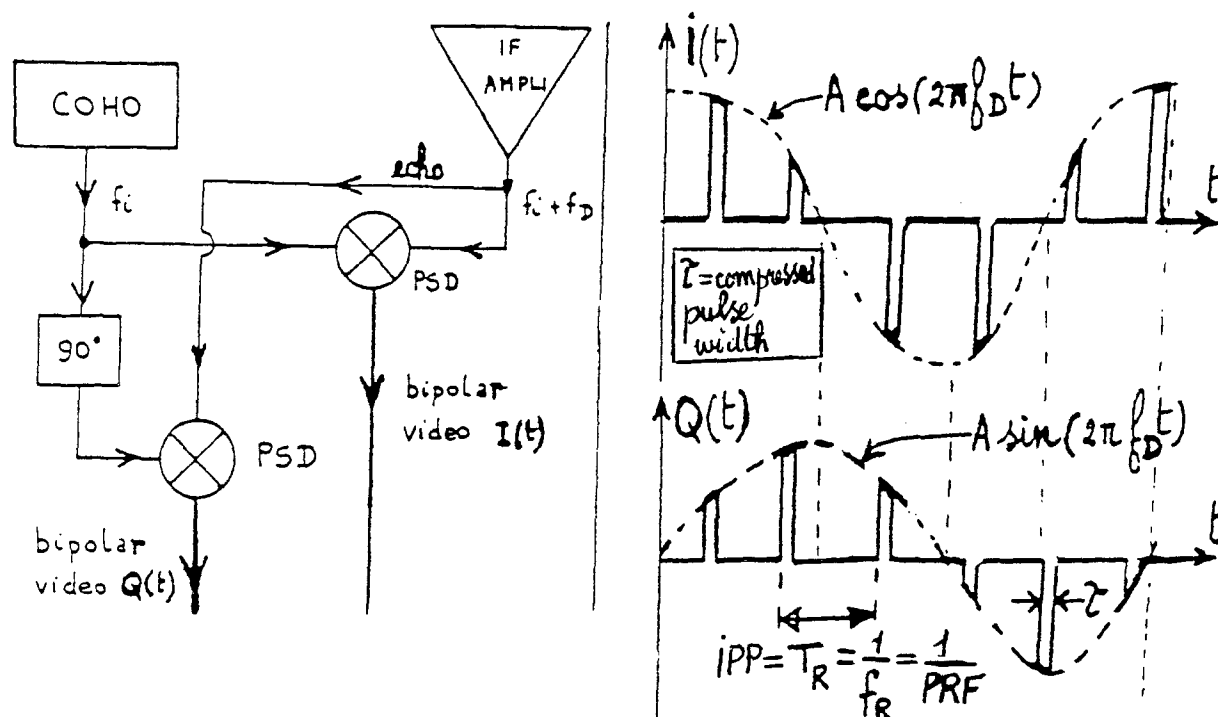


Fig. B27. The bipolar video's I and Q from a mobile target flying at constant velocity, in the base-band of a coherent radar (either a MTI radar or a pulse Doppler radar)

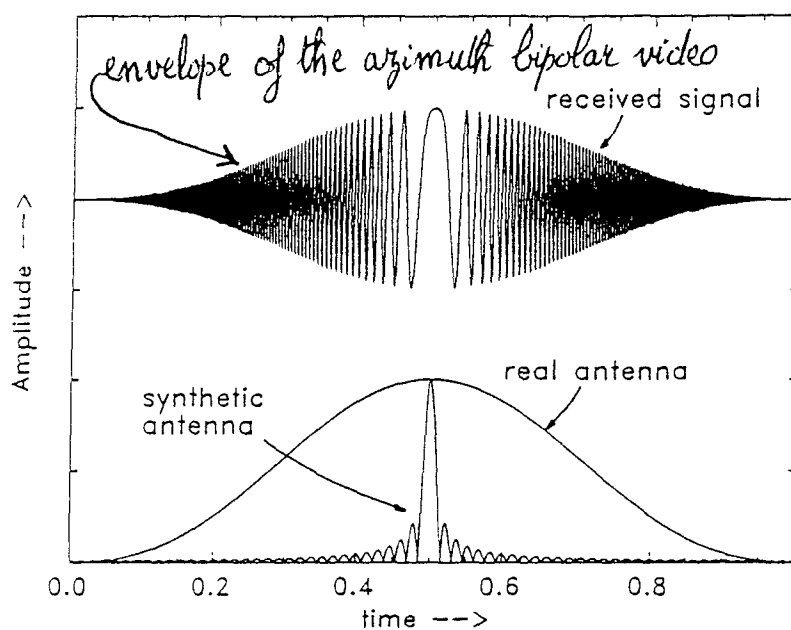


FIG. B28. Amplitude of a pulse response (upper curve) from a point target and SAR-processed impuls response of a point target (lower curve) following in comparison with point target response for a real aperture radar.

4. b. The easy way to establish the Doppler history along track for a point target

Since the SAR platform is flying along the point target P, conversely it can be considered that the SAR is stationary and that the point target P is moving in a straight path with a velocity $\bar{v}$ in the opposite direction : see fig. B29 where $PP_0 = x = vt$, P_0 being the target position at closest approach. Since fig. B29 represents the case of an opening target, the radial velocity v_R is negative and is thus given by $v_R = -v \sin \epsilon$. But ϵ is a small angle and thus $\sin \epsilon \approx \tan \epsilon = \frac{PP_0}{R_0}$

Hence $v_R = -v \cdot \frac{PP_0}{R_0} = -v \cdot \frac{x}{R_0} = -v^2 t / R_0$

and the corresponding Doppler frequency

$$f_D = \frac{2v_R}{\lambda} = -\frac{2v^2 t}{\lambda R_0}$$

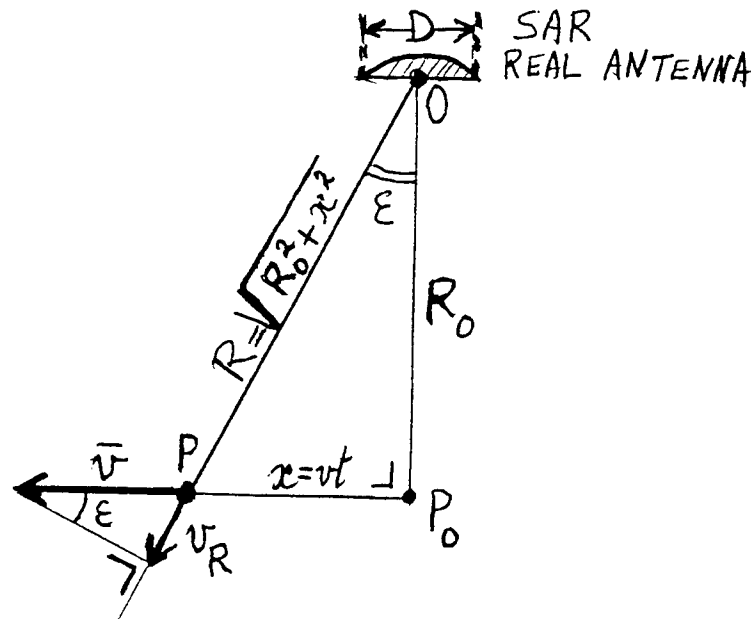


Fig. B29. Relative motion of the point target ; case of an opening target.

which is exactly the same as the Doppler shift f_D we have previously found through a hard and tedious way. From this point, it suffices to restart the whole reasoning following fig. B24 and the corresponding formula of the Doppler frequency f_D , knowing that the corresponding phase is given by

$$\phi = \int_0^t f_D(t) \cdot dt$$

and that $\cos \phi$ is the envelope of the azimuth bipolar video.

Since we have used a point target to establish the azimuth reference functions (given by $\cos \phi$), those azimuth reference functions are responses so that their complex conjugated Fourier

Transforms are the transfer functions of the matched filters which could be used alternatively to perform azimuth compression.

Fig. B30 displays the baseband echoes (as a function of time) corresponding to successive transmitted pulses. Out of those time varying echoes, how can we retrieve, for a particular point-target at a given range R_0 , the successive (alternatively positive and negative) pulses forming the azimuth bipolar video of fig. B26, the envelope of this bipolar video being the azimuth reference function of fig. B25 corresponding to the given range R_0 ?

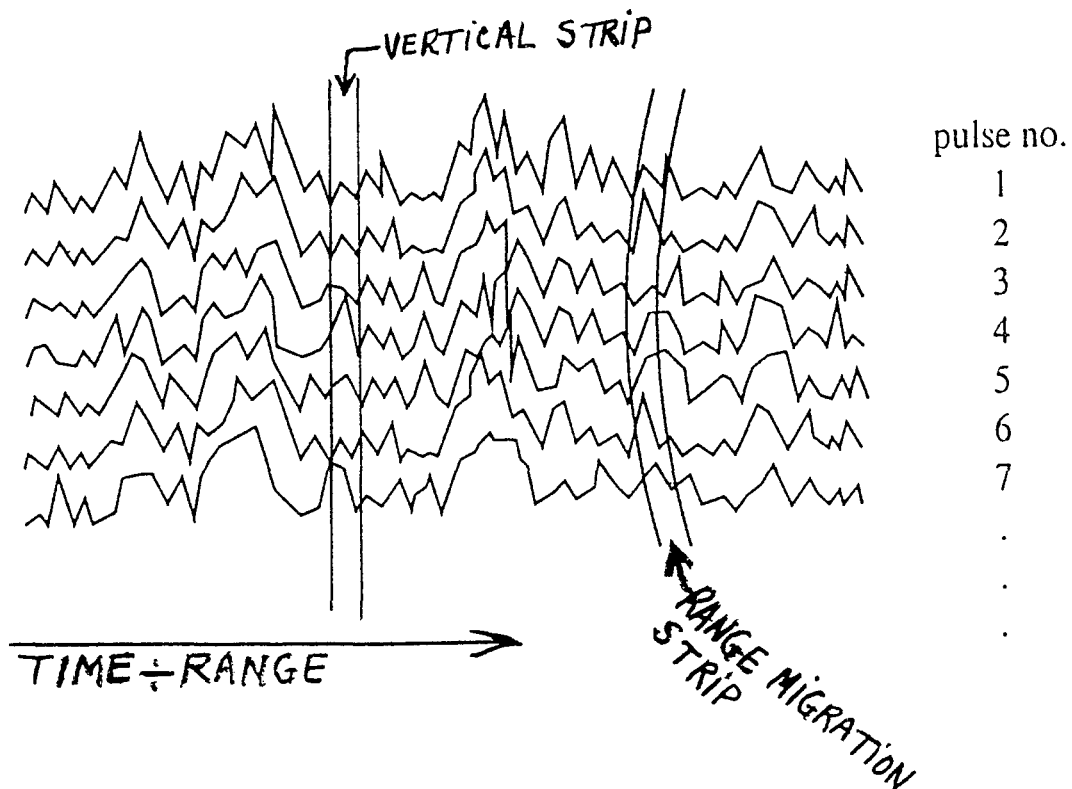


Fig. B30. Basis of SAR processing

Since the elapsed time is proportional to range in a pulse radar, those pulses forming the azimuth bipolar video should be found in the vertical strip of fig. B30. But because of the range migration (meaning that the range of each point target varies with time $t = \frac{x}{v}$) described by (see fig. B29)

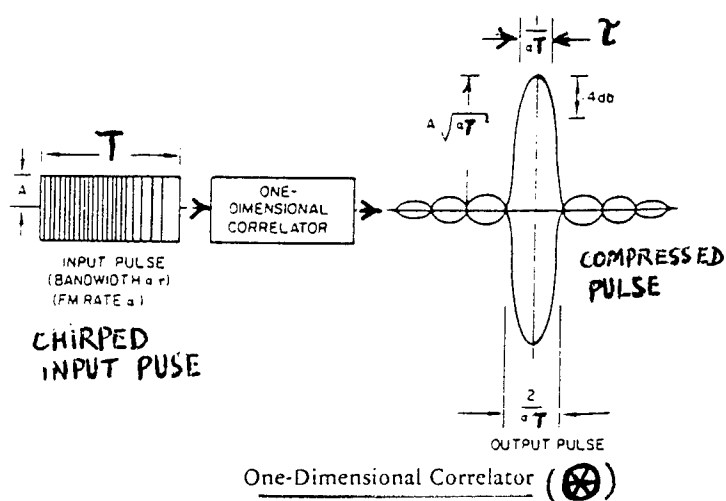
$$R = [R_0^2 + x^2]^{\frac{1}{2}}$$

those echo pulses belonging to the same point-target have to be retrieved from the range migration strip of fig. B30.

Each imaged cell is known as pixel. Each pixel corresponds to a particular target at a given range R_0 . To obtain the final video signal (going to the screen displaying the final high resolution image) corresponding to one pixel, the echoes at constant range R_0 from each curve of fig. B30 are taken, each echo is phase weighted (this phase weighting is performed by the correlations described in next paragraph) and the weighted echoes are summed. This process is repeated at all ranges and all azimuths and constitutes the sophisticated SAR signal processing.

B. 5. The Final Digital SAR Signal Processing.

Both range compression and azimuth compression are in fact pulse compressions. As known from the theory of the explanation of a PC (Pulse Compression) radar, each compression can be realized by the corresponding matched filter or, equivalently, by a correlation with the appropriated reference functions (those reference functions being replica of the expected returns from a point target at the same range). Fig. B31 represents the full digital SAR processing.



SAR Processing Schema

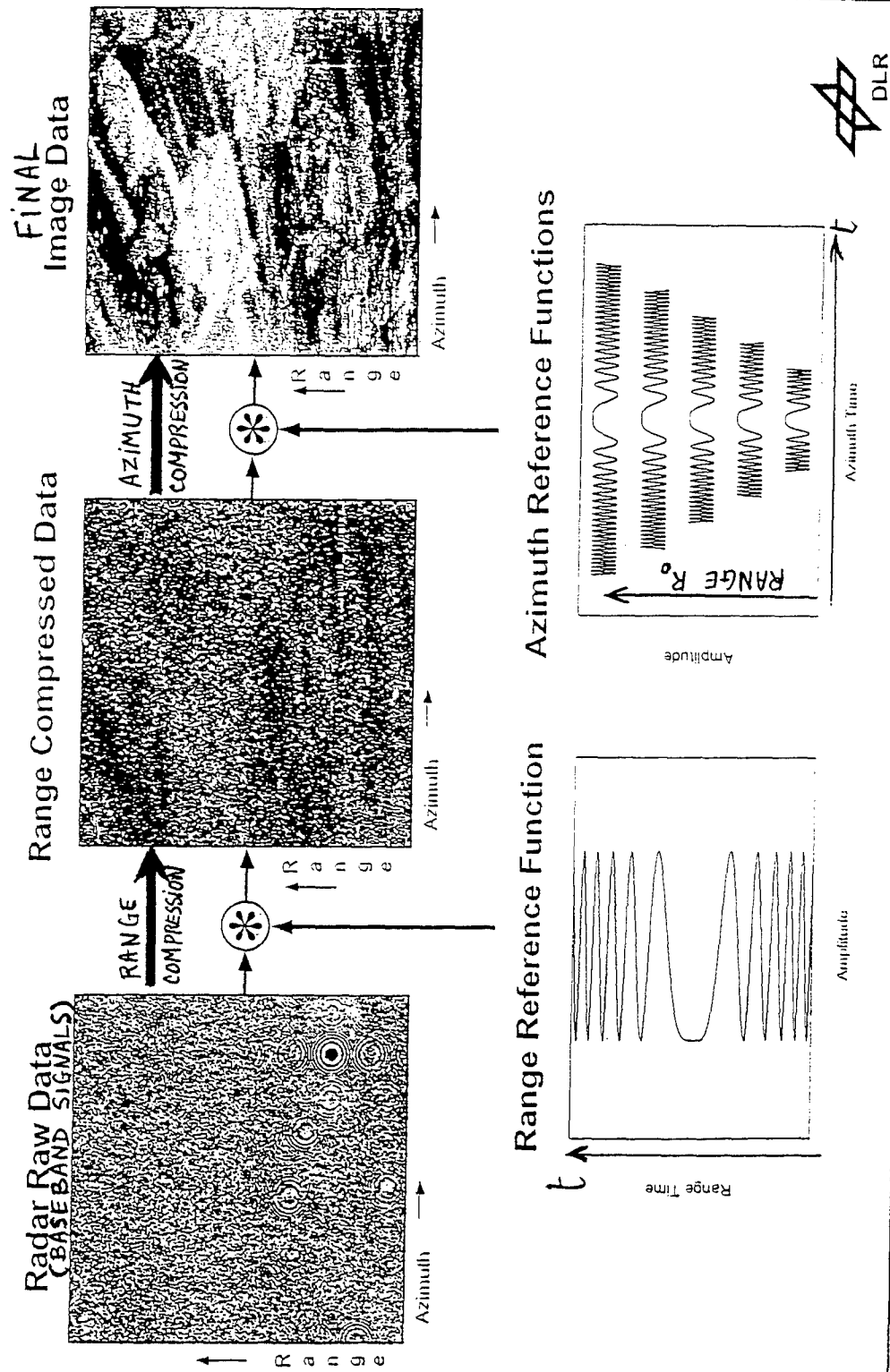


Fig. B31. SAR Processing Schema

As suggested by fig. B31 (this figure shows only baseband signals)

- the range compression is obtained by a correlation performed on each vertical strip of the radar raw data (so, after each transmitted pulse)
- the azimuth compression is obtained by a correlation performed on each horizontal strip (in fact, more accurately, on each range migration strip) of the range compressed data, i.e., after numerous transmitted pulses.

Since the azimuth reference functions are range dependent (they are depending on the range R_0 at closest approach), the range compression must be performed before the azimuth compression may be started. Those two subsequent operations are nicely summarized by fig. B32. The 1-D line of fig. B32 corresponds chiefly to the range migration which expresses that the range of a point target varies with $x = vt$.

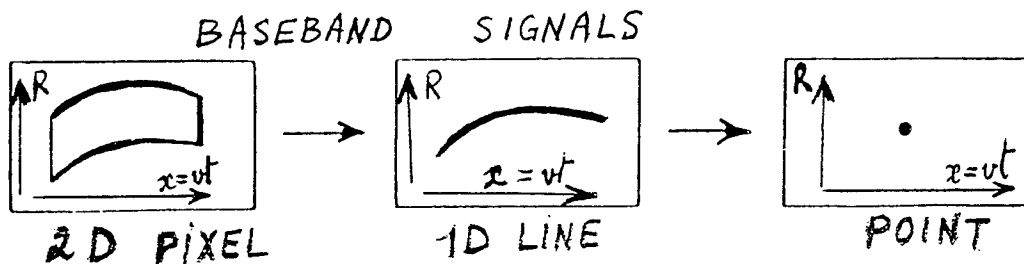


Fig. B32. Range processing followed by azimuth processing achieving an intrinsic 2-D (range and azimuth) data processing (the 1-D line corresponds to the range migration)

The intrinsic 2-D processing performed in 2 subsequent steps on fig. B32 can sometimes be done in one single step as illustrated by fig. B33 (depending on the algorithm approach)

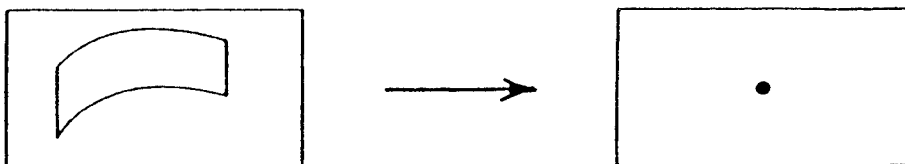


Fig. B33. 2-D processing : simultaneous range and azimuth data processing.

Notice : on fig. B31, the azimuth reference function has to be changed when the range R_0 (at closest approach) changes approximately by $2D^2/\lambda$ where D is the real aperture length along track. Notice also that the azimuth compression depends also on the SAR platform velocity because the magnitude v of this velocity clearly appears in the expression of the azimuth reference functions.

B. 6. SAR Imaging of Moving Targets

Up till now we implicitly considered only stationary targets (i.e., targets stationary with respect of the ground scene to be imaged). Considering the azimuth Doppler history (i.e., the sequence of echoes along track from the same stationary target) of fig. B24 , we remind that fig. B24 (as well as fig. B34) illustrates the relation (valid for a stationary target)

$$f_D = -\frac{2v}{R_0\lambda}x$$

which is a linearly varying Doppler history with a slope $\left(-\frac{2v}{R_0\lambda}\right)$. The last relation expressing this linear frequency ramp f_D as well as fig. B24 remind us that the exact position along track of the stationary target is given by $x=0$ (which corresponds to the range R_0 at closest approach , i.e., when the target is purely broadside), i.e., by the x value of the zero crossing point of the Doppler history of figures B24 and B34. Considering now a stationary target and a moving target at the same location (i.e., with the same $x=0$ coordinate along track). Their respective azimuth Doppler histories as a function of the x -coordinate along track are both given by fig. B34 showing that, with respect to the stationary target, the moving target (at the same location) exhibits a Doppler offset expressed by

$$\text{Doppler offset} = \frac{2v_R}{\lambda} = \frac{2v_R f_c}{c}$$

where v_r is the radial component of the velocity of the moving target with respect to ground and f_c is the radar carrier frequency.

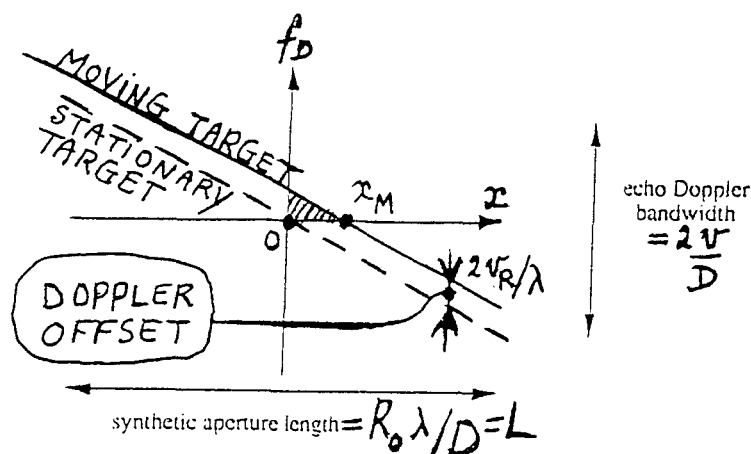


Fig. B34. A target with a radial velocity is matched filtered by the SAR processing with an azimuth shift (or azimuth displacement) x_M

Since the x-coordinate along track of the target is given by the zero crossing point of the Doppler history, it follows (see fig.B34) that the SAR processing will consider that the x-coordinate along track of the moving target is given by the azimuth shift (or azimuth offset) x_M instead of the true value $x=0$. This azimuth shift (or azimuth displacement) x_M satisfies (considering the hatched triangle) :

$$\frac{x_M}{\text{Doppler offset}} = \frac{L}{2v/D} = \frac{R_0\lambda/D}{2v/D} = \frac{R_0\lambda}{2v} \quad (= \text{inverse of slope})$$

$$\text{Or } \frac{x_M}{2v_R/\lambda} = \frac{R_0\lambda}{2v}$$

Hence

$$x_M = v_R R_0 / v$$

where

- v_R is the radial component of the velocity of the moving target with respect to the ground
- v is the SAR platform velocity with respect to the ground
- R_0 is the range at closest approach

Consequently, the SAR processing will result in a displacement x_M along track of the moving target with respect to its true location. This is nicely illustrated by figures B35 and B36.

Fig. B35 shows an image from an airborne SAR of a moving train. Because the train has a component of velocity in the range direction (i.e., in the radial direction) , the train image is shifted (i.e., displaced) with respect to the stationary railway track; so the train appears to be travelling not along the railway track, but displaced to the side. From the formula

$$x_M = v_R R_0 / v$$

it follows that it is possible to estimate the target radial velocity v_R , if one knows the platform velocity v , the geometry R_0 and the azimuth shift (to know x_M , it is necessary to know the true position of the target).

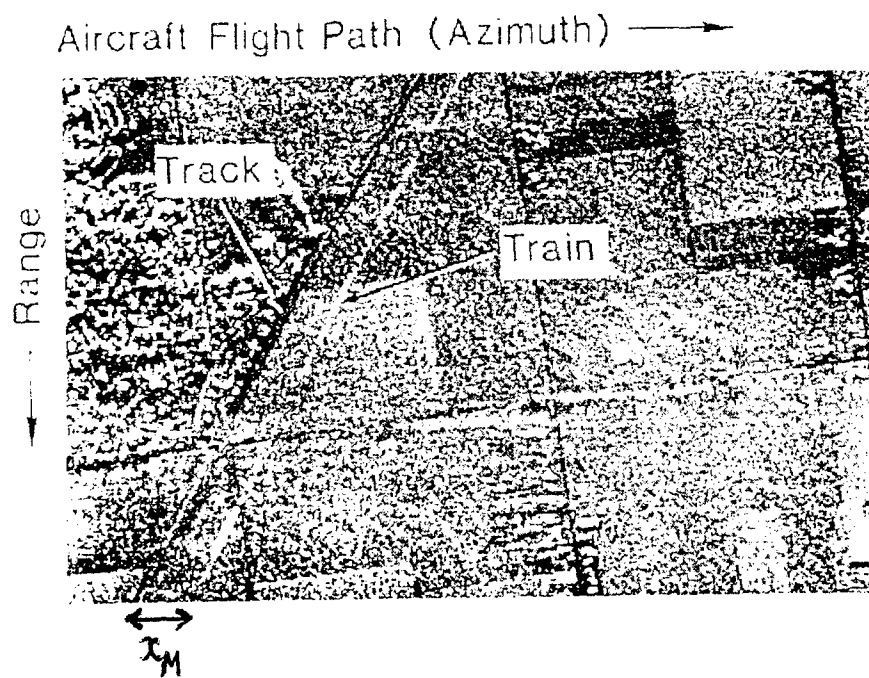


Fig. B35. SAR image of a moving train (after Rufenach et al.)

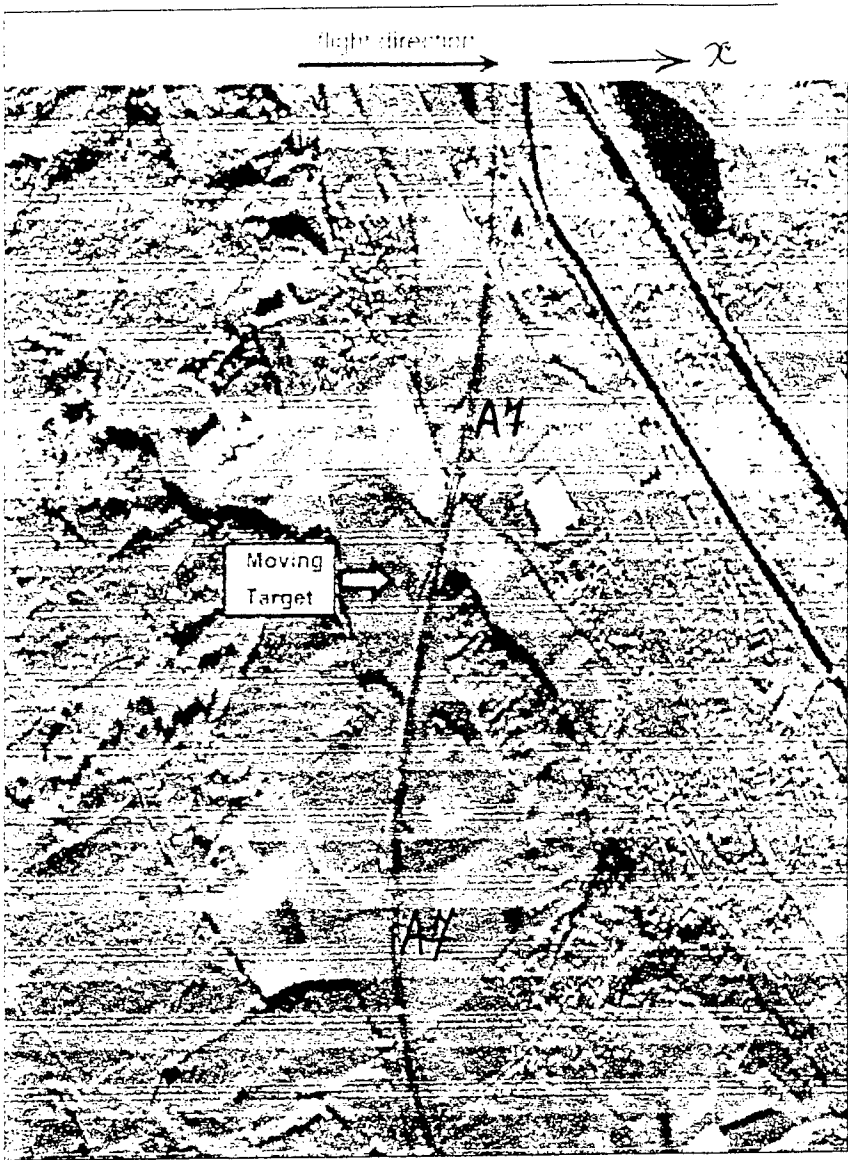


Fig. B36. C-band image of the high resolution real time processor. The imaged area shows the Iller river in the vicinity of Altenstadt, Bavaria, Germany.

Fig. B36 shows a real-time image from DLR airborne SAR during a test flight in the vicinity of Oberpfaffenhofen (location of DLR) The main SAR parameters were: aircraft altitude ca.3000m, ground speed ca.80 m/s and C-band. The image dimensions are approximately 3000m×3000m. On fig. B36, the indicated moving target is a car displaced off the highway A7. This displacement occurs due to the Doppler frequency offset $2v_R/\lambda$ introduced by the motion of the car with respect to the highway A7. It is worth to give here **the relevant parameters of the SAR's carried by the European Space Agency's ERS-1 and ERS-2:**

carrier frequency: 5.3 GHz

pulse bandwidth ($f_2-f_1 \approx 1/\tau$): 15.5MHz

pulse width T: 37.1us

peak transmitter power: 4.8kw

swath width: 80km

PRF: 1.64kHz < PRF < 1.72kHz

pixel size: 30m×30m

B. 7. MTI (Moving Target Indication)

In many situations, it is useful to be able to distinguish moving targets from stationary ones. In previous paragraph, we already saw that the image of a target with a component of velocity in the range (i.e., the radial) direction will suffer an azimuth shift x_M (i.e., an off displacement along track), but unless the true position is known, this does not allow the detection of moving targets, i.e. moving target indications (MTI).

There are essentially two ways to perform MTI with a SAR :

- the 1st way is DPCA (Displaced Phase Center Array) processing.
- the 2^d way is STAP (Space-Time Adaptive Processing)

a. STAP

By using an along-track array, it is possible to perform both spatial-domain (directional) filtering and time-domain (Doppler) filtering. This is known as STAP (Space-Time Adaptive Processing) and is the subject of considerable current research.

b. DPCA (Displaced Phase Center Array) processing (fig. B37)

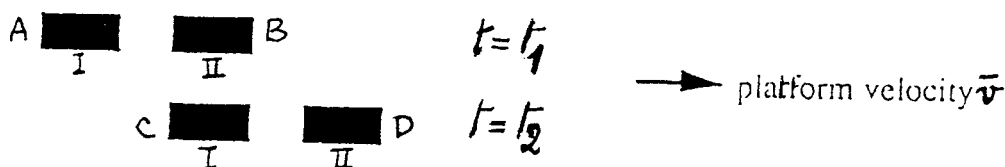


Fig. B37. DPCA processing

DPCA processing uses two (or more) antennas I and II arranged in tandem (fig. B37).

At time t_1 the 2 antennas occupy the positions A and B. At time t_2 the 2 antennas occupy the positions C and D. By controlling the radar PRF as a function of platform velocity v , it is possible that position B of antenna II at time t_1 coincides with position C of antenna I at time t_2 . So by transmitting alternately and controlling the radar PRF as a function of platform velocity v , it is possible to obtain echoes from identical positions, but displaced in time (fig. B37). If one echo is subtracted from the other (e. g. , subtract signal II at position B from signal I at position C), stationary targets should cancel, but echoes from moving targets will give a non-zero result from the subtraction, so should remain. Coe and White have carried out a theoretical and practical evaluation of this technique. They found that cancellation of the order of 25dB is possible. The major source of error are control of the PRF (to give exact spatial coincidence of the samples), and matching of the antenna radiation patterns.

B.8. Optical Processing.

The most general structure of the SAR depicted by fig. B38 remind us that the signal processing is always performed in the baseband, i.e. , on the bipolar video, and that the signal processing can be

- either a digital one, which we discussed in the previous paragraphs
- or an optical one (in that case the bipolar video is fed to a CRT) which we are going to describe now.

The VCO realizes both the AM (chopping the microwave signal from the STALO into pulses with the appropriate PRF) and the FM within each pulse.

Prior to the optical processing, the bipolar video needs to be stored on a moving film; but a film can only record light intensities and light intensities (out of CRT for instance) are always positive or zero! As shown by fig. B38 the echoes are downconverted (by means of the PSD) to baseband and the obtained bipolar video $s(t) = \cos\phi(t)$ at the output of the PSD is used to intensity-modulate a CRT display. Since the bipolar video $s(t) = \cos\phi(t)$ at the output of the PSD is essentially an unbiased one-i.e. , the average value is zero-as shown by fig. B39, the negative values of this unbiased bipolar video $s(t)$ would not be taken into account and would not be stored on the moving film.

Therefore, the baseband radar echoes are placed on a bias level s_b - giving thereby the biased bipolar video $[s(t) + s_b]$ depicted by fig. B40 - prior to recording on the moving film, in order always to have positive values at the input of the CRT.

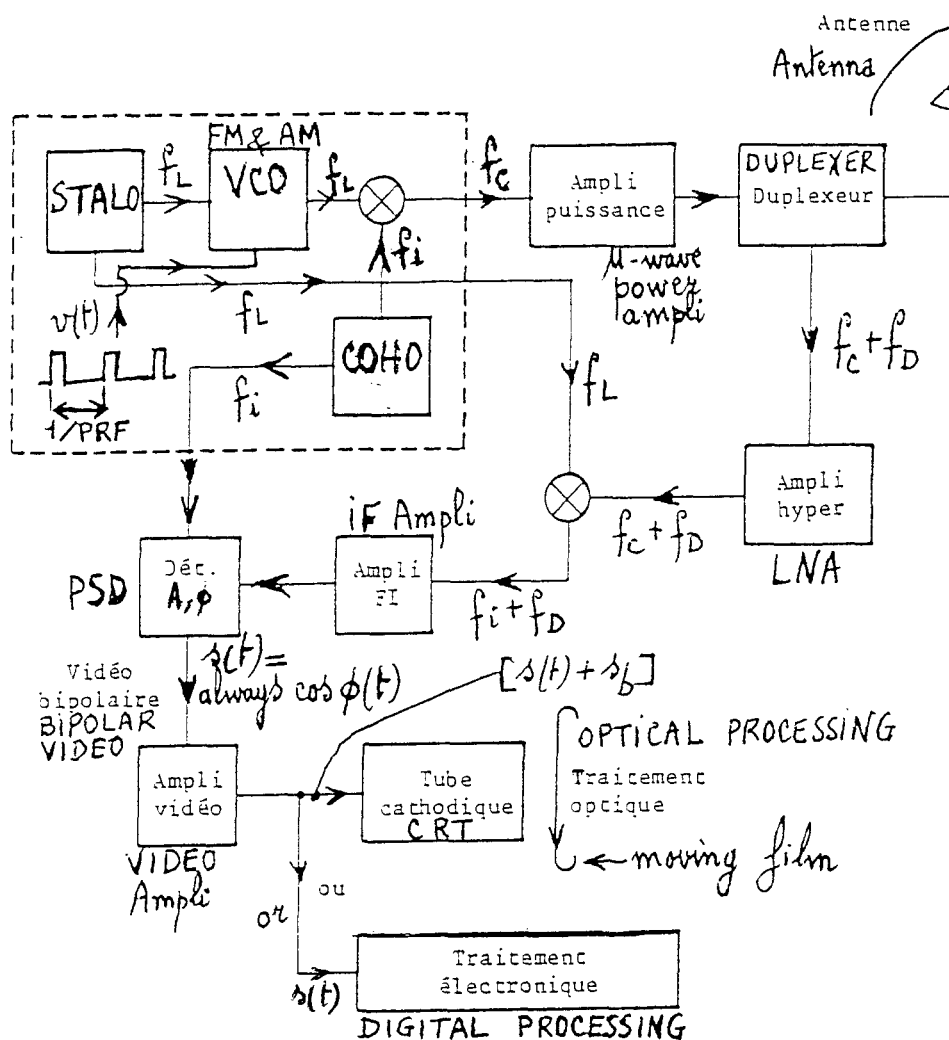


Fig. B38. General layout of the SAR hardware (p.295, THOUREL)

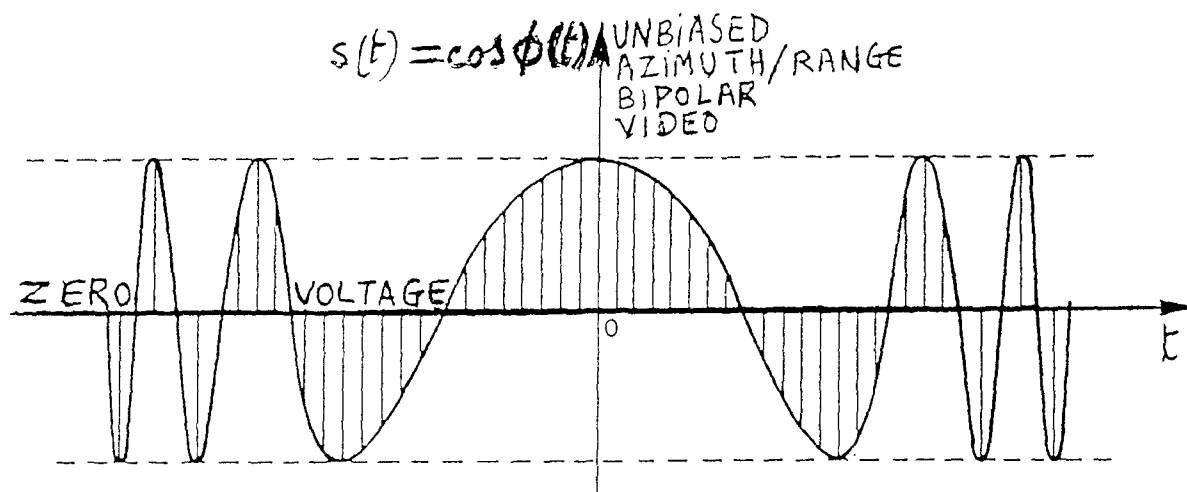


Fig. B39. Usual (unbiased) bipolar video at the output of the PSD of the radar layout of fig. B38.

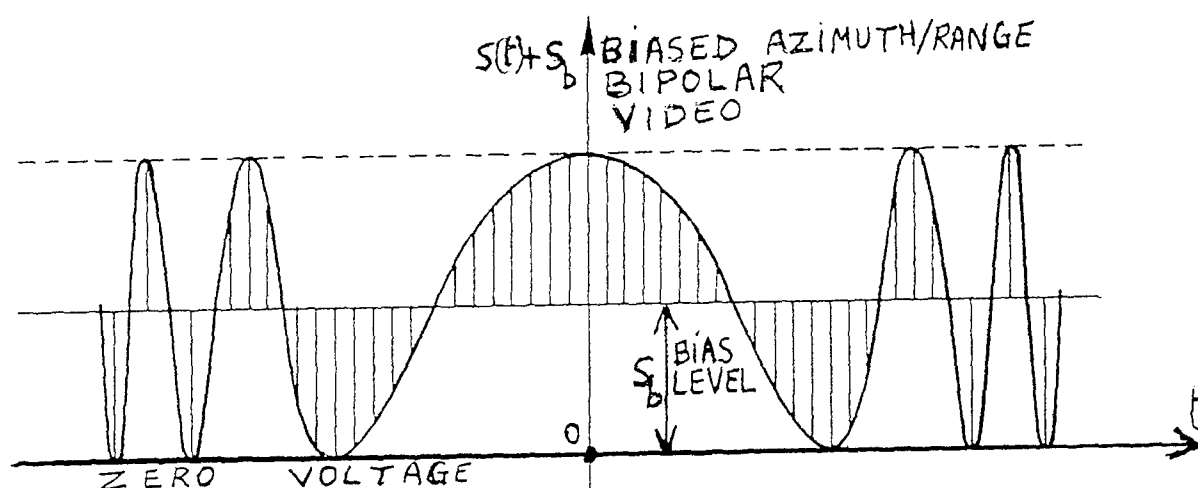


Fig. B40. Biased bipolar video to be applied to the CRT for recording on a film for later optical processing.

Therefore (see fig. B41) , the biased bipolar video $[s(t) + s_b]$ of fig. B40 is fed to the input of the CRT and intensity modulates the screen of the CRT in front of which the storage film is moving. So each echo is recorded as one line on a strip of the film, and successive echoes (emanating from successive transmitted pulses) are recorded as neighbouring lines.

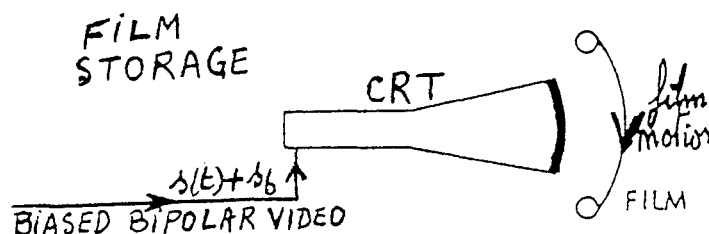


Fig. B41. Film storage through intensity-modulation of a CRT by the biased bipolar video
 $[s(t) + s_b]$

In other words, the radar receiver output is a sequence of reflected range pulses emanating from the transmitted pulses. The radar receiver output (i.e. , the output of the PSD on fig. B38) is placed on a bias level s_b , (cfr fig. B40) before being used to modulate the intensity of the sweep of the CRT (Cathode Ray Tube) , a single sweep being generated by each radar transmitted pulse, i.e. , the electron beam of the CRT (fig. B41)

is swept in synchronism with the returning pulses. The film moves (cfr fig. B41) across the screen of the CRT; the speed of the film is directly proportional to the radar platform velocity. Consequently, successive range traces are recorded side by side producing (cfr fig. B42) a two-dimensional format in which the dimension across the film represents range, and the dimension along the film corresponds to the along-track dimension, i.e., the azimuth coordinate x .

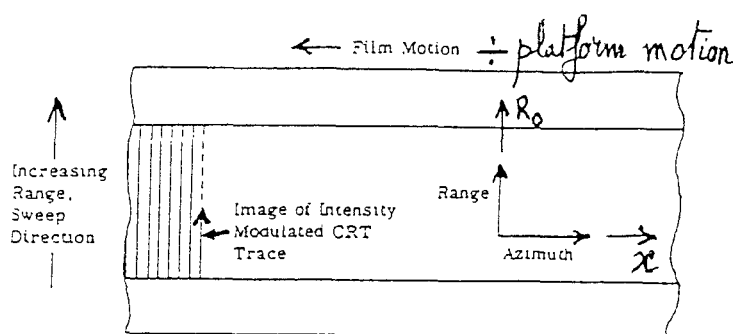


Fig. B42. Radar signal storage format

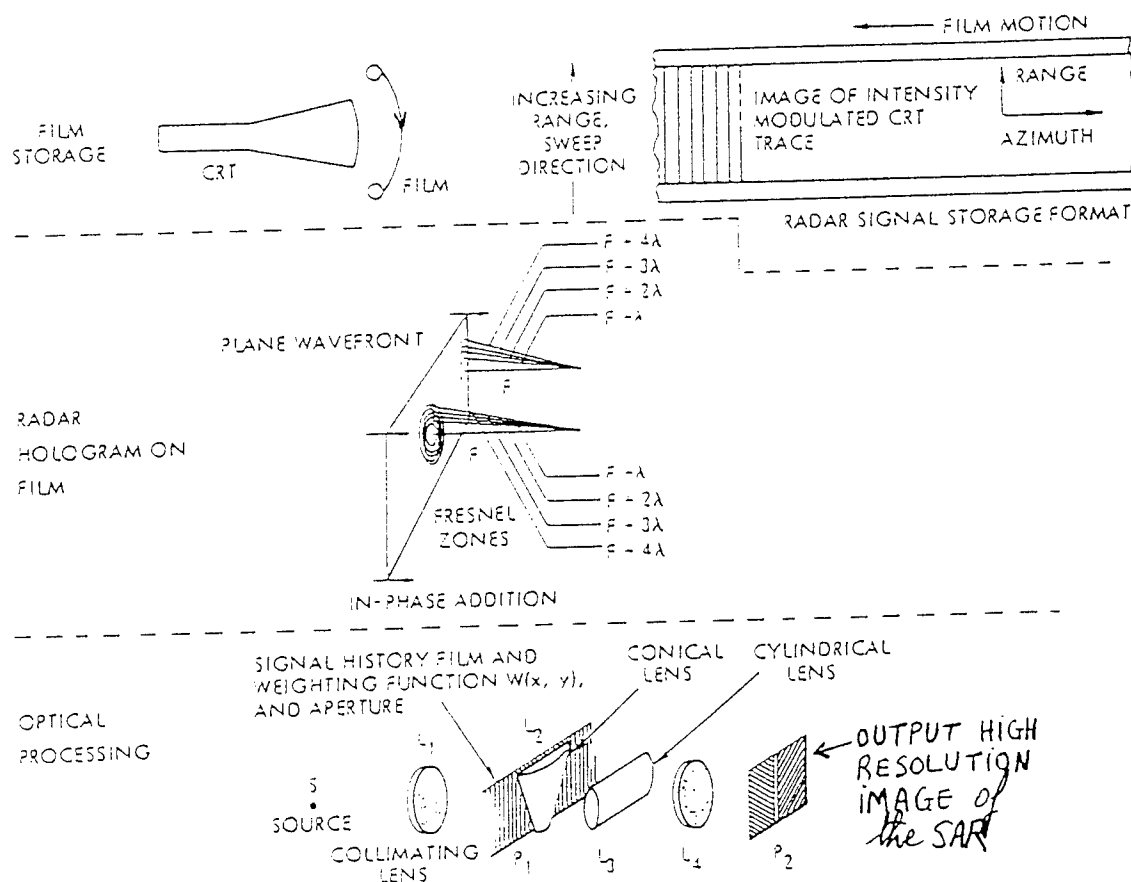


Fig. B43. Simultaneous Pulse Compression and Beam Sharpening (Optical Processing)

Fig. B43 both summarizes the essentials of optical type processing and introduces new terminology. A single sweep is generated across the CRT; the output of the radar's coherent detector PSD (see fig. B38) is used to modulate the intensity of this sweep.

The film moves across the tube; the speed of the film is directly proportional to the platform velocity. The lines are packed on the film so that there is some overlap. The phase history of the backscattered energy is stored on this film; the phase history within each echo pulse is stored in range dimension. When stored on the film, these phase histories have been characteristically called radar holograms owing to their similarity to optical holograms.

An image film from the signal film is produced in the optical processor, which typically is constructed with the elements shown in the lower diagram of fig. B43. A coherent source (thus, preferably, a laser source) of light (collimated to form a plane wave) shines on the signal film which contains the radar hologram; the conical lens L_2 has a (range) varying focal length to compensate for the phase variations in range. Because there is a different focal length for each

range element, a cylindrical lens L_3 is used to bring all of the range information into focus in one particular plane. Since the film recorded in front of the CRT of fig. B41 is processed on the ground, the optical processing is not a real time one. But, since 1995, several companies are giving a new life to optical SAR processing by replacing the film by spatial light modulators (also termed "light valves") enabling thereby a real time optical SAR processing.

Fig. B44 gives a better view on the whole optical SAR processor composed

- first of a conical lens L_2 performing the azimuth compression.
- secondly of a cylindrical lens L_3 performing the range compression

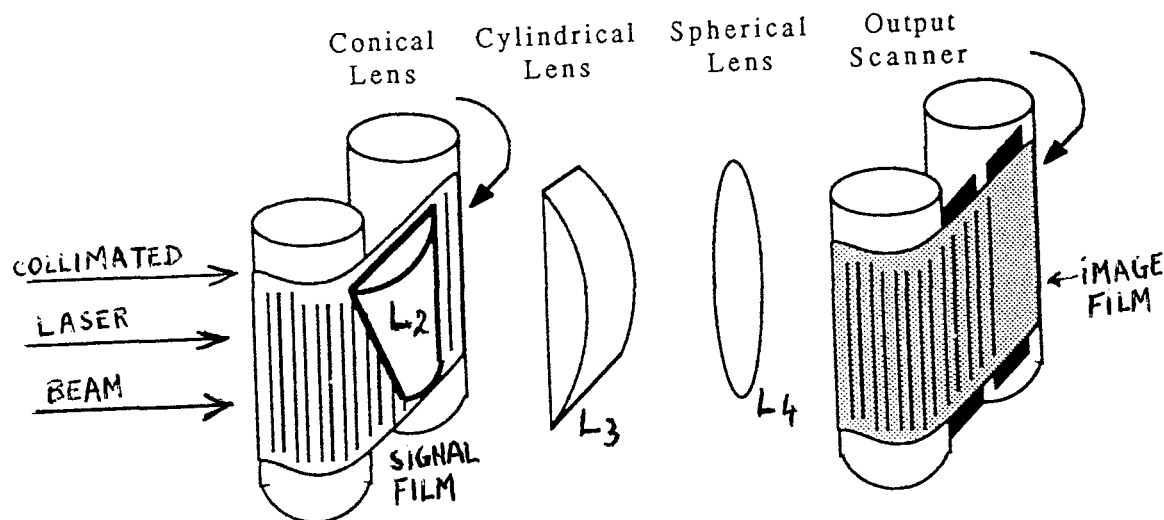


Fig. B44. Optical correlator for optical SAR processing (after BARBIER)

After the conical lens L_2 and the cylindrical L_3 , the rays are collimated, i.e. , the rays are parallel. Therefore a spherical lens L_4 is needed to focus all those rays into a very small pixel (ideally focusing into a single point, as shown by the upper image of fig. B17)

It can be easily understood why a conical lens is needed for azimuth compression by looking to figure B15 which shows that the recording of the phases of the successive echoes (emanating from the successive transmitted pulses) from the same point-target P is equivalent to a plane-concave diverging lens whose concave diopter is nothing else but the iso-range line of fig. B15. Clearly, fig. B15 shows that the focal length of this equivalent diverging lens is depending on R_0 , the target range at closest approach, which is also the radius of curvature of the concave spherical diopter of this diverging lens.

Since azimuth compression (i.e. , the realization of the synthetic aperture) requires the compensation of the phases $\left[-\frac{2\pi}{\lambda} \cdot 2R_i \right]$ recorded at the sequential positions E_i of the single SAR antenna, it is clear that this azimuth compression can be done by the opposite lens of the equivalent diverging lens, thus by a converging lens exhibiting a range depending focal length, which is nothing else but the conical lens L_2 of fig. B44.

A detailed working of the optical SAR processor of fig. B44 can be explained as follows. The three lenses L_2 , L_3 and L_4 form an anamorphic lens system. The conical lens L_2 , having a focal length equal and opposite to that of the signal films, operates on the azimuth focal plane and moves it to infinity. The azimuth focal plane thus becomes erected. The cylindrical lens operates only in the range dimension and images the signal film plane at infinity. The spherical lens, operating in both dimensions, takes the image at infinity and reimages it at its focal plane, where the image is now sharply focused in each dimension.

B. 9. Spotlight Mode and other High Resolution SAR's.

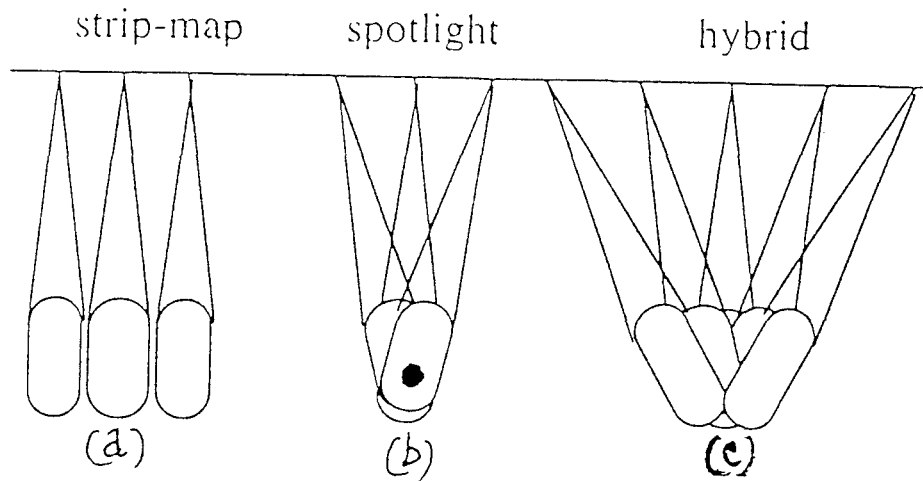


Fig. B45. Strip-map, spotlight and hybrid imaging modes (after Belcher and Baker)
where the big black dot of spotlight is the rotating point.

The conventional mode of SAR image formation as described in the previous paragraphs is known as STRIP-MAP mode. As we have seen, the strip-map mode uses (cfr fig. B45.a) fixed broadside antenna pointing, i.e. , in the strip-map mode the SAR antenna beam points perpendicularly to the platform path (fig. B46.a).

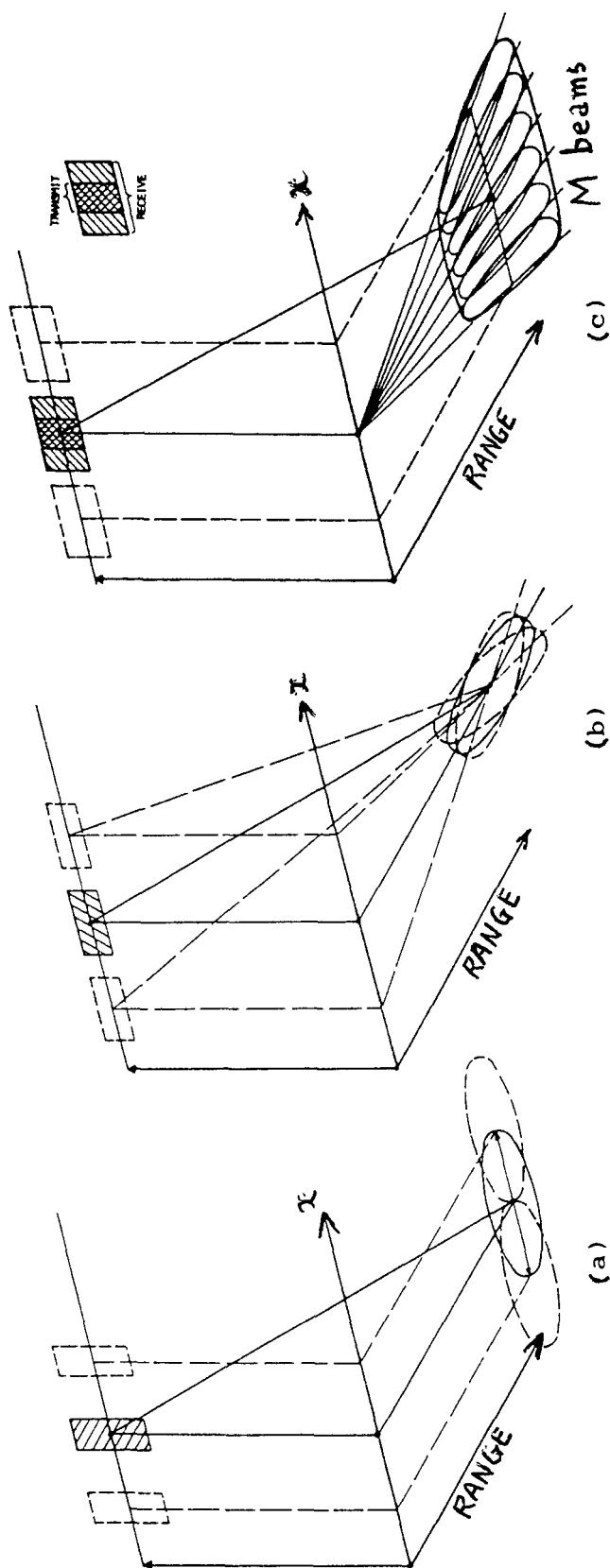


Fig. B46. High Resolution SAR concepts in azimuth

- (a) Regular SAR, i.e., strip-map mode ; (b) Spotlight SAR
(c) Multibeam SAR

If the SAR antenna is steerable in azimuth (either mechanically or electronically), it is possible (see figures B45.b and B46.b) to steer the antenna beam so that a particular target is illuminated for a longer period of time (i.e. , a longer dwell time), thereby allowing a longer aperture to be synthesized (since the synthetic aperture length is the path travelled by the sensor during the dwell time) and achieving better azimuth resolution δAZ than implied by the equation

$$\delta AZ = D/2$$

which is valid for a strip-map SAR.

It can be shown that the maximum azimuth resolution achievable in spotlight mode (if the antenna beam can be steered through 180°) is equal to

$$\delta AZ_{\text{spotlight}} = \lambda/4$$

In practice this is unrealistic (because the antenna beam can obviously not be steered through 180°) though considerably better resolution than strip-map mode can be achieved in spotlight mode.

It is possible (see fig.B45.c) to consider a hybrid mode of operation, intermediate between strip-map and spotlight modes.

Of course, spotlight mode demands that the target scene of interest is known beforehand, and only this scene is imaged.

As a summary, it may be claimed that a "spotlight" mode SAR directs the antenna to a single spot over and over again as the antenna passes high above. In the "spotlight" SAR mode, the synthetic aperture length L (which is the sensor path travelled during the dwell time) is drastically increased by pointing the steerable SAR antenna to the small area of interest during the fly-by.

Another method of achieving a high along track resolution (i.e., a high azimuth resolution) is to implement the well known multi-beam phased array radar as shown in figures B46(c) and B47. If M beams are formed simultaneously ($M = 3$ in the case of fig. B47), it is quite obvious from fig. B46(c) that the multi-beam SAR exhibits an M time finer along track resolution, i.e., for the multi-beam SAR (with M simultaneous beams)

$$\delta AZ = \frac{D/2}{M}$$

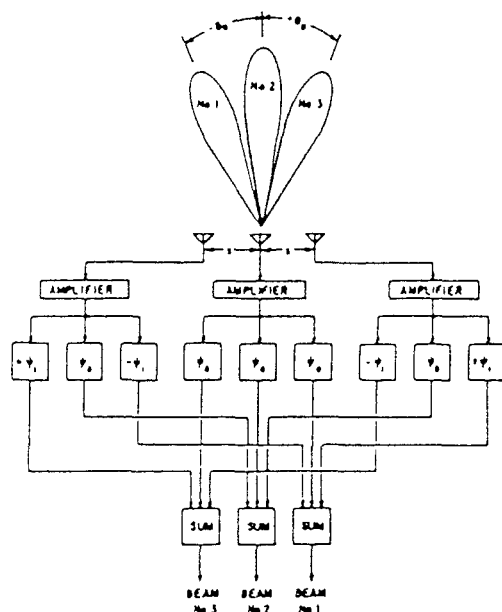


Fig. B47. Simultaneous postamplifier multibeam formation using array antenna.

B. 10. Swath Width.

As shown by fig. B48 the swath width AB is linked to the 3dB elevation beamwidth θ_{EL} of the SAR antenna and to the range R_0 at closest approach.

Since θ_{EL} is a small angle, $AC \approx R_0 \cdot \theta_{EL}$ and $AB \approx AC / \sin \Delta$ where Δ is the depression angle. Hence the swath width $AB \approx R_0 \cdot \theta_{EL} / \sin \Delta$

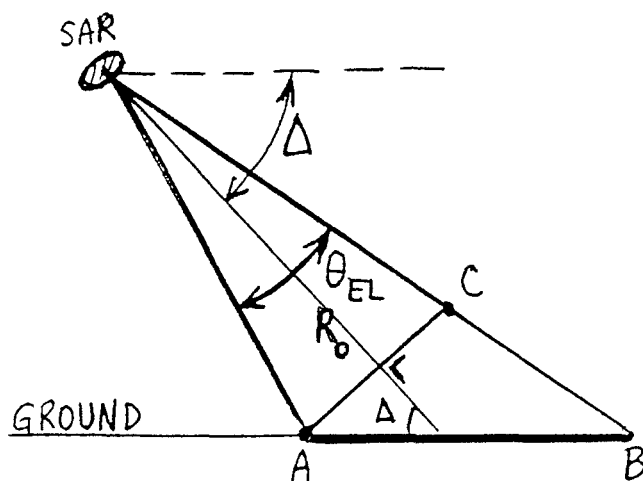


Fig. B48. Swath width AB and depression angle

B. 11. Special SAR's Characteristics and Working Modes.

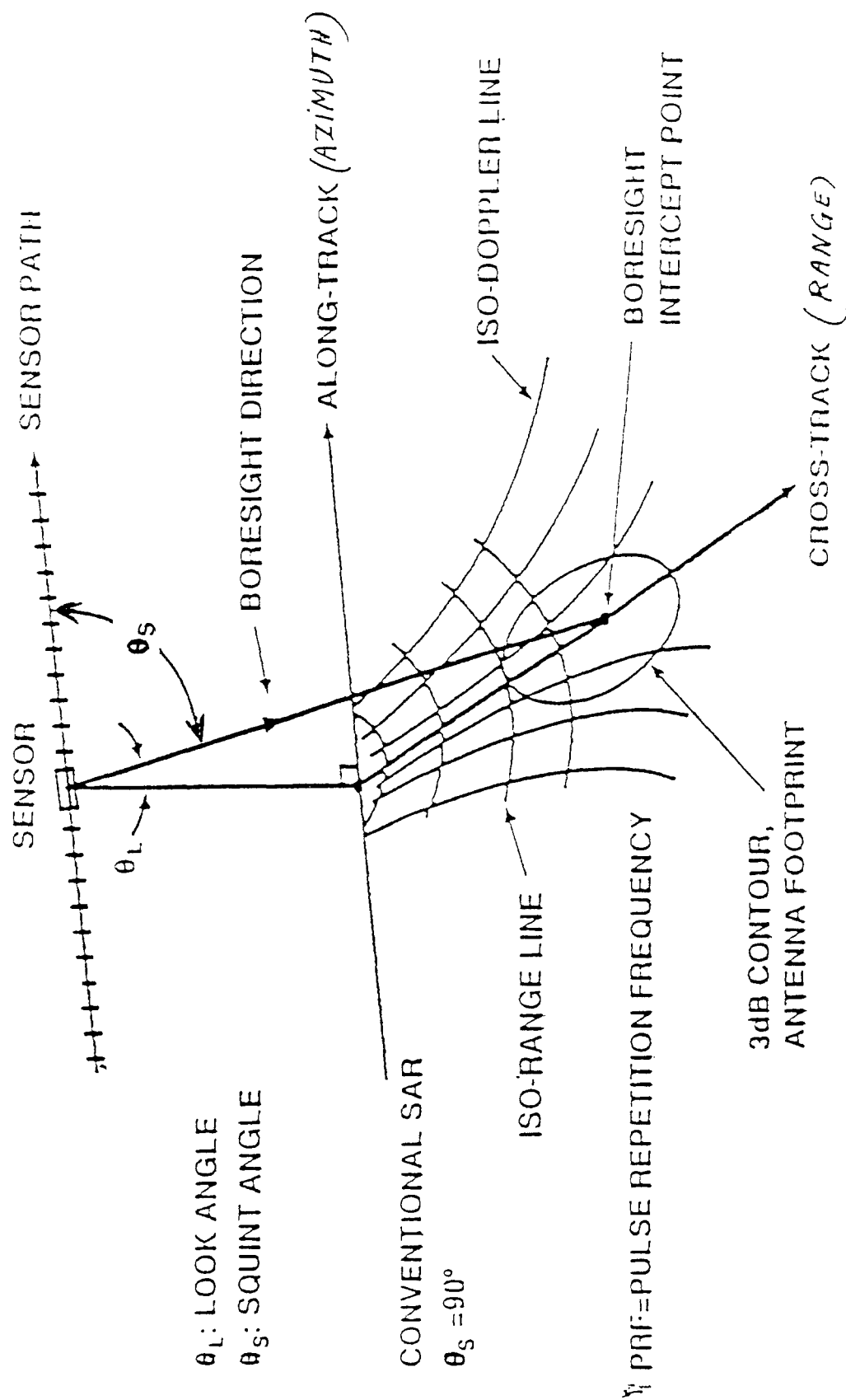
a. Speckle

Speckle is the "grainy" appearance of radar images. Speckle is caused by a combination of scattering from lots of small scatterers within a pixel, i.e., speckle is a scattering phenomenon which arises because the resolution of the sensor is not sufficient to resolve individual scatterers. It is not limited to SAR: speckle can also be seen by looking to the reflection of an expanded laser beam on a "smooth" wall.

b. Squint angle and squinted SAR mode

In the conventional strip-map SAR, the antenna is fixed, pointing purely broadside with respect to the platform (or sensor) path. In the squinted SAR mode, the antenna is also a fixed pointing one, however (see fig. B49 where the squint angle θ_s is different from 90°) the antenna pointing direction in azimuth is not longer perpendicular to the sensor path but squinted by a positive or negative angle θ_s , called squint angle (in a conventional strip-map SAR, $\theta_s = 90^\circ$). The squinted SAR mode could be of interest in a multi jammer environment to look through the individual jammers. Another advantage of this mode could be a "mapping ahead" capability.

SAR PROCESSING
 FIG. B49. SAR MAPPING GEOMETRY



c. Unfocused synthetic antenna

The conventional strip-map SAR with azimuth resolution $\delta AZ = \frac{D}{2}$ was previously called "focused synthetic antenna". Only from a historical viewpoint, it is worthwhile to mention the so-called

"unfocused synthetic antenna"

which is nothing else but the well known "Doppler beam sharpening" featuring an azimuth resolution

$$\delta AZ = \frac{1}{2} \sqrt{\lambda \cdot R_0}$$

(given by curve B of fig.B50)

to be compared with the azimuth resolution

$$\delta AZ = \frac{\lambda R_0}{D}$$

(given by curve A of fig.B50)

of a conventional SLAR
and the azimuth resolution

$$\delta AZ = \frac{D}{2}$$

(given by curve C of fig.B50)

of the strip-map SAR (or focused synthetic antenna)

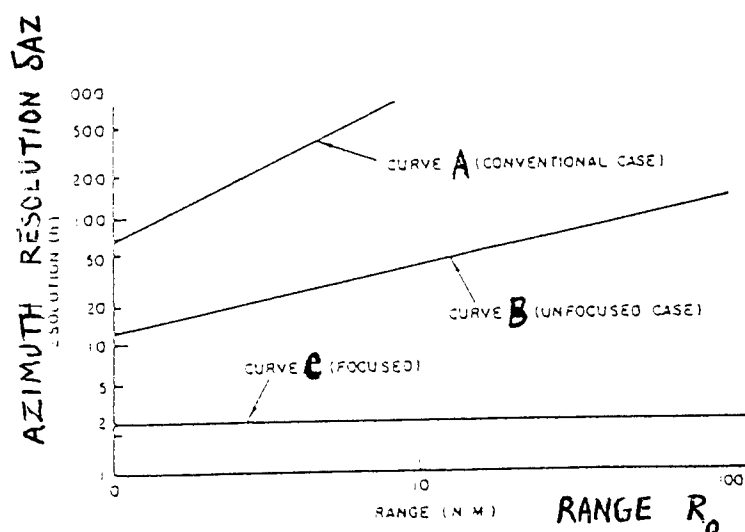


Fig. B50. Azimuth Resolution for Three Cases (A) Conventional (B) Unfocused, and (C) Focused

d. ISAR (Inverse Synthetic Aperture Radar)

Instead of using a moving radar to perform imaging of stationary targets, the radar can be held stationary and the target is moving. One distinguishes

- cooperative ISAR
- non-cooperative ISAR

Commonly the term ISAR is used to refer to a configuration where the radar is fixed and the target moves relative to the radar. A special case in this sense is the imaging of rotating targets without

transversal velocity components. In this configuration, the theoretical cross range resolution limit is given by

$$\rho = \frac{\lambda}{2} \cdot \delta\Omega$$

It depends on the wavelength λ and the observation angle $\delta\Omega$ only.

e. Synthetic aperture sonar

Sonars and radars have wavelengths λ of the same order. But the sonar wave velocity is 200,000 times smaller than the radar wave velocity. Additionally, water salinity and temperature influence the sonar wave propagation. Synthetic aperture sonar systems are about 15 years behind the SAR's; nevertheless they are interesting to find mines on the sea bed; usually those mines have a big shadow behind the mine.

Acknowledgements

The author expresses his sincere thanks to Mr Philippe Dechamps and Mr Michel Jacob for typing the manuscript.

References

- [BR] R. BRAHAM: "Aerospace & Military", IEEE Spectrum, January 1999, pp.73-78
- [BU] S. BUCKREUTH et al.: "New Aspects on SAR Signal processing", DLR-Nachrichten, Heft 86, June 1997, pp.42-50
- [DA] Flt Lt DAWSON: "The repair of aircraft fuel tanks in the RAF", RAF CSDE Annual Report 1998/99, p.54
- [TH] E.C. THIEDE & J.A. HOSCHETTE: "Sensor Fusion", SAL Seminar, Munich, Nov 1990
- [GR] H. GRIFFITHS: "Tutorial lecture about Synthetic Aperture Radar", RADAR'99, BREST, May 1999.
- [FR] A. FREEMAN: "SC 38; Synthetic Aperture Radar: Understanding the Imagery", SPIE's 11th Symposium on Aerosense, Orlando, 20-25 April 1997
- [SM] B.L. SMITH & M.H. CARPENTIER: "The Microwave Engineering handbook", Vol.1, Chapman & Hall, pp.416-420
- [AG] AGARD Lecture Series 182 "Fundamentals and Special Problems of SAR", August 1992
- [RTO] RTO-MP-12, AC/323 (SCI) TP/4 "Sensor Data Fusion and Integration of the Human Element", SCI Symposium, Ottawa, 14-17 Sept 1998.
- [LE] B.F. LEVINE, "Long wavelength GaAs QW Infrared Photodetectors", SPIE VOL. 1362, 1990, pp.163-167
- [CU] L.J. CUTRONA, E.N. LEITH, L.J. PORCELLO, W.E. VIVIAN, "On the application of coherent optical processing techniques to synthetic-aperture radar", Proc. IEEE, Vol. 54, No.8, Aug 1966, p.1026