

J.R. ABRAIAL

STRUCTURES DE
DONNÉES

1970-71

Chapitre 1

Point de vue existentiel

1	Point de vue existentiel	15
1.1	Exemple de la vie courante	15
1.2	Donnée enregistrée et Donnée calculée	15
1.3	Les Espaces d'une donnée	17
1.3.1	Espace des Lieux	32
1.3.2	Espace des Noms	42
1.3.3	Espace des Valeurs	47
1.3.4	Récapitulation et Exemples	49
1.3.4.1	Exemple 1: Data Set	49
1.3.4.2	Exemple 2: Enregistrement	51
1.3.4.3	Exemple 3: Allocation	51
1.3.4.4	Exemple 4: "Base of Data"	51
1.3.4.5	Exemple 5: "Record Class"	60
1.3.4.6	Exemple 6: Système de l'ajout	62
1.3.4.7	Exemple 7: "Entités" dans le système JOCRATE	63
1.4	Opérations sur une donnée	66

1.4.1. Définition	67
1.4.2. Opérations de création (Allocation, baptême, initialisation)	72
1.4.3. Opérations de suppression (deallocation, débaptisation, dévalorisation)	77
1.4.4. Opérations d'accès	81
1.4.4.1. Accès physique	81
1.4.4.2. Accès logique	86
1.4.4.3. Recapitulation et autres opérations	96
1.5. Bibliographie	



1 Point de vue existentiel

1.1.2 Exemple de la vie courante

Nous allons faire une analogie entre une donnée et une personne vivante.

Une personne habite (on est d') un certain lieu, elle porte un nom et elle a des caractéristiques qui la différencient des autres personnes - Ces caractéristiques ont certaines valeurs - Par ailleurs avant sa naissance une personne était dans le néant, elle y retourne à la mort - Enfin, au cours de sa vie une personne noue toute sorte de relations avec d'autres personnes -

La vie d'une personne est essentiellement une série d'actions qui consistent à

- la faire changer de lieu
- la faire changer de valeurs
- éventuellement la faire changer de nom -

La naissance d'une personne (1.1)
est une opération qui consiste à

- lui donner un nom
- lui fournir un lieu
- lui donner quelques premières valeurs

La mort d'une personne est une
opération qui la rend inexistante (sauf
dans les mémoires) donc ^(constant) à lui
supprimer :

- nom
- lieu
- valeurs

Lorsque deux personnes parlent d'une
troisième elles peuvent la désigner de
diverses façons. Par exemple des
Napoleon sera désigné par

- l'Empereur des Français (de France)
- Bonaparte
- le petit laporal

Voici d'autres exemples :

Désignation par liens :

(12)

Roi de France

Duchesse de Lamballe

Pépé le Mocco

l'homme de la Magnon

le locataire du 5^e

le lit 104 - B (hôpital)

Désignation par Nom

Louis XIV

Buffalo - Bill

Le fus

Dupont

1-38-11-78-646-010

Désignation par valeur

l'homme à l'oreille cassée

l'homme au Mas par de Fer

la résine latière (hôpital)

Sur ces exemples nous pouvons voir apparaître les caractéristiques de ces différents

ins des de désignation :

(13)

La désignation par lien est non
ambigue à un instant donné mais
elle n'a pas cours indéfiniment

Exemple : Un roi de France peut
être déchu, un malade peut guérir
le 15 104 - B

La désignation par valeur est très
ambigue et non stable dans le
temps -

La désignation par nom est (en
principe) la seule qui soit toujours
valable - En fait le nom est
souvent la juxtaposition d'un certain
nombre de valeurs ou d'éléments du
lien de telle sorte qu'ils forment
un tout suffisamment représentatif
pour être non ambigue : Ceci est
typiquement le cas du N° de Sécu-
rité sociale qui est un mélange de la
date de naissance, du sexe et d'un numéro
d'ordre d'une part et du lieu de

naissance d'autre point -

(16)

Un dernier exemple montre à quel point ces modes de désignation sont constants : On trouve bien un Botin par nom, un Botin par me et un Botin par profession.

On peut enfin désigner une personne par une relation par rapport à une ou plusieurs autres :

Exemple :

l'amant de Lady Chatterley

le Père du Régiment

Cette manière peut être ambiguë ou non suivant la nature de la relation

On peut aussi désigner une personne par l'intermédiaire d'un lien :

Exemple : lorsque l'on trouve l'avis suivant devant une porte

chose : " Changement d'adresse - (15)
la nouvelle adresse est "

on peut alors désigner une personne
par l'adresse indiquée à une
certaine adresse etc - - -

1.2 Donnée enregistrée et Donnée calculée

Avant d'étudier systématiquement
la structure d'une donnée nous
allons éliminer un cas particulier.

En informatique une donnée peut
être enregistrée (stockée en mémoire)
ou calculée (si l'on donne le mode
de calcul).

Exemple : soit la donnée dont
le nom est $\sin(30)$.

Si sin est un tableau, $\sin(30)$
indique le lieu où est stockée la
valeur correspondante.

Si sin est un programme $\sin(30)$

est un appel du programme avec (16)
comme paramètre actuel la valeur
30 - le résultat de l'exécution de
ce programme est la valeur de
 $\text{fin}(30)$ -

En fait, que l'on soit dans l'un
ou l'autre cas, les problèmes sont
les mêmes : Dans le 1^{er} cas on
recherche une donnée particulière
le programme fin que l'on désigne
par son nom.

Par la suite, sauf cas explicite,
nous considérerons que toutes les
données sont enregistrées -

1.3 Les espaces d'une donnée

Comme une personne une donnée a un
nom, un lieu (celui où elle est stockée)
et une valeur (stockée dans le lieu) -

Nous allons définir avec précision
ces trois notions -

1.3.1 Espace des liens

(17)

Nous parlons d'un lien des que nous pouvons stocker de l'information -

L'unité de lien est donc la position binaire -

La mémoire (ou espace des liens) est donc une réserve de positions binaires.

Une donnée peut avoir besoin de plusieurs positions binaires pour être stockée. Le lien d'une donnée est donc caractérisé par sa taille (en nombre de position binaire) - La taille d'un lien pour une même donnée peut varier dans le temps -

Une autre caractéristique d'un lien d'adresse est le nom de la donnée qui l'occupe - Si ce nom n'existe pas on dira que le lien est libre - On dit q aussi que c'est un trou -

Un lien ^{est} a un nom qui "représente" ~~l'adresse~~ ce lien par rapport à tous les autres - le nom d'un lien

est son adresse. L'adresse d'un @
lib. est en général un paramètre
pour d'un algorithme de sélection.

Cet algorithme de sélection peut
être purement hardware.

Exemple:

circuit de sélection, des registres

circuit de sélection, de la mémoire
centrale -

ou un mélange software - hardware

Exemple

handler + armoire de contrôle

pour une mémoire secondaire -

L'espace des lieux dans une installa-
tion donnée est généralement

- non homogène
- dynamique

Il est non homogène car il existe
~~donc~~ plusieurs algorithmes de
sélection différents

Il est dynamique si sa taille globale ⁽¹⁰⁾ peut varier -

On appelle support le sous ensemble d'un espace des lieux qui est homogène

A tout support correspond un temps d'accès moyen et une durée de séjour moyen de l'information

Le tableau ci-dessous récapitule les diverses caractéristiques

	Registre	^{mémoire} Tôres	Tambour	Disque	Bande
Taille (en octet)	100	100 K à 1000 K	1000 K. à 10.000 K	∞	∞
Temps d'accès moyen	ms	μs	10 ms	100 ms	1 min
Durée de séjour moyen	ans	10 ms à min	10 ms à heure	∞	∞
dynamique	non	non	non	oui	oui
Prix					

On voit sur ce tableau que les ⁽²⁾ supports disque et bande sont les seuls à pouvoir stocker de l'information en permanence et que leur structure dynamique leur permet de suivre les pulsations (naissance, mort) des données.

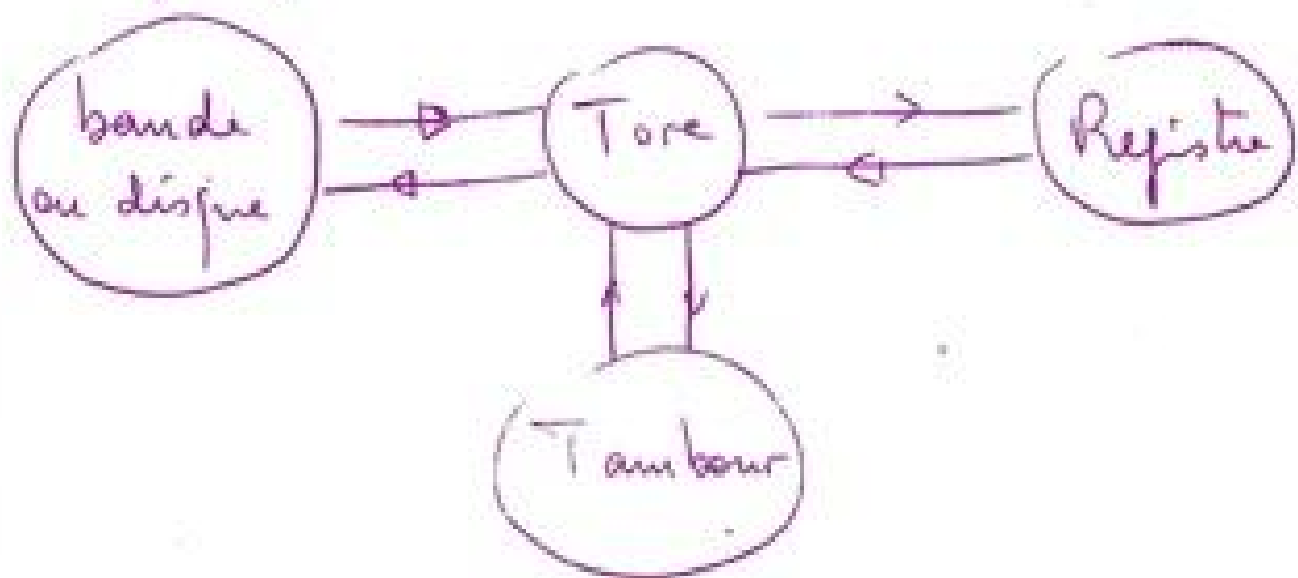
Les supports Tambour et Tore servent à stocker momentanément de l'information venant des bandes ou disque et qui servent de reservoir sous pen.

Le support Register contient l'information en cours de traitement par l'unité de calcul.

On voit donc que, en fait, le véritable espace des lieux est le support disque (ou bande) et qu'il existe ensuite une hiérarchie des supports

de plus en plus rapide, de plus en ⁽²⁾
plus petit et de plus en plus cher.

Suivant l'usage que l'on fait
d'une donnée elle va donc subir
des déplacements correspondant
au graphe suivant



Pour une même donnée on pourra
par exemple avoir le voyage suivant

Disque - Tore - Registre - Tore - Registre -
Tore - Tambour - Tore - Registre - Tore -
Tambour - Tore - Disque

La Gestion de l'espace des lieux ⁽²⁾
est donc en fait en plusieurs
tâches bien séparées :

~~Le fait de lier un lieu à un support~~
Comme nous l'avons vu un lieu
peut être caractérisé par

- son nom (son adresse)

- ses caractéristiques physiques

 - sa taille

 - le fait qu'une donnée

l'occupe ou non

 - le nom (éventuellement)

de la donnée qui l'occupe

En fait pour tous les lieux d'un
support, ces éléments sont rassemblés
et constituent un ensemble de donnée
en soi que l'on appelle la table

d'occupation du support...

(43)

généralement le nom de la donnée qui occupe un lien est le nom du lien qui occupe cette donnée dans ~~un~~^{le} support situé à un niveau juste inférieur dans la hiérarchie.

Exemple :

En mémoire centrale on pourra trouver une table des buffers (un buffer est un lien en mémoire centrale.)

Contenant pour chaque buffer

- son adresse (en mémoire centrale.)
- sa taille
- l'adresse tampon de la donnée

affiliée.

Sur le tambour on pourra alors trouver une table d'occupation (cette table pourra être en mémoire centrale) contenant pour chaque ~~adresse~~^{lien} son le

tambour :

(26)

! son adresse (sur le tambour)

sa taille

l'adresse disque de la donnée
associée - (ou la bande)

En fin on trouvera sur le disque
une table d'occupation du volume
contenant pour chaque lien de ce
disque (ou de cette bande)

son adresse (disque)

sa taille

son nom (son "adresse" dans
le monde extérieur)

La gestion des liens sur un support
consiste donc à gérer la table
d'occupation correspondante de façon
à ce qu'à chaque naissance (venue
d'une nouvelle donnée dans les support
de niveau inférieur dans la hiérarchie)
la nouvelle donnée puisse trouver

un lien fin lui conviendra (de taille ≥ 1 suffisante) -

Deux cas peuvent se présenter

a) En cherchant dans la table d'occupation on trouve un lien ^{libre et} de taille suffisante - C'est ce lien que la donnée nouvelle va maintenant occuper

b) En cherchant dans la table d'occupation on ne trouve pas un lien de taille suffisante - On peut alors reorganiser les liens sur le support (par exemple compactage) de façon à ^{transformer} ~~passer~~ de petits trous ^{en} un trou unique plus grand - Après réorganisation 2 cas peuvent encore se présenter :

(31) La reorganisation donne satisfaction ; c'est à dire qu'elle a pu produire un lien de taille suffisante pour la nouvelle donnée -

(12) de l'organisation on donne (26)
pas satisfaction (saturation du
support) - Si le support n'est pas
dynamique c'est qu'il occupe une
position haute dans la hiérarchie
(tore ou tambour) - On va alors
dans ce cas dégrader ~~des~~ une ou
plusieurs données de support
en les faisant rejoindre un support
de niveau inférieur jusqu'à ce que
l'on obtienne un lien de table suffi-
sant pour la nouvelle donnée.

Exemple : - Dans le cas d'une
saturation les tores on fait descendre
une donnée sur tambour

- Dans le cas d'une saturation
des tambours on fait descendre
une donnée sur disque

- Si le support qui sature
est dynamique on peut l'agrandir

en mettant la nouvelle donnée sur (2)
un nouveau volume (nouveau disque,
nouvelle bande) -

À ce stade il se pose un nouveau
problème : En effet si le support est
dynamique il n'en est pas de même
du support du support

Exemple : Dans une installation
de disque le nombre de disques conte-
nant de l'information est a priori
quelconque ; par contre le nombre de
tourne disque est lui constant -

L'extension dynamique est donc une
opération longue qui nécessitera final-
lement le démontage d'un disque
(ou d'une bande) et le remontage d'une
autre -

On peut dire qu'au niveau des
support de support on a une par

analogie avec celle des supports une (2)
hiérarchie entre

le tambour disque (ou le dérouleur)

et l'armoire à disque (ou à bande)

Les tables d'occupation correspondantes
sont respectivement

les étiquettes sur les disques/ou tambour;

la liste collée sur la porte de l'armoire

Quant au "programme" de gestion ~~des~~
et de sélection des supports son rôle
est joué par l'opérateur -

d'organigramme qui suit récapitule
les opérations de gestion postées
à la naissance d'une donnée sur un
support

Il existe de la place d'op. sur la table d'occupation? (A)

non

Reorganisation de la table d'occupation

Il existe de la place d'op. sur la table d'occupation? (A)

non

Le support est capable d'extension dynamique. (B)

non

Degradation d'un certain nombre de données que l'on envoie sur le support de hiérarchie immédiatement inférieure pour faire de la place. (A)

Il existe de la place d'après
la table d'occupation du
support du support (il y a un
dérivateur ou un bonnet de
libre) ?

oui

Reorganisation du support du
support en démontant un volume
(dérivateur ou ~~bonnet~~^{bonnet}) pour le mettre
dans l'anneau de façon à
faire de la place

Montage d'un nouveau
volume et fabrication
d'une table d'occupation
pour ce volume (formatage)



Nous avons donné ici un schéma
très général des opérations qui peuvent
se passer à la naissance d'une donnée.

Nous verrons plus loin au § 1.4 les
schémas correspondants à d'autres
opérations : mise à jour, mort etc...

Les solutions techniques données à
ce schéma général sont extrêmement
variables d'un système à l'autre.

Il n'est pas question ici d'en faire une ⁽³²⁾ étude exhaustive; nous donnerons cependant quelques détails au § 2.2.2.

Le terme "table d'occupation" que nous avons employé ici est à prendre dans un sens très général - la table d'occupation et le système de données servant à décrire les liens d'un support -

Le terme "table d'occupation" sera à mettre au même niveau que "catalogue" lorsque nous étudierons l'espace des noms au § 1.3.2 ou "table des symboles" lorsque nous étudierons l'espace des valeurs au § 1.3.3.

Table d'occupation et hiérarchie des supports sont les notions les plus importantes de ce §.

1.3.2 = Espace des Noms

d'espace des noms a longtemps (et est encore) confondu avec l'espace des liens - Cette confusion vient du fait

que les données ont ~~été~~ d'abord été ⁽³³⁾
statiques et à durée de vie limitée :

elles naissaient toutes en même temps
elles mourraient toutes en même temps -

Par ailleurs, les données ne subissaient
pas d'extension dynamique (leur "taille"
restait constante) -

Dans ces conditions on comprend
que, le ~~lien~~ lien d'une donnée étant
unique pour chaque donnée, on pouvait
leur donner comme nous ~~à~~ leur adresse.

Au paragraphe précédent nous avons
un comment

- à un même instant une donnée
pouvait occuper plusieurs lieux différents
(sur différents supports)

- à deux instants différents une
même donnée pouvait occuper deux
lieux différents sur un même support
(ceci est du essentiellement aux

possibilités de l'organisation) --

(34)

Il apparaît donc souhaitable d'avoir
"quelque chose" d'absolu qui représente
~~nécessaire~~ une donnée (l'existence
d'une donnée) indépendamment de
l'espace (les différents supports) et
indépendamment du temps - le "quel
quelque chose" est le nom de la donnée
par analogie avec les propriétés
(presque toujours vraies) du nom
d'un être vivant -

Quelquefois on ne donne pas de nom
aux données, c'est le ^(nom de) ~~trien~~ qui en
tient compte et ceci même pour des
données dynamiques - La programma-
tion de tel système est très délicate
On imagine volontiers toutes les précau-
tions qu'il faut prendre lorsque l'on
change une donnée de place ~~sur~~
sur un support ou qu'on la fait

évoluer dans la hiérarchie des supports (35)

~~*****~~

A propos du nom ~~de~~ une donnée, on peut se poser la question suivante :
qui et comment donne-t-on son nom à une nouvelle donnée qui vient de naître ?

Rappelons d'abord la propriété essentielle d'un nom - la donnée qui le porte est la seule à le porter pendant toute la durée de sa vie -

~~*****~~

Un humain peut être à l'origine de la naissance d'une donnée - (Par exemple création d'un fichier par un programmeur) - le nom est alors proposé par cet humain au système de gestion des noms qui regarde si ce nom proposé n'est pas déjà porté par une autre donnée

Les noms donnés par les humains ③
aux données sont généralement des
chaînes de caractères alphanumériques.
Le test par lequel on valide un nom
se fait en consultant un catalogue.
Le catalogue est l'équivalent de
la table d'occupation ~~provisionnelle~~ des
lieux d'un support, pour les noms.

Généralement, outre le nom, le catalo-
gue contient pour chaque donnée l'adresse
correspondante ~~dessinée~~ dans la table
d'occupation sur le support de plus
bas niveau (support permanent) où
peut se trouver cette donnée.

Ainsi on voit comment, grâce à
la table d'occupation, on peut trouver
connaissant l'adresse, le ~~nom~~ nom
d'une donnée et inversement comment
grâce au catalogue on peut trouver

face au nom, l'adresse d'une donnée (39)
sur son support de plus bas niveau -

~~Un programme peut être à l'origine~~

Un programme peut être à l'origine
de la naissance d'une donnée -
Au lieu de proposer un nom alpha
numérique on peut alors chercher
directement dans ^{un} ~~le~~ catalogue un
nom non utilisé - Comme par
ailleurs un programme mais plus
naturellement de l'information binaire
que de l'information alphabétique
les noms donnés par programme sont
de préférence binaire -

Comme par ailleurs l'espace des
noms ressemble assez à un espace
des lieux homogène * (cf les
expressions sur l'analogie entre
catalogue et table d'occupation),

On appelle souvent ~~le~~ l'espace des (38)
nous binaires, l'espace virtuel -

Ceci est une abstraction très commune
de ~~la~~ Et nous permet de donner les
propriétés d'un espace virtuel

(a) d'adresse virtuelle d'une donnée est
unique pendant toute la vie de la donnée

(b) il existe un moyen permettant de
savoir si une adresse virtuelle correspond
~~ou non~~ ou non à une donnée vivante

(c) l'espace virtuel ne doit jamais
être saturé - Comme il est homogène
(par définition) il doit pouvoir être exten-
sible (ou alors être extrêmement foud)

(d) Il doit exister une correspondance
entre l'espace virtuel et la table
d'occupation des supports -

Dans la pratique on trouve beaucoup
de faux espaces virtuels qui n'ont pas
la propriété d'être des espaces de

norme en particulier parce que la propriété ③
est fondamentale ② n'est pas
satisfaite - Ces faux espaces virtuels
servent uniquement à faire une homopé-
sition artificielle de l'ensemble

(forc + tambour) - Ils rendent
ependant de grands services -

Il existe 2 sortes de vrais espaces
virtuels :

a) Espaces virtuels pour lesquels
la correspondance ④ est une fonction
enregistrée - Le système MULTICS sur
le GE-645 est un prototype d'un
tel système - La mémoire virtuelle est
en fait une mémoire à 2 dimensions :
chaque adresse virtuelle est composée
d'une adresse de segment et d'une
adresse de ligne à l'intérieur d'un
segment - La correspondance ④
se fait alors au moyen d'une

Table des segments puis pour chaque segment d'une table des pages. Nous reviendrons sur ces techniques au § 2.2.2.2. Le test ⑥ s'effectue alors soit en détectant un désordre dans les Tables précédentes. Soit en trouvant une ligne vide dans la table des segments --

β) Espaces virtuels pour lesquels la correspondance ④ est une fonction calculée. Le système SOCRATE est un prototype d'un tel système. Le "calcul" se fait en affectant ~~une~~ un hash-coding sur l'adresse virtuelle. Nous reviendrons sur ces techniques au § 2.2.2.2. Le test ⑥ s'effectue en consultant une chaîne de bits ou en détectant une donnée indéfinie. Dans beaucoup de cas une donnée si elle est créée par un programme peut avoir ^{aussi} un nom alphabétique.

donnée par un humain ... Dans ce (41)
cas il est fréquent d'avoir 2 noms
pour une donnée

un nom externe alpha numérique
un nom interne (son adresse virtuelle)

Il est souhaitable bien entendu que l'on
puisse passer de l'un à l'autre et
vice-versa -

Par extension on appellera catalogue
légende la correspondance entre le nom externe
et l'adresse virtuelle.

Dans MULTICS le catalogue ^(étendu) est
enregistré ; dans LOCRAE le catalogue
est calculé (un hachcoding sur le nom
externe donne le nom interne) -

1.3.3 Espace des valeurs

(43)

La valeur d'une donnée est l'information qui est stockée dans le lieu de cette donnée - Au chapitre 3 (Point de vue relationnel) nous élargirons cette notion de valeur par celle de relation -

Une donnée n'a de sens que par rapport à l'interprétation que l'on fait de sa valeur.

Exemple : Dans un système automatisé de gestion du personnel d'une entreprise, la couleur des cheveux des gens est une information qui n'intervient pas - (Cette donnée n'est pas interprétable par le programmeur de paie par exemple)

Cette remarque paraît très triviale - Elle n'est cependant pas toujours faite par les responsables de la structure des fichiers de gestion qui font stocker

des informations sans chercher s'il (4)
existe vraiment un interprète (humain
ou programme) capable de régir à leur
diverses valeurs.

La manière par laquelle les diverses
valeurs d'une donnée peut être représentée
sur son lien est une convention
qui est fait établir avec les interprètes
de la donnée. Cette convention s'appelle
la codification de la donnée.

Exemple: le sexe d'une personne peut
être masculin ou féminin. Le responsable
d'un fichier peut décider que la
codification est 1 pour masculin, 2
pour féminin. S'il n'a pas passé de
convention (de préférence écrite) avec
le programmeur qui se sert de ce fichier
dans un programme (interprète) dont
le comportement dépend (conditions)
du sexe des personnes, on imagine bien

que le résultat de l'exécution du ④
programme risque d'être complètement
faux. Le cas d'erreurs est très fréquent.

Parmi les éléments de la spécifi-
cation d'une donnée il est fondamental
de fournir la convention correspondant
à une information indéfinie. Bien
des erreurs arrivent parce que des règles
précises n'étant pas prescrites : souvent
"indéfini" est codifié par zéro par
"zéro".

La valeur d'une donnée peut changer
au cours de la vie d'une donnée.

Le changement de la valeur d'une donnée
peut être éventuellement accompagné d'une
variation (augmentation ou diminution)
de la taille de son bien.

On voit donc en résumé que toute
donnée s'accompagne nécessairement

d'un certain nombre de régles définies p(4)
saut :

- le domaine de variation des valeurs de la ~~partie~~ donnée
- la codification de ces valeurs
- la manière par laquelle une donnée peut se développer ou se rétrécir

L'ensemble de ces règles constitue ce que l'on appelle le type de la donnée - le type d'une donnée peut, par analogie, se comparer au "programme génétique" d'un être vivant

Le chapitre 2 (Point de vue morphologique) consistera en une étude détaillée des types de données -

Les données qui sont d'un même type portent un nom génétique -

le nom générique et le type d'une donnée sont définies bien entendu avant la naissance de cette donnée. Cette opération de définition s'appelle une déclaration ; nous y reviendrons au § 1.4.1. (Définition).

Généralement on confond le nom générique d'une donnée avec son nom. Ceci vient du fait qu'à un instant donné dans un système il n'y a qu'une seule donnée correspondant à un certain type : on peut alors lui donner pour nom son nom générique. Dans d'autres cas (Algol), seule la dernière donnée ~~créée~~ ^{crée} d'un certain type est accessible, on lui donne donc par abus le nom générique comme nom.

L'ensemble des ~~des~~ types d'un système constitue un système de données en soi que l'on appelle généralement

la Table des symboles

(17)

1.3.4 Récapitulation et Exemples

Vous pouvez récapituler les 3 derniers paragraphes par le schéma suivant

Élément de la
Table de symboles

Élément de
l'atlas

Nom générique
Codification et (ou) règle de déclassement

Nom de la Donnée





lieu de la
Donnée

Élément de la
Table d'occupation

Les $\rightarrow$ indiquées sur le schéma (4) ont uniformément le sens d'une relation (non forcément réalisée par un pointeur)

Les réalisations pratiques correspondant à ce schéma sont extrêmement diverses : parfois certains éléments du schéma sont totalement absents ; parfois ces éléments sont plus fréquemment réunis ; parfois des intermédiaires viennent compliquer (souvent de façon inextricable)

ce schéma de base - (à croire que leurs auteurs ont voulu réaliser une véritable "jungle" de pointeurs et de liens de façon à ce que les "exploiteurs" s'y perdent)

Nous allons étudier maintenant quelques exemples typiques

1.3.4.1 Exemple 1: Data set

(49)

(Données de l'"operating system" (BOS))

a) Table des symboles

La Table des symboles est fragmentée. Chaque élément en est constitué par le Data Control Block (DCB)

b) Nom générique

Le nom générique est le DDNAME du Data set

c) Catalogue

Le catalogue est représenté par le Job File Control Block (JFCB) qui contient les DD statements et en particulier l'association entre

le nom générique	(DDNAME)
le nom	(DSNAME)
le lien	(Volume)

Parfois si le Data set est catalogué (CATLG), l'association entre le nom et le lien se fait dans ce qui est appelé effectivement le CATALOGUE.

③ Nom de la Donnée

(50)

nom de la donnée est son DSNNAME

④ Lien et Table d'occupation

Lien de la Donnée se trouve sur un volume. La table d'occupation s'appelle

Volume Table of Occupation (VTOC)
après élément de la VTOC, s'appelle

Data Set Control Block (DSCB)

DSCB contient le DSNNAME, la position du lien du data set sur le volume et les caractéristiques (taille, fragmentation etc...) du lien.

⑤ Création et Suppression

opération de création se fait à l'aide d'un DD Statement (NEW)

le Data set ne doit plus servir
indiquera dans le DD statement
à mention DELETE

⑨ Références

(5)

Les Documents suivants donnent tous les détails sur l'organisation de ce système de données :

IBM System / 360 Operating System

Job Control Language page 24

(D Statement)

IBM System / 360 Operating System

Supervisor and Data Management

Services page 65 ... (Interface with the Operating System).

1.3.4.2 Exemple 2: Enregistrement dans un fichier indexé séquentiel

① Table des symboles

La Table des symboles contient la (ou les) structure(s) des Φ enregistrements du fichier. Si le fichier est créé et utilisé par COBOL, Φ cette structure est définie dans la DATA Division du programme COBOL.

Si le fichier est utilisé en P2/P1, la (1) structure de l'enregistrement est finalement défini par une STRUCTURE (déclaration de donnée composée)

La table des symboles sera donc dans ces 2 cas, la "Table des symboles" des compilateurs -

Si le Fichier est utilisé en assembleur la table des symboles n'existe pas : la structure de l'enregistrement est donné de façon externe au programmeur -

(b) Nom générique

Le nom générique est le nom du fichier

(c) Catalogue

Le catalogue est représenté par les INDEX qui contiennent à différents niveau les listes des clés (keys) de différents enregistrements ainsi que des pointeurs vers leur position sur le Data Set

d) Nom de la Donnée

(53)

La donnée est ici un enregistrement (RECORD) - le nom de la donnée est sa clé (KEY) -

e) Lien et Table d'occupation

Le lien de la donnée se trouve sur une Data set (Voir { précédent})

En toute rigueur la table d'occupation n'existe pas puisque l'index donne une correspondance entre nom et lien - cependant l'index ne donne pas une correspondance partielle (toutes les clés ne figurent pas dans l'index : on cherche d'abord dans l'index où les clés sont classées puis systématiquement dans le Data set à partir de l'enregistrement dont la clé est la plus voisine de celle de l'enregistrement cherché : d'en le voir

indépendant) -

(5)

En fait au niveau de chaque enregistrement (dans le lien lui-même) on trouve des informations concernant l'organisation de ce lien: la taille de l'enregistrement (si cette taille est variable) et le fait qu'il correspond à un ~~enregist~~ enregistrement inexistant (X FF : delete code)

(f) Création et suppression

L'opération de création se fait à l'aide d'une opération WRITE indiquant une clé qui n'existe pas déjà dans le fichier

L'opération de suppression se fait à l'aide d'une opération de mise à jour en mettant la valeur du "delete code". Dans cette organisation la suppression n'entraîne pas la deletion

tion automatique des liens -

(5)

③ References

Le document suivant donne les détails sur l'organisation de ce système de données :

IBM System /360 Operating System Supervisor and Data Management Services page 111... (Processing an Indexed Sequential Data set)

1.34.3 Exemple 3 : Allocation automatique des blocs en PL/1

(a) Table des symboles

Le type d'un bloc PL/1 est l'ensemble des déclarations faites à l'intérieur de ce bloc - C'est la Table des symboles du compilateur PL/1 qui contient ces informations

(b) Nom générique

Si le bloc correspond à une procédure le nom générique du bloc est le ~~nom~~ nom de la procédure -

Dans le cas contraire le bloc n'a pas de nom généré -

(c) Catalogue

Il n'y a pas de catalogue car il n'y a qu'un seul bloc qui a un nom : le dernier créé -

(d) Nom

Le nom d'un bloc n'existe que pour le dernier créé : c'est R13 (le Registre 13 contient l'adresse absolue de l'allocation du bloc) -

(e) lien et Table d'occupation

Le lien d'un bloc s'appelle une Dynamic Storage Area (DSA) -

Il se trouve dans une zone de mémoire centrale réservée à l'expansion dynamique des programmes

PL/1 (subpool) -

Des éléments d'organisation de la DSA se trouvent dans la DSA elle-même (longueur, chaînage) - Par

ailleurs une liste des éléments libérés
de la "subpool" indiquent les éléments
qui peuvent constituer une DSA à
la création d'un bloc.

(f) Création suppression

La création d'un bloc se fait
par une opération appelée PROLOGUE.

Cette opération appelle une routine
de la bibliothèque PL/1 appelée
nommée IHESADA, laquelle se
sert de la macro OS : getmain.

La suppression d'un bloc se fait par
l'opération EPILOGUE faisant
appel à IHESAFA laquelle utilise
la macro OS : freemain.

(g) Référence

On trouvera de plus amples détails sur
l'organisation de ce système de données
dans :

IBM System /360 Operating System

PL/1 Programmer's Guide page 130
(Run Time Stack)

et dans:
IBM System/360 Operating system
supervisor and Data Management
services page 44 ... (Main Storage
Management) -

13.4.4 Exemple 4 : Données basées
en PL/I (Based Data)

(a) Table de Symboles

Celle du compilateur ou si la donnée
est une structure on a un "DOPE VECTOR"
- qui décrit la structure de la valeur
de la donnée - (ou mieux la classe de
la donnée)

(b) Nom générique

Ce nom générique est le nom de la
donnée basé à sa déclaration

(c) Catégorie

Il n'existe pas car on ne donne pas
(malheureusement) de nom à la
création

(d) Nom of (e)

(e) Lien et Table d'occupation (59)

Le lien pour une donnée basée de PL/1 se trouve à l'intérieur d'une zone appelée AREA -

La Table d'occupation est dispersée dans l'AREA. Chaque élément ^(de l'AREA) est

- soit occupé par une donnée basée (Allocated) dont la taille est indiquée

- soit libre (free) - les éléments libres de l'AREA forment une liste libre dont les éléments sont chaînés entre eux par taille croissante.

(f) Création et suppression

La création se fait à l'aide de l'opération ALLOCATE et la suppression à l'aide de l'opération FREE -

Une variable POINTE 2 (faisant office de nom) contient l'adresse de la donnée basée lors de sa création. Cependant on ne peut pas dire free

le nom de la variable POINTER (c)
est le nom de la donnée basée
car un motant donné plusieurs POINTER
peuvent être assignés à une même donnée
et de plus un POINTER peut changer
d'assignation.

① References

De plus amples détails sur ce système
de données peuvent être trouvés dans
- IBM System /360 Operating System
PL/1 Programmer's guide - page 145...
(Allocation and Release of Storage in an
Area)

- IBM System /360 Operating System
PL/1 Language Reference Manual - page 172...
(Based Variables and List Processing)

1.3.4.5 Exemple 5 : Record class en Algol W

Ce type de donnée est très proche
d'une donnée basée de PL/1

Les variables de type REFERENCE, (6)
jouent le rôle des variables de type
POINTER de PL/I

La différence vient des opérations
de création et de suppression.

En Algol W la création est
explícite comme en PL/I mais
il n'existe pas d'opérateur spécial.
c'est l'affectation d'une "Record
class Identifier" à une variable
de ~~type~~ référence qui en tient lieu.

La suppression explicite en PL/I
(FREE) est ici implicite : une
record d'Algol W est implicitement
supprimé lorsqu'il n'existe plus
aucune variable de type référence
pointant vers lui.

1.3.4 (Exemple) : Système de pagination (b)
décrit dans l'article : Use of Fast
and slow memory in List Processing
language par J. Cohen.

(a) Table de symboles

Il n'y a pas de table de symboles.
Toutes les "données" sont de même type :
ce sont des pages de longueur fixe.

(b) Nom générique

Pas de nom générique.

(c) Catalogue

Le Catalogue est un Tableau indexé par
l'adresse disque de la page.

(d) Nom

L'adresse disque de la page constitue son
nom.

(e) Lien et Table d'occupation

Le lien d'une page se trouve en mémoire
centrale dans une zone fixe pouvant
contenir un certain nombre de pages.

L'Table d'occupation donne pour chaque position dans la zone dynamique des pages, le nom de la page qui l'occupe et d'autres renseignements annexes -

④ Création et suppression

Ces opérations correspondent à l'entrée d'une page de la mémoire secondaire dans une position en mémoire centrale et à la sortie d'une page en mémoire secondaire (ou éventuellement le recouvrement d'une position par ~~sa~~ l'entrée d'une autre page)

⑤ Références

Pour ^{de} plus amples renseignements voir l'article Use of Fast and slow memory in List Processing Language par J. C. C. A.C.M. Vol 10. N°2 - Feb. 67

1. 3. 4. 7 Exemple 7: Entités dans le système Isolate -

① Table des symboles

La table des symboles est la structure

Interne de la Banque de Données - (1)
Une entité a une description complète
de ses composants dans la structure
Interne -

(b) Nom générique

le nom générique est le nom de l'entité
donnée lors de la définition de la Banque
de Données

(c) Catalogue

Le catalogue n'existe pas explicitement : c'est l'ensemble des adresses
virtuelles - via correspondance entre
une entrée du catalogue $\frac{1}{2}$ et le lien
se fait par hashing - On peut
dire donc que le catalogue a
bien d'être enregistré comme habituel
lement est calculé par le programme
de hashing -

(d) Nom

Le nom d'une donnée est son
adresse virtuelle -

(c) Lien et Table d'occupation (6)

Le lien est l' Espace Reel -
la Table d'occupation des lieux
est rassemblée (Technique du diction-
naire rassemblée) en dispersée -

Le lien d'une donnée est ce fait
éclaté en ^{diverses} sous-pages (blocs de
longueur fixe) dans l'espace réel -
la Table d'occupation fixe les sous-
pages.

(d) Création / suppression

Les opérations sont explicites (b, t)

(e) Références

Pour de plus amples détails Voir
Projet sonate - Université de
Jocoble -

1.4 Opérations sur une donnée (61)

Nous venons d'introduire pour une donnée trois espaces :

- Noms
- lieux
- Valeurs

Si nous introduisons de plus l'espace vide $\emptyset$, nous pouvons alors envisager les opérations sur une donnée comme des opérations faisant passer d'un espace au même ou à un autre, le ou les valeurs dans les ou les ^{autres} espaces restant constantes.

Nous pouvons alors schématiser ces opérations sur la matrice de transition suivante.

	L	N	V	A
L	^{logique} Déplacement physique	^{logique} Désignation physique	^{logique} Accès physique	^{logique} Deallocation
N	^{logique} Localisation logique	^{logique} Changement de nom ?	^{logique} Accès logique	^{logique} Suppression
V	^{logique} Localisation associative	^{logique} Désignation associative	^{logique} Titre à four	^{logique} Dévaluation
A	^{logique} Allocation	^{logique} Baptême	^{logique} Initialisation	^{logique} Suppression

Création

Avant d'étudier ces diverses opérations nous allons d'abord considérer l'opération de définition.

1.6.1 Définition

L'opération de définition est celle qui consiste à "déclarer" un nouveau type (une nouvelle classe) de donnée - C'est une opération de création

dans le système de données constitué (61)
par - nom générique (Nom)
- Table des symboles (Lieu)
- Organisation et codification (Vale)

Exemple
lorsque l'on fait la déclaration
suivante en PL/1 :

DECL TOTO FIXED BINARY (31) ;

~~On~~ on crée une nouvelle classe de
données de nom générique TOTO
et d'organisation FIXED BINARY (31)

C'est bien pour le système de
données explicité plus haut la
création d'une donnée puis ~~puis~~ l'on
a une opération de baptême et d'initia-
lisation - c'est le compilateur qui
fait l'allocation de cette "méta donnée"
dans la Table des symboles.

Lors de la définition on fixe généra-
lement la codification des données futu-
res de cette classe : ainsi dans

l'exemple ci-dessus les données de la classe TOTO seront des nombres binaires entiers codés sur 31 bits -

La codification n'est pas toujours aussi explicite - Ainsi en ALGOL on écrit par exemple "INTEGER A;" c'est alors la version du compilateur ALGOL avec laquelle on travaille qui définit implicitement la codification des futures données de la classe -

Lorsque les données de la classe ~~possèdent~~ doivent porter chacune un nom, la codification de celui-ci peut être défini lors de la déclaration de la classe

Exemple:

Lorsque l'on veut définir un fichier sous OS on emploie la macro DCE. Parmi les ^(nombreux) paramètres de cette macro figure bien entendu le DDNAME

qui est le nom générique et aussi (F)
un paramètre KEYLEN qui donne
la longueur (en caractère) de sa clé
des enregistrements. On a donc
bien ici, lors de la définition,
une spécification de la codification
des noms des futurs données de la
classe -

La définition d'une classe de données
peut être une opération ponctuelle (qui
se fait en une seule fois) : c'est
le cas des données définies dans un
programme écrit dans un langage de
haut niveau (c'est la déclaration)

Par contre, malheureusement, une
classe de données ^{peut être} ~~est~~ définie en plusieurs
fois dans certains cas. Ceci complique
le processus de définition
très et amène ~~des erreurs~~...

Exemple - en fait (cf 1374). (3
peut en effet être partiellement défini
dans une macro DEC, mais aussi
dans le DD STATEMENT des cartes
de contrôle (paramètre introduit par
DEC) ~~et les autres~~
~~parfois~~ ~~et les autres~~ enfin
d'autres informations de définitions
peuvent être contenues dans la
Table d'occupation. Bien entendu
ces diverses sources d'information définis-
sant une même classe de données (un
même fichier) peuvent être incohérentes
et les auteurs ont donc défini des
règles de précedence - (Pour de
plus amples détails voir le document
IBM system/360 Operating System
Supervisor and Data Management
Services page 65 - (Interface with the
operating system))

Pouvoir définir une donnée en plusieurs fois, c'est évidemment plus souple mais cela complique peut être arbitrairement l'organisation des données.

1.4.2 Opérations de création
(Allocation, baptême, initialisation)
De ces trois opérations, c'est l'allocation qui est la plus importante et qui existe toujours.

Créer une donnée suppose qu'il en existe un nom générique. Dans la pratique on trouve deux classes de création :

- ou bien la Table des symboles est présente (c'est le cas pour les fichiers en général)

- ou bien la table des symboles n'est pas présente (c'est le cas pour les ~~app~~ données définies dans un programme)

pour l'auto-évaluation.

Dans tous les cas l'allocation consiste à chercher dans la Table d'occupation des lieux, un lieu qui est apte à contenir la nouvelle donnée - sa d'information concernant la taille du lieu se trouve soit dans la table des symboles, soit dans les éléments de valeurs fournies par l'initialisation. (Si la Table des symboles n'est pas présente lors de la création d'information concernant la taille du lieu est indiquée dans un programme de création généré par le compilateur pour cette classe de données)

L'allocation d'une donnée peut être une opération automatique ~~non~~ (c'est le cas des données déclarées à l'intérieur d'un bloc en PL/s ou en ALPOL) ou manuelle (intervention humaine ~~et~~ explicite ou par l'intermédiaire

l'aire d'un programme : c'est le cas, @
des données bases de PL/1, des
enregistrements de fichiers et des
data sets) -

On appelle souvent dans la littérature
le moment où l'on "accroche" une
donnée nouvelle à son nom fini ou fin
le "binding-time"

Exemple : x les données dites STATIc
de PL/1 ont leur "binding-time" au
début de l'exécution du programme
dans le quel elle sont définies - (quel
queroient leur position dans le programme)

x les données dites AUTOMATIC
de PL/1 ont leur "binding-time" au ~~ed~~
à l'entrée dans le bloc où elle sont
définies -

Le baptême, s'il a lieu, ~~se~~ se passe
toujours "en même temps" que l'allocation.
Comme l'allocation, le baptême peut être
explicite (le nom de la donnée est fourni

soit par une intervention humaine) : le baptême est refusé si le nom fourni existe déjà pour une autre donnée - c'est le baptême ~~est aussi~~ le cas de la création d'un ~~fichier~~ enregistrement de fichier muni d'une clé d'accès -

Le baptême peut être automatique ; dans ce cas c'est le système lui-même qui fournit un nom (une adresse virtuelle) à la donnée nouvellement créée - Dans ce cas le système est obligé de connaître les noms déjà affectés pour donner à la nouvelle donnée un nouveau nom - Dans le système SOCRATE une chaîne de bits sert à indiquer quels sont les noms affectés et ceux qui ne le sont pas - (Pour de plus amples détails voir *Projet Socrate* page 197 (Technique Statique) d'initialisation (~~donnée~~ ~~fourniture~~ d'une première valeur à la donnée) est une opération qui peut avoir lieu en même temps que l'allocation et

le baptême : c'est le cas par exemple
des données dont la déclaration porte
l'attribut INIT en PL/1. Il n'en
est pas toujours ainsi - très souvent l'ini-
tialisation a lieu après la création
proprement dite - Dans ce cas entre
le moment de la création et le moment
de l'initialisation il se passe un
certain temps pendant lequel la
donnée existe (puis qu'elle a été
~~allouée~~ ^{allouée} et baptisée) ~~et~~, est accessible
mais n'a pas de valeur - que se passe-t-il
alors si l'on demande sa valeur -

Dans la pratique 2 cas se présentent
dans le 1^{er} (le plus fréquent) cas, on
a une erreur (le programme s'arrête):
il est interdit de manipuler des données
indéfinies (c'est le cas des données PL/1)
dans le 2^e (Banque de Données converses
trouvées) une donnée indéfinie (non
initialisée) est une donnée licite au
même titre qu'une autre - Il peut en

effet être intéressant de connaître les raisons des données qui n'ont pas de valeurs de façon à pouvoir repérer cette lacune...

Dans ce 2^{ème} cas, l'opération inverse de l'initialisation, la dévalorisation est donc une opération licite qui n'entraîne pas automatiquement la suppression de la donnée. On peut alors avoir pour une même donnée plusieurs initialisations si l'on a en plusieurs dévalorisations:-

1.4.3 Opérations de suppression (Deallocation, débaptisation)

La suppression d'une donnée est une opération ^(toujours) beaucoup plus complexe que sa création.

Nous verrons en effet au chapitre 3 (Point de vue relationnel) qu'une donnée n'est jamais isolée comme elle apparaît dans ce chapitre-ci. En fait, entre les valeurs mortes d'une donnée, & ils

de nous au cours de la vie du programme (donnée toute entrée de données -

les relations se manifestent dans la pratique par des pointeurs (des données dont la valeur est une adresse : ^{(le nom d'un} lieu) ou par des références (des données dont la valeur ~~est~~ est le nom d'une autre) (Attention! le sens que nous donnons ici au terme référence est assez particulier et ^{ce} n'est pas toujours celui-là que l'on trouve dans la pratique où l'on confond souvent référence avec pointeur. cf Algol W)

Bien entendu, lorsqu'une donnée meurt, son lien et son nom doivent disparaître - Or il se trouve que par le jeu des relations que nous venons d'évoquer ces noms et liens sont invoqués par des ~~pointeurs~~ références et des pointeurs

Ainsi avant de supprimer une donnée ^(de décompiler) il est indispensable de rendre indéfini tous les pointeurs ou références qui

caractère de cette donnée - Toute la (??)
complexité de la suppression vient de
cette obligation

Comme pour la création, la suppression
peut être explicite (cas d'un enregistre-
ment d'un fichier) ou implicite (cas
d'une donnée AUTOMATIQUE de PL/1 supprimée
automatiquement à la sortie du bloc où
elle est définie) - Il existe un autre
cas de suppression implicite intéressant.
Une donnée est supprimée lorsqu'elle
n'est plus accessible par aucun pointeur.
C'est le cas des couples (CAR-COR) de
LISP - Nous y revenons au chapitre

3 -

Lorsqu'une donnée est supprimée, sa
deallocation peut être immédiate ou
différée - (le 2^{ème} cas est celui des
enregistrements indexés séquentiel ou
la suppression consiste seulement à
marquer sur l'enregistrement qu'il
est supprimé - X'FF' en tête de

l'enregistrement) - lorsque la deallocation est difficile le lien de toutes les données est reorganisé lorsque la fragmentation devient trop importante (Voir § 1.3.1. Espace des Lumps) -

C'est au cours de cette réorganisation qu'on fait le lien en même temps toutes les deallocations qui avaient été difficiles depuis la dernière réorganisation - Au lieu de faire les deallocations coup par coup on les fait toutes ensemble -

Les techniques de garbage collection que nous verrons au Chapitre 3 font partie de ces méthodes de deallocation difficiles -

La déréférencement consiste, s'il y a lieu, à supprimer le nom de la donnée du Catalogue (s'il existe) - ~~Données~~

1.4.4 Opérations d'accès

(1)

(accès physique, accès logique, accès associatif)

Les opérations d'accès ^(1 logique ou physique) sont celles qui à partir d'un nom ou d'une adresse donnent la valeur d'une donnée.

Par extension on appellera accès associatif l'opération qui à partir de la connaissance d'une partie seulement de la valeur d'une donnée fournit la totalité de la valeur de la donnée.

Les opérations d'accès sont donc les suivantes

$L \rightarrow V$ accès physique

$N \rightarrow V$ accès logique

$V \rightarrow V$ accès associatif

1.4.4.1 Accès physique

L'opération $L \rightarrow V$ est toujours faite par le hardware.

En fait si l'on réfléchit au processus exact de cette opération, on s'aperçoit qu'il s'agit plutôt d'une opération $L \rightarrow L$.

C'est à dire un déplacement -

L'opération $L \rightarrow V$ est la production d'un certain nombre de déplacements pour amener la donnée (ou une partie de celle-ci) à être sur un support où un interpréteur peut la prendre en compte -

Si l'interpréteur est humain, le lien final peut être un papier d'imprimante, un écran cathodique, une carte perforée interprétée etc...

Si l'interpréteur est un programme, le lien final est un registre -

Dans ^{presque} tous les cas la mémoire centrale est la plaque tournante de l'information : on ne peut pas malheureusement faire communiquer directement un disque avec une imprimante ou un perforateur de carte (mais peut être cela existe-t-il après tout ? ...) ~~par contre~~ ...

Par contre ~~par contre~~ les opérations RD, WD (write direct, read direct) permettent

en principe de passer directement d'un (8)
équipement périphérique spécialement
connecté à un registre (et inversement).
(Ceci au prix d'un ralentissement consi-
dérable ~~de~~ de l'Unité Centrale) -

Le rôle de plaque tournante jouée vicie-
sairement par la mémoire centrale a
amené les constructeurs à faire
des développements techniques allant
dans le sens de la multiprogrammation
ou dans le sens du ~~traitement~~ multitrai-
tement -

En effet les successions de dépla-
cements : Disque - Tore - Imprimante ou
Disque - Tore - Perforateur ou
Imprimante - Tore - Disque ou
Perforateur - Tore - Disque ou
les mêmes avec un Tambour (ou un
disque à tête fixe) sont des
opérations de routine extrêmement
fréquentes et longues - On aimerait

donc qu'elles perturbent le moins possible le fonctionnement de l'Unité Centrale. Avec la zone génération, ces opérations peuvent se faire en parallèle avec le fonctionnement de l'Unité Centrale grâce au Canal qui est indépendant. Il est cependant nécessaire d'initialiser ces opérations et de fermer la zone Tore servant d'intermédiaire.

La solution multi programmation consiste à écrire de petits programmes ~~spécifiques~~ spécifiques les syndicats qui

- sont spécifiques d'un couple (par exemple des/na - perforateur)
- possèdent chacun une petite zone de travail en mémoire centrale.

Les syndicats sont réglés par le jeu des interruptions de fin de transfert : ils opèrent par "bouffée" successives.

Le fonctionnement normal de l'Unité Centrale est donc interrompu de temps en temps pour permettre à un syndicat d'exécuter

travailler le transfert d'une bouffée (1)
d'un desfe (ou tambour) vers la
zone Tore ou de sa zone Tore vers
un périphérique lent ou l'inverse.

Certains constructeurs (peut être même
tous) ont pensé que l'action des symbiotes
perturbait par trop le fonctionnement
de l'Unité Centrale et ils ont préféré
mettre un ^(ou plusieurs) petit(s) ordinateur(s) chargé(s)
de faire marcher les symbiotes.
C'est par exemple la solution choisie
par CDC sur la série 6000 - c'est
aussi le rôle du 360/40 lorsqu'il
"pilote" un 360/65 ~~dans~~ ^{avec} le système
ASP/MVT.

En ce qui concerne l'opération
 $L \rightarrow V$ pour laquelle la destina-
tion finale est un registre ~~sur~~ nous
avons déjà fait au § 1.3.1 (Espace
des lieux) un certain nombre de dévelop-
pements -

1.4.2.2 Accès logique
d'accès logique (opération $N \rightarrow V$)
est toujours le produit d'une opération
~~XXXXXX~~ $N \rightarrow L$ par une opération
 $L \rightarrow V$

Il y a lieu de distinguer tout de
suite 2 cas :

(a) le catalogue n'existe pas sous
quelque forme que ce soit (les données
n'ont pas de noms) : La seule manière
de faire un accès logique est alors de
designer la donnée par son nom gé-
nérique (pour lui dire où il est et où il n'est pas)
(pour lui dire où il est et où il n'est pas)
exister sinon il n'est plus possible
de faire un accès dit logique) —

Ce cas est celui des enregistrements
dans un fichier séquentiel ou bien
des données d'Algol ou de P-1 —

Puisque l'on designe la donnée
par son nom générique et non par son
nom il y a la plupart du temps une
ambiguïté puisque à une classe

de donnée ($\bar{x}$ = nom générique) et (d)
correspond plusieurs données -

La seule façon de lever cette ambiguïté consiste à établir une convention de façon à ce qu'un nom générique il ne corresponde qu'une seule donnée de la classe -

Dans le cas d'un fichier séquentiel la convention est la suivante : l'invocation du nom générique pour accéder à une donnée (opération GET) fournit automatiquement, la donnée qui suit dans la séquence celle qui ~~précède~~ a été accédée la plus récemment -

Dans le cas des données d'Alfol ou des données AUTOMATIC de PL/1, l'invocation du nom générique fournit la dernière donnée née de ce type (portant ce nom générique)

Dans ces 2 cas il existe donc une convention qui évolue dynamiquement au cours du temps - Par extension du 1^{er}

Cas on peut ^{dire} qu'il y a un point (13)
constant qui établit la liaison
entre le nom féminin et l'unique
donnée accessible ^{appartenant à cette classe} ~~particulière~~

Dans le cas d'un ~~seul~~ fichier
séquentiel ce point courant se trouve
dans l'élément de la Table des symbo-
les (c'est à dire dans le D.C.B.)

Dans le cas d'une donnée PL/1 on
Alfol le ^{rôle du} ~~point~~ courant est joué par
un registre (qui en fait pointe sur
la zone dynamique correspondant
au dernier bloc créée - La donnée elle-
même est représentée par un déplacement
dans cette zone)

⑥ Le Catalogue existe -

Dans ce cas l'accès binaire consiste
à chercher l'élément du catalogue
qui contient le nom, puis suivant
la réalisation ^(technique) de la liaison entre

l'élément du catalogue et le lien de la donnée à ~~passer~~^{accéder} plus ou moins directement à celle-ci -

Rappelons que le catalogue peut être • totalement incorporé à la donnée (le nom ^{de la donnée} apparaît alors comme une valeur particulière et l'accès logique comme un cas particulier d'un accès associatif) -

• séparé de la donnée mais donnant accès à celle-ci sans passer par une table d'occupation (celle-ci est en fait dans ce cas confondue avec la donnée : la taille de la donnée est en tête du lien de celle-ci -) - le cas est celui d'un enregistrement d'un fichier indexé séquentiel -

- séparé de la donnée mais ne donnant accès à la donnée que par l'intermédiaire d'une Table d'occupation des liens séparée des données. C'est le cas des Data sets -

Dans tout le cas la première (90)
opération à effectuer est une
recherche dans le catalogue -

En fait si l'on considère le
catalogue comme un système de
donnée en soi et ~~les~~ ^{un} nom comme
la valeur d'une donnée, cette opéra-
tion est une ~~relation~~ ^{opér} associative
dans le catalogue -

Nous revenons au chapitre 2
(§ 2.2.2.1.2) sur les techniques
de recherche -

Etant donnée l'importance (fréquence
d'emploi, temps de mise en œuvre) des
opérations de recherche ~~sur~~ ^{en} catalogue,
on a assisté depuis les débuts de l'in-
formatique à une prolifération
extraordinaire des méthodes de recher-
che associatives -

Nous allons cependant donner ici un
aperçu rapide de la manière dont on
peut classer ces méthodes -

Notation Trois éléments sont évidemment ⁽⁹⁾

pour analyser les méthodes de recherche

- le temps moyen de recherche $\bar{T}$

- le ^{encombrement} ~~taille~~ du catalogue : E

- le nombre d'éléments actifs du Catalogue : N

On introduit souvent en plus le coût

du catalogue $C = E \cdot \bar{T}$

Ce coût est intéressant car en général les valeurs de E et $\bar{T}$ varient en sens inverse d'une méthode à l'autre -

En effet si l'on veut construire un catalogue ayant un temps de recherche $\bar{T}$ minimum (on a alors évidemment $\bar{T} = 1$) il est nécessaire d'avoir un catalogue dont la taille est au moins ~~est~~ $\sqrt{N}$ ^{ou au moins} ~~est~~ les valeurs possible que peut prendre un nom - (On imagine la taille d'un ^{tel} catalogue si le nom est un numéro de téléphone local 10^3) - le catalogue est évidemment manifestement vide

Si l'on veut par contre construire un (a)
catalogue qui a un encombrement mini-
mum on disposera les noms qui le
constitue en séquence - on aura alors

$$\begin{cases} \overline{E} = N \\ \overline{E} = \frac{N}{2} \end{cases}$$

On peut grossièrement diviser les
méthodes de recherche en 3 classes

(a) Celles pour lesquelles le temps de
recherche est indépendant de N . Parmi
ces méthodes on trouve l'indexation
directe ou plus haut, celles qui utilisent
des mémoires associatives (dans l'état
actuel de l'art, le prix de ces
mémoires, l'usage en est réduit à de
très petits catalogues cf la mémoire
associative du 360/67), celles qui
simulent par software les mémoires
associatives : ce sont les méthodes
de hash coding. Un hash code est
une fonction qui fait passer du nom
cherché à un index dans le catalogue

l'indexation pure est un hachage dont (93)
la fonction est l'identité - Le problème
du hachage est donc de choisir suivant
la structure des noms une fonction qui
minimise l'encombrement -

La fonction de hachage est donc
une fonction "many-one" (c'est à dire
que plusieurs noms peuvent donner le
même index) - Lorsque deux noms
donnent le même index on a une collision.

Les variations qui existent entre les
diverses méthodes de hachaging correspon-
dent à l'emploi de diverses ~~fonc~~ classes
de fonctions de hachaging et aux
diverses méthodes de règlement des collisions.

En toute rigueur, à cause des collisions,
le temps moyen ~~d'accès~~ de recherche $\bar{T}$
depend de N et on met le Coefficient
de remplissage : $\frac{N}{E}$ -

On peut montrer que si les noms sont
régulièrement disposés dans leur domaine
de variation on a : $\bar{T} = \frac{3}{2} \frac{N}{E} \leq \frac{3}{2}$ ^{d'un catalogue}

Il existe une littérature consi dérable

sur les méthodes de hash coding - Les méthodes de hash coding ont d'abord été développées pour les compilateurs qui ont expérimenté souvent à rechercher les noms génériques des données dans la Table des symboles -

Pour de plus amples détails sur les méthodes de hash coding on peut consulter les articles suivants

- Scatter Storage Techniques par MORRIS
(CACM Vol 11 N1 Janvier 1968)
- An Improved hash code for scatter storage ~~code~~ par MAURER
(CACM Vol 11 N1 Janvier 1968)
- File structures using hashing functions par E. COFFMAN
(CACM Vol 13 N7 July 1970)

et le chapitre 4 (Table) de Elementary Techniques of Compilation par F.R.A. HOPGOOD (Dunod - Traduction 1970)

⑥ Celles pour lesquelles le temps ^①
d'accès est proportionnel à $\log N$
le type même de ces méthodes est
la recherche binaire bien connue (on
dit aussi souvent recherche dichotomi-
que) - Ces méthodes nécessitent
~~pour~~ que les éléments du catalogue
soient classés - Ceci explique pourquoi
on présente souvent dans la littérature
en même temps les méthodes de recherche
et les méthodes de classement -

Nous reviendrons au § 2.2.3 (Arbres)
sur les méthodes de recherche utilisant
la construction d'un arbre -

Pour de plus amples détails sur ces
méthodes voir

- l'ouvrage de I. FLORES - Computer
Sorting

- le Numéro spécial des Communica-
tions de l'A.C.M. (Vol 6, N° 5 May 63) consa-
cré aux problèmes de classement et de
recherche -

③ celles pour lesquelles le temps de recherche est proportionnel à N (méthodes purement séquentielles) - Elles sont comme nous l'avons dit d'un très mauvais rendement en temps mais d'un encombrement optimum - De plus ces méthodes sont obligatoires si le catalogue est sur un support séquentiel. par essence (bande magnétique) - Elles sont donc souvent les seules possibles pour des installations utilisant de petits ordinateurs (~~avec~~ temps d'unité communément peu cher) avec ~~des~~ mémoires secondaires des bandes magnétiques - C'est la méthode du pauvre - ...

1.4.3.3 Récapitulations et autres opérations

Au cours des paragraphes 1.4.3.1 et 1.4.3.2 nous venons d'étudier les opérations suivantes

- $L \rightarrow V$ accès logique (dont nous avons vu qu'elle était en fait une opération $L \rightarrow L$)

• $N \rightarrow V$ a ces plays fin dont (1)
nous avons vu fin elle était en fait
le produit d'une opération $N \rightarrow L$ par
une opération $L \rightarrow V$

• $N \rightarrow L$ (cf $N \rightarrow V$)

• $V \rightarrow L$ que nous avons étudié
à propos de l'opération $N \rightarrow L$ qui a
été est très proche -

Il reste à étudier l'opération

$L \rightarrow N$ et l'opération $V \rightarrow N$.

(Celle dernière peut se décomposer en
 $V \rightarrow L$ puis $L \rightarrow N$.)

l'opération $L \rightarrow N$ ne pose aucun
problème si le nom de la donnée est
stocké dans le lien de celle-ci (c'est
ce qui se passe la plupart du temps même
s'il y a un catalogue) - Si ce n'est
pas le cas il existe toujours soit direct-
ment soit par l'intermédiaire de la table
d'occupation un moyen pour "remonter"

depuis le lieu jusqu'à l'événement
catalogue (donc jusqu'au nom) (98)

1.5 Bibliographie

- Another Look at Data

G. H. MEALY F.J.C.C. 1967 (pp 525-536)

- Data structure and their representation in Storage. M. D'IMPERIO - Annual Review in Automatic Programming Vol 5. Pergamon Press New York - Spring 1969 -

- IBM System/360 Operating System - Supervisor and Data Management Services

- IBM System/360 Operating System - PL/I (F) Language Reference Manual -

- A use of fast and slow memories in list processing languages. J. COHEN I.C.A.C.M. Vol 10. N.2 Feb 67

- Scatter Storage Techniques

R. MORRIS C.A.C.M. Vol 1. N1. Janv. 1968

- Projet SOERATE

J. R. ABRIL, G. BEAUME etc.... Université
de Grenoble - Août 1970

2. Point de vue morphologique

(29)

Dans le chapitre précédent nous avons étudié les structures de données sous un angle existentiel sans préciser ce qui était la "valeur" et le "lien" d'une donnée.

Nous avons schématisé à l'extrême ces 2 espaces en laissant sous entendre que le lien d'une donnée était une boîte et sa valeur, le contenu de cette boîte.

Dans la pratique il en est tout autrement : le lien d'une donnée n'est pas souvent isolable en une seule boîte mais de plus il peut varier dynamiquement au cours de la vie de la donnée.

Dans ce chapitre nous allons étudier la structure des données d'un point de vue morphologique ; c'est à dire analyser les différentes formes que peuvent prendre les ~~données~~ valeurs des données et voir

comment techniquement on peut réaliser ⁽¹⁰⁾
au mieux les liens correspondants -

Plutôt que de commencer par une
étude formelle nous allons d'abord
considérer deux "études de cas" fonda-
mentales : les listes linéaires puis les
arbres - A propos des listes linéaires
nous ferons d'assez longues digressions ^{concernant}
~~sur~~ des piles et des files d'attente -

Nous terminons le chapitre par une
étude plus théorique en nous serons
amenés à comparer la morphologie
des données avec celle des langages
de Programmation.

2.1 Liste linéaire (Introduction)

Une donnée est une liste linéaire si
l'on peut décomposer sa valeur en ~~en~~ une
suite ordonnée d'éléments -

Nous parlerons de liste linéaire par opposition
au terme "liste" : dans une liste linéaire
chaque élément est un tout insécable à

la structure de quel on ne s'intéresse (10)
pas! -

Une liste linéaire est la plus simple
des données à structure de valeurs dynami-
que; elle possède un degré de liberté:
son nombre d'éléments -

On voit donc que en plus, des opéra-
tions que nous avons étudiées au
chapitre précédent (création, suppression,
accès) il vient s'ajouter un certain
nombre de nouvelles opérations spéci-
fiques de la structure de liste linéaire

- les opérations de croissance (extension,
diminution ~~et~~ : ajouter ou enlever des éléments)
- les opérations de sélection

Ces dernières servent à isoler un ou plusieurs
éléments de façon à pouvoir les lire, chan-
ger leur valeur ou les enlever - Nous allons
donc les étudier d'abord -

Le tableau ci-dessous résume un certain

Sélection	$\left\{ \begin{array}{l} \text{du premier} \\ \text{du dernier} \\ \text{du } n^{\text{me}} \\ \text{de n'importe quel} \\ \text{de tous} \end{array} \right\} \text{élément } \underline{\text{conditionné}}$
Condition	$\left\{ \begin{array}{l} \text{rien} \\ \text{condition sur la valeur} \\ \text{condition sur la position} \end{array} \right\}$
Condition sur la position	$\left\{ \begin{array}{l} \text{juste} \\ \text{en } n^{\text{me}} \text{ position} \\ \text{n'importe où} \\ \text{_____} \end{array} \right\} \left\{ \begin{array}{l} \text{avant} \\ \text{après} \end{array} \right\} \text{élément } \underline{\text{sélectionné}}$

Un système général de gestion de liste linéaire devrait permettre de réaliser simplement toutes les opérations sur une liste ^{linéaire} désignée par exemple par son nom.

A noter que l'opération d'ajout d'un élément se fait bien entendu par rapport à la position d'un élément existant.

403

Operation	ET ENTS				Commentaire
	de nom de liste	de rang	de valeur	de	
Créer(l)	avec liste	/	/	X	Le résultat de cette opération est la création d'une nouvelle liste c'est l'opération (boîte + allocation de chaque l_2)
Supprimer(l)	Le nom de la liste pas	/	/	/	Suppression de la liste de nom l (dépassement + deallocation) cf ch. 1
Valeur(l, r)	"	Il y a une liste pour l'élément de rang r on doit avoir $1 \leq r \leq \text{card}(l)$	/	/	Le résultat de cette opération est la valeur de l'élément de rang r de la liste de nom l
Minimum(l, r, v)	"	On doit avoir $1 \leq r \leq \text{card}(l)$ et $r \geq r_2$	Il n'existe pas d'élément	/	Cette opération a pour résultat la valeur minimum des éléments de la liste l ayant pour valeur v et tel que $1 \leq r \leq r_2$
Maximum(l, r, v)	"	On doit avoir $1 \leq r \leq \text{card}(l)$ et $r_1 \leq r_2$	"	/	Cette opération a pour résultat la valeur maximum des éléments de la liste l ayant pour valeur v et tel que $r_1 \leq r \leq r_2$
Insérer(l, r, v)	"	On doit avoir $1 \leq r \leq \text{card}(l)$	Il ne peut être inséré un élément	/	Cette opération remplace la valeur de l'élément de rang r de la liste l par la valeur v
ajout(l, r)	"	On doit avoir $1 \leq r \leq \text{card}(l)$	Il n'existe pas d'élément	X	Cette opération ajoute un nouvel élément à la liste de nom l - cet élément prend la rang r et a une valeur

Opération	Liste				Commentaires
	de nom de liste	de 1 rang	de valeur	de position	
insérer (l, x)	(le nom de la liste) à insérer	on doit avoir $1 \leq x \leq \text{card}(l)$	/	/	Cette opération insère le x -ième de nom l l'élément de rang x . $\text{card}(l) \leftarrow \text{card}(l) + 1$
Card(l)	"	/	/	/	Cette opération a pour résultat le nombre d'éléments de la liste l

Ces ensembles d'opérations représentent un répertoire possible pour un système général de gestion de listes linéaires -

Les opérations les plus fréquents pour l'enregistrement dans la liste sont peut-être les variables pouvant la manière dont les listes sont allouées en mémoire. Mais il est assez rare de ~~trouver~~ trouver un système permettant efficacement toutes ces opérations.

La manière dont nous avons organisé ces opérations est assez typique de la façon dont on devrait travailler en programmation listée.

Si l'on part en effet d'une machine ⁽¹⁰⁵⁾
munie pourvue d'une mémoire centrale,
et d'un certain nombre de mémoires secon-
daires ainsi que d'un répertoire d'instruc-
tions et si l'on implémente à partir de
ces possibilités les 9 opérations ci dessus,
on s'est fabriqué une nouvelle machine
munie d'une mémoire ^{centrale} extensible qui
n'est pas un simple vecteur mais un
ensemble quelconque de liste linéaire.

Ainsi la génération des programmes
(et donc des programmes) qui se servent
de cette nouvelle machine ont l'impression
d'avoir à leur disposition un hardware
plus élaboré que le hardware de base.

Ils peuvent à leur tour construire de
nouvelles opérations etc. . . .

On voit que les erreurs que nous avons
indiquées ici sont l'équivalent pour cette
nouvelle machine des erreurs que l'on
peut faire en utilisant le hardware

de base commune : une division par zéro, ⁽¹⁰⁶⁾
une tentative d'atteindre une adresse
inexistante etc. . . .

Dans la pontifine on traite ces erreurs
d'utilisation du hardware par un dévrou-
lement vers un ~~point~~ programme implanté
en mémoire en une zone constante.

Il est souhaitable que les erreurs
~~testées~~ dans l'emploi des instruc-
tions de la nouvelle machine soient
traitées d'une façon analogue.

A titre d'exemple on peut citer à
l'aide des opérations précédentes l'algo-
rithme : dans (l_1, l_2) qui a
pour résultat (dans un registre r_0) le
^{minimum} rang dans la liste de noms l_1 à
~~partir~~ partir duquel on trouve
une suite d'éléments consécutifs iden-
tiques à ceux de la liste de noms l_2

dans (l_1, l_2)

$r_0 \leftarrow 1$

$\alpha: v_2 \leftarrow \text{valeur}(l_2, 1)$

$r_0 \leftarrow \text{rangmin}(l_1, v_2, r_0, \text{card}(l_1))$

$r_1 \leftarrow r_0 + 1$

$r_2 \leftarrow 2$

$\beta: v_2 \leftarrow \text{valeur}(l_2, r_2)$

$v_1 \leftarrow \text{valeur}(l_1, r_1)$

if $v_1 = v_2$

then $\left[\begin{array}{l} r_1 \leftarrow r_1 + 1 \\ r_2 \leftarrow r_2 + 1 \\ \text{if } r_2 > \text{card}(l_2) \\ \text{then goto fin} \\ \text{else goto } \beta \end{array} \right.$

else $\left[\begin{array}{l} r_0 \leftarrow r_0 + 1 \\ \text{goto } \alpha \end{array} \right.$

fin:

Si l'on sort de l'algorithme avec une erreur (détectée par l'une des opérations de base) c'est que l_2 n'est pas compris dans l_1

Dans certains septimes de tableau de liste linéaire, seuls un certain nombre des opérations du tableau ci dessous sont permises et constituent des cas particuliers fondamentaux

Pile

Une pile est une liste linéaire sur laquelle seules les opérations suivantes sont permises

opérations générales	opération - pile	
créer	créer	créer
supprimer(l)	supprimer(l)	supprimer
enlève(l, card(l))	retire(l)	retire
ajoute(l, card(l)+1)	pousse(l)	ajoute
insérer(l, card(l), v)	met(l, v)	insère
valeur(l, card(l))	valpile(l)	valeur

File d'attente

Une file d'attente est une liste linéaire sur laquelle seules les opérations suivantes sont permises :

Opérations Générales	Opération - File	
Créer	Créer	card & FA
Supprimer (l)	Supprimer (l)	supprimer P & FA
enleve (l, 1)	tirer (l)	enlever l & FA
ajoute (l, card(l) + 1)	pousse (l)	ajouter un l'élément
misajour (l, card(l), v)	met (l, v)	mettre à jour de données
valeur (l, 1)	valfile (l)	demande 1 ^{er}

2.2 Pile

La pile est un des mécanismes les plus largement répandus en programmation système -

Nous allons illustrer l'emploi d'une pile pour l'écriture de programmes réentrants

Un programme réentrant est un
programme qui n'est pas modifié par
sa propre exécution - Que peut-on
 modifier dans l'exécution d'un
 programme ? - On peut modifier
 a) les instructions de l'algorithme

À u début de l'informatique, il ⁽¹¹⁰⁾ était courant de voir certains "artistes" écrire des programmes qui modifiaient leurs instructions en cours d'exécution (aiguillage télécommandé etc. . .)

Le genre d'acrobatie a à peu près (heureusement) complètement disparu car on s'est rendu compte que les programmes ainsi écrits étaient illisibles et très difficiles à mettre au point; & de plus ils ne pouvaient pas être rendus recitants.

(b) Les zones de travail de l'algorithme. Elles - ci sont au nombre de

trois : (b1) la zone où ~~est~~^{sont} disposés les paramètres d'appel de l'algorithme

(b2) la zone où sont disposés les paramètres de résultat de l'algorithme

(b3) la zone où sont disposées les variables de travail proprement dites.

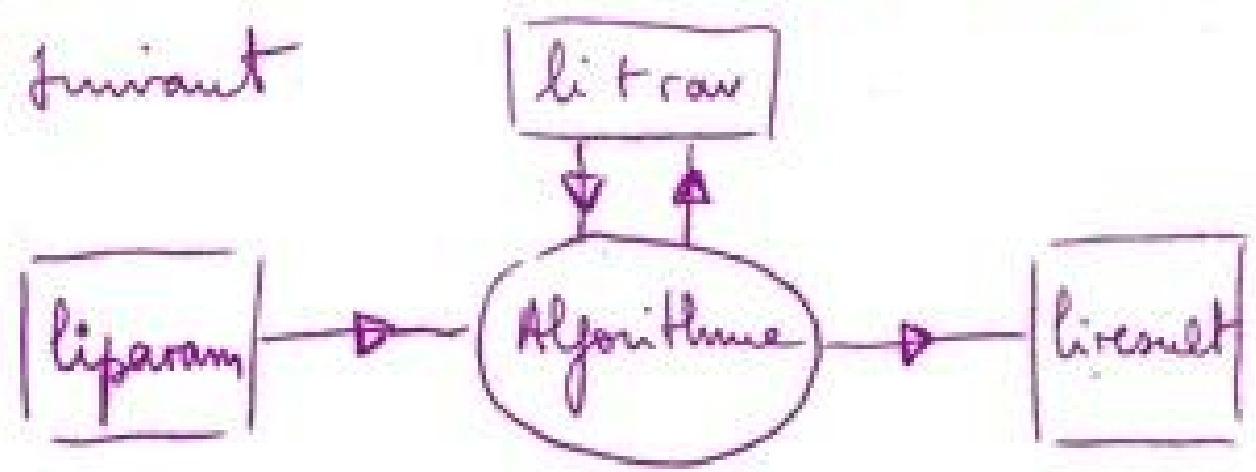
Bien entendu, suivant les cas, ces zones b_1, b_2, b_3 peuvent être inexistantes.

Rendre un programme réentrant c'est donc faire en sorte que les 3 zones b_1, b_2, b_3 ne soient pas ~~des~~ contenues dans l'algorithme lui-même.

En fait ces zones sont des listes, nous les appellerons par la suite :

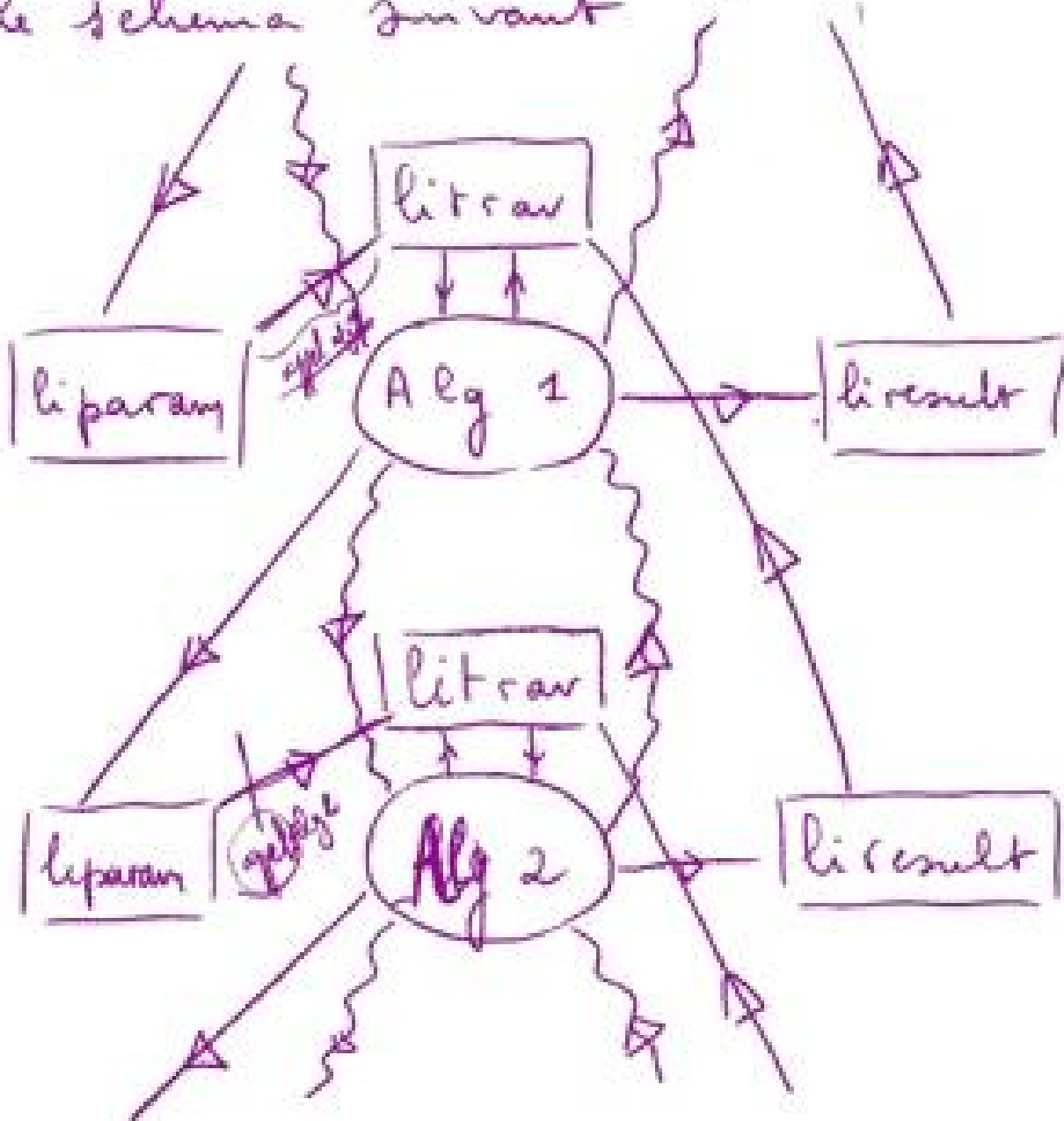
- li param
- li resultat
- li trav

Le schéma d'exécution d'un algorithme réentrant doit donc être le suivant



En fait ce schéma se complique (1.12) car un algorithme peut en appeler un autre avec retour (sous-programme) ou sans retour (passage de contrôle).

Dans le 1^{er} cas on peut dessiner le schéma suivant



Sur ce schéma nous avons supposé (113)
que

a. au début de l'exécution d'un algorithme sa liste de paramètres était transférée dans lienv de façon à libérer la liste liparam pour un autre appel

b. Au cours de l'exécution, l'algorithme travaille avec lienv (en la faisant grandir dynamiquement éventuellement) et il élabore ses résultats dans liresult; s'il desire appeler un autre algorithme il ~~élabore la liste de paramètres~~ élabore la liste de paramètres dans liparam avant de faire l'appel proprement dit

c. à la fin de l'exécution de l'algorithme on doit enlever de lienv l'environnement de son exécution c'est à dire les paramètres

que nous avions transféré (en ②) et ⑪k
les variables de travail que nous y
avons mis (en ④) - Puis l'on
doit mettre dans litrav le résultat
(^{litrav}~~litrav~~) et enfin ~~se~~ redonner le
contrôle à l'algorithme, appelant.

On voit donc que litrav fonctionne
comme une pile -

Dans la pratique litrav est
généralement constituée par les
registres de l'Unité centrale qui
ne sont pas en nombre suffisant
pour constituer une pile de taille
suffisante. On doit donc constituer
une pile indépendante -

Dans les conventions de l'IBM
(IBM) la pile est formée par une
suite de SAVEAREAs doublement
chaînées entre elles; nous verrons sur

ces techniques au § 2.3 (Allocation des listes linéaires)

Nous allons maintenant à l'aide des opérations générales définies au § 2.1 pour les piles définir de nouvelles opérations permettant de réaliser le mécanisme du schéma ci-dessus.

Rappel des opérations sur les piles

$\text{pousse}(l)$: ajoute un ~~nouvel~~ élément à la fin

$\text{retire}(l)$: enlève le dernier élément

$\text{met}(l, v)$: donne la valeur v au dernier élément

$\text{valpile}(l)$: ~~donne~~ a pour résultat la valeur du dernier élément

Très souvent on fait ~~consecutivement~~ ^{successivement} les 2 opérations $\text{pousse}(l)$ et

$\text{met}(l, v)$ - Nous introduisons donc l'opération :

insert (l, v)

pousse (l)

met (l, v)

Nous avons aussi besoin des deux opérations $\text{sauve}(l_1, l_2)$ et $\text{restaure}(l_1, l_2)$ que la première vide la pile l_2 dans la pile l_1 en ajoutant en plus dans l_1 la taille initiale de la pile l_2 . La 2^e est l'opération inverse (où l'on suppose donc que le dernier élément de la pile l_1 contient le nombre d'éléments à insérer dans la pile l_2).
 À la fin de $\text{sauve}(l_1, l_2)$ la pile l_2 est vide -

sauve (l1, l2)

trav ← card (l2)

for i = (1, 1, trav)

do [insert (l1, valpile (l2))
 retire (l2)

insert (l1, trav)

restaune (l_1, l_2)

1 $\text{trav} \leftarrow \text{valpile}(l_1)$
 $\text{retire}(l_2)$

 for $i = (1, 1, \text{trav})$

 do $\left[\begin{array}{l} \text{insert}(l_2, \text{valpile}(l_1)) \\ \text{retire}(l_1) \end{array} \right]$

A l'aide de ces 2 nouvelles opérations nous allons maintenant construire les opérations d'appel et de retour. Dans chaque cas nous considérons 3 variantes

- cas général (li param ou li result)
- cas special 1 (un seul paramètre, un seul résultat)
- cas special 2 (~~un seul~~ pas de paramètre pas de résultat)

La liste l est la pile. Elle est implémentée -

Nous avons donc les opérations suivantes.

appel ret liste (adprog, adret)

saue (l, ltrav)

insert (l, adret)

saue (ltrav, lparam)

goto adprog

l	ltrav
adret	lparam

appel ret (adprog, adret, v)

saue (l, ltrav)

insert (l, adret)

insert (ltrav, v)

insert (ltrav, 1)

goto adprog

adret	ltrav
1	v

appel sans (adprog, adret)

saue (l, ltrav)

insert (l, adret)

~~insert (ltrav, 0)~~

goto adprog

adret	ltrav
	0

A noter que dans les 2 cas spéciaux nous avons mis 1 ou 0 dans ltrav de façon à être cohérent avec le cas général

Adret, l'adresse de retour est stockée (119)
juste derrière les variables de travail
et va permettre de "remonter" à l'algorithme
appelant

Pour les opérations suivantes nous
avons besoin de ~~lister~~ l'opération
suivante qui vide complètement une
pile

nettoie (l)

trav ← card(l)

for i = (1, 1, trav)

do [retire (l)

retour liste

trav ← valpile (l)

retire (l)

nettoie (litrav)

restaure (l, litrav)

saure (litrav, liresult)

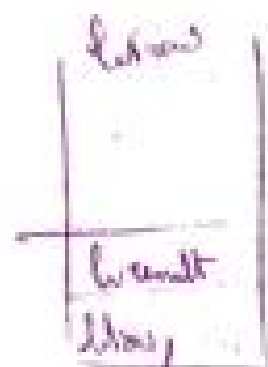
goto trav

val pour déterm
l'adresse de la pile
de liste

litrav ← trav - 1

liresult ← l + 1

liresult ← l + 1



utaur (v)

1 trav $\leftarrow$ valpile (l)

utire (~~l~~)

nettoie (litrav)

restaure (l, litrav) $\left\{ \begin{array}{l} \text{trav} \leftarrow \text{valpile}(l) \\ \text{nettoie}(l) \\ \text{entre } l \text{ - objets sur la} \\ \text{long de trav} \end{array} \right.$

insert (litrav, v)

insert (litrav, 1)

foto trav

utaursans

trav $\leftarrow$ valpile (l)

utire (l)

nettoie (litrav)

restaure (l, litrav)

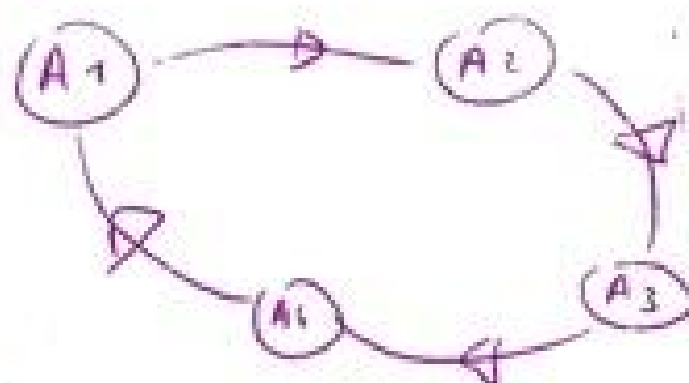
insert (litrav, 0)

foto trav

Le tableau ci dessous donne une
recapitulation générale de toutes
les opérations que nous venons d'in-
troduire

Operation		Comme suit ci-dessous
insert (ℓ, v)		ajoute un nouvel élément à une pile avec initialisation de cet élément à la valeur v
retire (ℓ)		vide complètement une pile
peek (ℓ_1, ℓ_2)		vide la pile ℓ_2 dans la pile ℓ_1 si ℓ appartient à la taille initiale de ℓ_1
peekmax (ℓ_1, ℓ_2)	opération inverse de la précédente	
appelret-liste (adprog, adret)		appel de l'algorithme adprog et retour plus tard dans l'algorithme adret (avec une liste de paramètres dans l'param)
appelret (adprog, adret, v)		même opération avec un seul paramètre v
appelparam (adprog, adret)		même opération pour paramètre
retour liste		retour dans un algorithme avec une liste de résultats dans le résultat
retour (v)		même opération avec un seul un seul v
...		même opération pour valeur

En monoprogrammation on a besoin (ici) de procédures (algorithmes) réentrantes ~~par~~ lorsque l'on a des algorithmes récurifs c'est à dire tels que la série des appels que l'on peut trouver forme une boucle.



On a l'impression que de tels appels vont à l'exécution produire un cycle infini : ceci n'est pas sûr car les appels sont conditionnés (Remarque : A noter que ceci est une condition nécessaire pour qu'une série d'algorithmes ne boucle pas indéfiniment, mais ~~non~~ suffisante (cf ch 0) car elle n'est pas suffisante).

En multiprogrammation (Nous y revenons au § 2.3) on a besoin d'algorithmes réentrant des fois on veut les partager entre plusieurs utilisateurs -

(123)

Notons enfin que en mono programmation on a besoin d'une pile des
fonct. l'on a des appels avec retour
entre algorithmes (car il faut stocker
quelque part les points de retour) - En
fait dans bien des cas la pile
"accompagne" l'exécution des
algorithmes - Un bon exemple
de ceci est l'étude de l'exemple 5
page 12 (Non-reentrant save Area
chaining) de l'exemple 9 page 16
(passing control with return) et de
l'exemple ^(page 18) 13 (establishing a return
code) du document IBM System/360
Operating System - Supervisor and
Data management Services -

~~Après allouer un point de retour dans les exemples d~~

A noter que dans ce dernier cas on
a une pile mais elle est stockée
dans les algorithmes eux même - Si

l'on veut rendre les algorithmes récursifs⁽¹²⁶⁾
on doit fabriquer la pile à l'extérieur
des algorithmes en utilisant les macros
get main et free main (cf exemple
page 12 - Recursible save area chaining)

Nous allons maintenant étudier
2 exemples de programmes récursifs
pour illustrer l'emploi des opérations
que nous venons d'introduire

Exemple 1

Soit l'algorithme suivant
 $\text{fact}(N)$

if $N = 1$ then 1

else $N \times \text{fact}(N-1)$

On le "traduit" comme suit

fact: if valeur(litrac, 1) = 1

then retourne (1)

else appelle fact, α , valeur(litrac, 1) - 1

α : retour (valeur(litrac, 1) $\times$ valeur(litrac, 3))

l'appel de l'algorithme se fait
comme suit

insert ~~l~~ (l, trav, o) → met o de l dans
appel ret (fact, fin, v) → met l dans l
fin: → met o et l dans l

Trace

si l'on fait $v = 4$ on a la trace
suivante

passage	l	l.trav
		l
→ fact	o; fin	(4, 1)
→ fact	o; fin, 5, 4, 3, 2, 1	3, 1
→ fact	o; fin, 4, 1, 2, 3, 1, 2, 1	2, 1
→ fact	o; fin, 4, 1, 2, 3, 1, 2, 3, 1, 2, 1	1, 1
← fact	o; fin, 4, 1, 2, 3, 1, 2, 1	2, 1 → 1, 2
→ d	o; fin, 4, 1, 2, 1	3, 3, 2, 1
← fact	o; fin, 4, 1, 2, 1	4, 1 → 6, 1
→ d	o; fin, 4, 1, 2, 1	24, 1
fin		

Exemple 2: Ctte exemple correspond (26)
à l'analyse descendante et à l'inter-
prétation du langage ~~correspondant~~
défini par la grammaire suivante

$$\begin{aligned} \text{Expr} &::= \text{Terme} \mid \text{Terme} + \text{Expr} \\ \text{Terme} &::= \text{Facteur} \mid \text{Facteur} \times \text{Terme} \\ \text{Facteur} &::= (\text{Expr}) \mid \text{Nombre} \end{aligned}$$

On suppose que la chaîne d'entrée
se trouve dans une file appelée source
l'écriture "Algol" est la suivante

Expr
trav ← Terme
if symbole = '+'
then [nouveau
résultat ← trav + Expr
else [résultat ← trav

Terme
trav ← Facteur
if symbole = 'x'
then [nouveau
résultat ← trav x Terme
else [résultat ← trav

Facteur
if symbole = '('
then [nouveau
trav ← Expr


```

      nouveau
      resultat ← trav
    else [
      trav ← symbole
      nouveau
      resultat ← trav
    ]

```

La "traduction" de cet ensemble d'algorithmes est la suivante :

```

Expr : appelle sans (Travaux, α)
α : if valfile (source) = '+'
    then [
      lire (source)
      appelle sans (Expr, β)
    ]
    retour (valeur (litrav, 1))
β : retour (valeur (litrav, 1) + valeur (litrav, 3))
Travaux : appelle sans (Facteur, γ)
γ : if valfile (source) = 'x'
    then [
      lire (source)
      appelle sans (Travaux, δ)
    ]
    retour (valeur (litrav, 1))
δ : retour (valeur (litrav, 1) x valeur (litrav, 3))
Facteur : if valfile (source) = '('
    then [
      lire (source)
      appelle sans (Expr, ε)
    ]
    insert val (litrav, valfile (source))
    lire (source)
    retour (valeur (litrav, 1))
ε : lire (source)
    retour (valeur (litrav, 1))

```

La trace de cet algorithme pour la chaîne source $(a + b) \times c \perp$ est la suivante :

	valp	Source
→ Expr	{	
→ Terme	{	
→ Facteur	{	
→ Expr	a	
→ Terme	a	
→ Facteur	a	
& Facteur	+	
→ γ	+	
& Terme	+	
→ α	+	
→ Expr	b	
→ Terme	b	
→ Facteur	b	
& Facteur	)	
→ γ	)	
& Terme	)	
→ α	)	
& Expr	)	
→ β	)	
& Expr	)	
→ ε	)	
& Facteur	x	
→ γ	x	
→ Terme	c	
→ Facteur		
& Facteur	↓	
→ γ	↓	
& Terme	↓	
→ δ	↓	
& Terme	↓	
→ α	↓	
& Expr	↓	

$(a \times b) + c. \perp$

2.3 File d'attente

(129)

Comme la pile, la file d'attente est un des mécanismes les plus répandus en programmation système -

Au paragraphe précédent nous avons vu comment à l'aide d'une pile nous pouvions assurer la récursion d'un ~~algorithm~~ algorithme -

Au cours de l'exécution du dernier exemple nous avons rencontré plusieurs fois l'"instruction" Tire (Source) -

Il se peut que la ~~liste~~ liste linéaire "source" ne soit pas en mémoire :

il est nécessaire alors d'effectuer une opération d'entrée pour aller chercher le premier élément de "source"

Pendant ce temps l'exécution doit évidemment s'arrêter : l'exécution attend - Au lieu de stopper l'Unité

Centrale, on peut mettre à profit cette (130)
attente pour lui faire exécuter un
autre ensemble d'algorithmes (par exemple
le même mais pour une autre liste
d'entrée fournie) - Cette nouvelle
exécution va elle-même être bloquée
à son tour et on peut mettre à profit
etc...

On voit donc qu'il va se former
des files d'attente d'exécution bloquée.

Nous allons formaliser un peu ces
phénomènes dans les pages qui suivent
en introduisant d'abord un certain
nombre de définition puis de nouvelles
opérations élémentaires.

Définitions

① C Client : Nous appelons CLIENT
l'exécution d'un certain nombre
d'algorithmes. Un client porte un
nom est associé à une pile. Cette

Pde s'appelle le contexte du client (12)

(b) Point de service Un POINT DE SERVICE est un "endroit" où un client peut être servi. Plusieurs clients peuvent vouloir se faire servir au même point de service. Un client est servi à un point de service par un SERVEUR (voir c). Il peut y avoir plusieurs serveurs par point de service. En un point de service on aura 3 files d'attente

(b1) une file nouveau client (point de service) qui est la file d'attente des nouveaux clients arrivant à ce point de service

(b2) une file vieux client (point-de-service) qui est la file d'attente de clients qui avaient commencé de se faire servir par un serveur du point de service puis qui étaient partis se faire servir à un autre point de service

et qui maintenant reviennent ~~se faire~~ ⁽⁹²⁾
continuer de se faire servir ~~à~~ à
ce point initial -

(b3) Une file serveur libre (point de service)
c'est la file des serveurs de ce
point de service qui n'ont rien à faire
(qui n'ont pas de client à servir)

Un serveur d'un point de service peut
n'avoir rien à faire pour ^(attendre) d'autres

(d) le nombre de clients ^(attendant) à ce
point de service est ~~différent~~ nul +
 $N(\text{point de service}) = 0$

(e) le nombre de clients dont on a
déjà commencé le service à ce point de
service a dépassé un seuil fixé à
l'avance - le seuil s'appelle le
degré de multiprogrammation max.

mm : $DM(\text{point de service})$ c'est
une constante du point de service,
le nombre de client dont on a déjà
commencé le service (et qui n'est
pas fini) s'écrit $D(\text{point de service})$

Les 2 files (a) et (b) (file serveur et file client) sont situées dans une salle d'attente. Pour des raisons qui apparaîtront plus clairement par la suite, cette salle d'attente peut être fermée à des moments ou temps pour certaines périodes très courtes.

On a donc :

Salle d'attente (point de service) $\begin{cases} \rightarrow 0 & \text{si ouverte} \\ \rightarrow 1 & \text{si fermée à des} \end{cases}$

(c) serveur Un serveur est associé par construction à un point de service. Le nombre de serveurs associés à un point de service est une constante qui s'appelle le degré de multi-traitement maximum : DT (point de service)

Un serveur peut être en train de dormir s'il appartient à la file serveur de son point de service.

On a donc :

dort (serveur) $\begin{cases} \rightarrow 1 & \text{si il dort} \\ \rightarrow 0 & \text{si il sert un client} \end{cases}$

Le point de service auquel est associé un serveur s'appelle service (serveur) -

On a donc la règle

si $\text{dort}(\text{serveur}) = 1$ alors

$\text{serveur} \in \text{file_serveurs_libres}(\text{service}(\text{serveur}))$

Pour des raisons qui apparaîtront plus clairement par la suite, chaque serveur veut ~~pouvoir~~^{pouvoir} être tranquille (non dérangé) pendant certaines périodes de temps très courtes.

On a donc

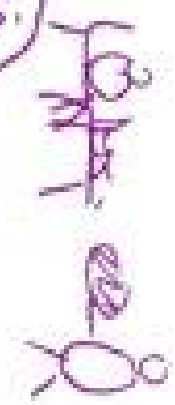
$\text{attention}(\text{serveur}) \begin{cases} 0 & \text{peut être dérangé} \\ 1 & \text{ne dérange sous aucun prétexte} \end{cases}$

Enfin lorsque un serveur ne dort pas c'est qu'il sert un client. Le client est $\text{client}(\text{serveur})$ -

Le tableau ci dessous récapitule la structure de données que nous venons d'introduire -

Élément	Propriété	Règles associées
client		Conteste d'exécution d'algorithme
point de service	Salle d'attente	o/s ouverte ou fermée
	file pour clients	clients nouveaux en attente
	file pour clients	clients anciens en attente
	file serveur libre	serveurs tels que doit (serveur)
	DM	Nombre maximum de clients que l'on peut commencer de servir. Dépend de Multiprogrammation
	DT	Nombre de serveurs. Dépend de multiprogrammation
service	D	Nombre de clients que l'on a commencé de servir
	N	card (file pour clients) + card (file pour clients) + nombre de client servi
serveur	attention :	o/s on peut déranger ou ne pas déranger
	doit	o/s selon client en doit
	client	client du serveur
	service	point de service du serveur

Le schéma ci-dessous concrétise quelques-unes des activités à un point de service



Pe de service

filantroupe

filantroupe

aller d'abord

appel aux actions

Vous devez aller à un autre point de service. Vous rendez obligés dans la file des vieux clients. Ce ne peut être pas moi pour vous servir de nouveau. Mais un collègue - ça marche au mieux. En fait c'est le vrai dernier à venir. Mais il y a des clients à servir en bien si on a pas commencé déjà trop de clients.

On va se faire. Vous allez vous faire à votre précédent point de service - d'abord y en a pas le vrai suppresseur. Quand à moi le vrai dernier à venir. Mais il y a des clients à servir - ça fait un client en moins.

Notre fils a envie d'aller à faire servir ailleurs. C'est en fait, mais on ne s'occupe plus de lui car moi le continue à vous servir. Ça fait un client allé plus dans la système.

appel aux actions

Avant de définir avec précision les opérations de base correspondantes aux "conversations" du schéma ci-dessus, il convient de prévenir l'usage abusif des salles d'attente ~~abstraites~~ : un seul client à la fois n'est autorisé à entrer dans une salle d'attente -

De façon à garantir ceci nous définissons deux opérations :

blq (porte) et dblq (porte)

blq (porte)

K:

cc ← porte
porte ← 1

Test & set

if cc = 1 then goto K

Les deux opérations $cc \leftarrow porte$ et $porte \leftarrow 1$ sont faites de façon ininterrompible c'est à dire qu'un seul serveur ne peut la faire ~~à~~ à la fois -
En 360, ces 2 opérations accolées s'appellent Test & set - Sans de plus amples détails

voir le document IBM System /360 (138)
Operating System : Principle of Opera-
tion page -

d'opération d'blk est triviale.

id blk (porte)

porte ← 0

Nous allons maintenant écrire les
algorithmes des 3 opérations décrites
dans le schéma.

Appel avec retour (serveur, service appelé)

Sortie (serveur, service retour)

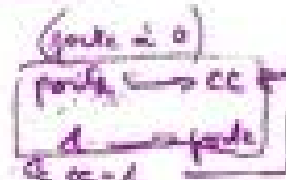
Appel sans retour (serveur, service appelé)

On suppose que dans les 3 cas, un
certain nombre d'opérations de routine
du genre save ou restore ont eu
lieu juste avant ces opérations et
ont rendu le contexte du client
cohérent avec elles -

Appel avec retour (serveur, service appelé)

attention (serveur) ← 1

blq (salle d'attente (service appelé))



insert (file nouveau client (service appelé), client (serveur))

insert
dans la liste
+ 1
à la fin

$N(\text{service appelé}) \leftarrow N(\text{service appelé}) + 1$

if $D(\text{service appelé}) < DM(\text{service appelé})$
and $\text{card}(\text{file serveur libre}(\text{service appelé})) \neq 0$

Then $\text{dort}(\text{val file}(\text{file serveur libre}(\text{service appelé})))$

← 0

$\text{time}(\text{file serveur libre}(\text{service appelé}))$

dlq (salle d'attente (service appelé))

goto 2

dlq (salle d'attente (service appelé)) *pointe suivante*

2: blq (salle d'attente (service (serveur)))

! dans ce cas $N(\text{service}(\text{serveur})) \leftarrow N(\text{service}(\text{serveur})) - 1$

! dans ce cas
on a un
nouveau
client
à insérer
à la fin

if $N(\text{service}(\text{serveur})) = 0$ or $D(\text{service}(\text{serveur})) = DM(\text{service}(\text{serveur}))$

Then $\text{dort}(\text{serveur}) \leftarrow 1$
 $\text{insert}(\text{file serveur libre}(\text{service}(\text{serveur})), \text{serveur})$
 $\text{dlq}(\text{salle d'attente}(\text{service}(\text{serveur})))$

on doit passer

fin

attention (serveur) $\leftarrow 0$
goto scheduler
goto scheduler *un pl pour scheduler
pour les jobs si change
de serveur pour
de minute.*

fonction (serveur, service retour)

attention (serveur) $\leftarrow 1$

if service retour = 0

then [supprimer (client (serveur))
goto β

blq (salledattente (service retour))

insert (filemapclient (service retour), client (serveur))

$N(\text{service retour}) \leftarrow N(\text{service retour}) + 1$

if card (fileserveurlibre (service retour)) $\neq 0$

then [dort (valifile (fileserveurlibre (service retour)))

$\leftarrow 0$
time (fileserveurlibre (service retour))

dblq (salledattente (service retour))

goto β

dblq (salle d'attente (service retour))

β : blq (salledattente (service retour))

$N(\text{service retour}) \leftarrow N(\text{service retour}) - 1$

$D(\text{service retour}) \leftarrow D(\text{service retour}) - 1$

if $N(\text{service}(\text{serveur})) = 0$

in the file server

then $\left[\begin{array}{l} \text{dort}(\text{serveur}) \leftarrow 1 \\ \text{insert}(\text{file}(\text{serveur}), \text{serveur}) \\ \text{dble}(\text{salle}(\text{attente}(\text{service}(\text{serveur}))) \\ \text{attente}(\text{serveur}) \leftarrow 0 \\ \text{goto scheduler} \end{array} \right]$ *end of server*

else goto scheduler

appel sans retour(serveur, service appelle)

client $\leftarrow$ creer

attente(serveur) $\leftarrow 1$

ble(salle(attente(service appelle)))

insert(filenom client(service appelle), client)

$N(\text{service appelle}) \leftarrow N(\text{service appelle}) + 1$

if $D(\text{service appelle}) < DM(\text{service appelle})$

and card(fileservantibn(service appelle)) $\neq 0$

then $\left[\begin{array}{l} \text{dort}(\text{valfile}(\text{fileservantibn}(\text{service appelle}))) \\ \leftarrow 0 \end{array} \right]$

trie(fileservantibn(service appelle))

dble(salle(attente(service appelle)))

goto γ

dble(salle d'attente(service appelle))

γ : ~~attente(serveur)~~ $\leftarrow 0$

fin
 [Scheduler (serveur)]

blq (dort (serveur))

attention (serveur) ← 1

blq (salle d'attente (service (serveur)))

you get an
 email
 (confirmation #)
 for

chdsin : dort (serveur) ← 0

used
 blq 8° = done
 (scheduling)

if card (filierclient (service (serveur))) = 0

then client ← choix (service (serveur),

filerclient (service (serveur)))

enleve (filerclient (service (serveur)),
 client)

client (serveur) ← client

D (service (serveur)) ← D (service (serveur))
 + 1

else client ← choix (service (serveur),

filerclient (service (serveur)))

enleve (filerclient (service (serveur)),
 client)

blq (salle d'attente (service (serveur)))

attention (serveur) ← 0

sert (client (serveur))

Analyse de Scheduler

Nous allons "suivre" un serveur -

Lorsque $\text{dort}(\text{serveur}) = 1$, ce serveur tourne en rond sur l'instruction $\text{dq}(\text{dort}(\text{serveur}))$. Dès que $\text{dort}(\text{serveur})$ passe à la valeur 0 ce serveur sort de la boucle -

- $\text{Dort}(\text{serveur}) \leftarrow 0$ est une instruction qui fait se passer ~~soit~~ ¹ dans appel retour ou appel saisi retour (lorsque $D < D_{\text{ms}}$ et que la file serveur libre n'est pas vide) pour le service ¹appel soit dans sortie pour le service retour.
- Dès que le serveur peut devenir actif il bloque la salle d'attente de son service et choisit un client soit dans file nouveau client soit dans file vieux client (il essaie d'abord d'en

choisir un dans file vieux client) (146)

Le choix est une question de politique
qui peut varier d'un ~~service~~ point de
service à l'autre -

Avant de servir son client le serveur
débloque la salle d'attente -

Si le serveur a pris un client
dans file nouveau client il augmente de 1
le nombre de clients commencés pour
son point de service -

Si au cours du travail, le client doit
s'en aller (avec retour) vers un nouveau
point de service, son serveur le place
dans la file d'attente des nouveaux clients
de ce ~~service~~ point de service et si
les conditions sont remplies il réveille
l'un des serveurs ^(endormis) de ce point de service
Celui-ci peut alors prendre en charge immé-
diatement le client - Quant au
serveur qui vient de perdre son client

il peut éventuellement aller s'endormir (145)
ou ~~se~~ au contraire reprendre immédia-
tement un ~~autre~~ client (scheduler)

À la fin du travail arrive (sortie)
le serveur renvoie éventuellement (sinon
il le tue) son client à son point de
service précédent et il va lui-même
s'endormir ou prendre un autre client.

Les opérations consistant à endormir
ou à réveiller un serveur sont standards
nous pouvons donc les écrire :

endort (serveur) ~~se~~

dort (serveur) $\leftarrow$ $\uparrow$
insert (file d'attente (serveur), serveur);
delq (file d'attente (serveur))

attention (serveur) $\leftarrow$ o

job scheduler

et de même :

reville (point de service; retour)

dort (valfile (file serveur libre (point de service)))



tire (file serveur libre (point de service))

delq (salle de détente (point de service))

foto retour

On peut introduire de nouvelles opérations
commodes :

passer contrôle (serveur, service appelé)

file
à appeler
qui est à l'arrêt
et qui est appelé
à l'instant

[première partie de Appel au retour
(jusqu'en α)]

à l'arrêt
à l'arrêt
à l'arrêt
à l'arrêt

[deuxième partie de sortie
(à partir de β)]

sortie spéciale (serveur, service retour)

à l'arrêt
à l'arrêt
à l'arrêt
à l'arrêt

[première partie de sortie
(jusqu'en β)]

à l'arrêt
à l'arrêt
à l'arrêt
à l'arrêt

[deuxième partie de Appel au retour
(à partir de α)]

→ $V(\text{point de service})$

blq (salledakute (point de service))

$D(\text{point de service}) \leftarrow D(\text{point de service}) - 1$

if $N(\text{point de service}) > 0$ and

$\text{card}(\text{files errant libre}(\text{point de service})) \neq 0$

then $\text{veille}(\text{point de service}, \delta)$

δ :

On reconnaît ici l'implémentation de l'opération V de Dijkstra -

Pour l'opération P il suffit de faire

$\text{Appliquer}(\text{serveur}, \text{point de service})$

le travail fait à ce point de service étant alors sortie spéciale ou plus haut

Pour plus de détail sur les opérations P et V de Dijkstra voir l'article

Co-operating sequential processes du livre

Programming Languages (F. Jennings) (148) (Academic Press 1968)

Tout ce que nous avons dit jusqu'à maintenant supposait que les ~~seus~~ serveurs des pu ils avaient un client ~~pouvait~~ travailler - En fait dans la réalité ce n'est pas vrai. Car pour travailler un serveur ~~et~~ quel qu'il soit a besoin d'un hardware -

De plus il serait vraiment peu économique de mobiliser un hardware pour un serveur qui bouclerait sur la première instruction de scheduler (~~blg~~ (dont (serveur)))

Nous sommes donc obligé d'organiser les serveurs en file d'attente (en attente du hardware)

En fait nous n'avons pas fait un hardware mais ~~une~~ plusieurs classes de hardware qui regroupent ~~et~~ plusieurs hardware

identiques -

(149)

Nous allons donc introduire les nouvelles structures suivantes

(a) classe de hardware

Une classe de hardware est constituée par une file hard (classe de hardware)

qui est une liste des hardware qui constituent la classe - Les serveurs ^(de la classe) qui attendent l'un quelconque des hardware forment une file de garçons (classe de hardware)

Cette file de garçons est dans une salle des gardes dont on veut pouvoir bloquer la porte pendant un temps très court -

En fin il existe un ~~PMU~~ organe chargé de faire fonctionner cette classe de hardware c'est le Regulateur (classe de hardware)

Pour travailler il a besoin d'un indicateur appelé tourner rond (classe de hardware)

(b) hardware

Chaque hardware lorsqu'il travaille est

emporté par un serveur fini est ⑪
emporteur (hardware) et de plus
il a un compteur d'adresse (hardware) -

⑫ serveur On ajoute à la structure
de serveur que nous avons déjà consti-
tuée les éléments suivants :

adr(serveur) fini contient un point
de reprise pour ce serveur (une adresse)
et hard(serveur) fini indique le hardware
emporté par ce serveur - On a :

$\text{emporteur}(\text{hard}(\text{serveur})) = \text{serveur}$

$\text{hard}(\text{emporteur}(\text{hardware})) = \text{hardware}$

⑬ Point de service On ajoute à la
structure de point de service que nous avons
déjà constitué l'élément

classe(point de service) fini indique
la classe de hardware dont se servent les
serveurs de ce point de service -

Les différents programmes que nous
avons écrits sont inchangés sauf
en doit et reville fini de nouveau :

endort (serveur)

dort (serveur) $\leftarrow 1$

insert (fileserveurlibre (service (serveur)), serveur)

dbq (salledattente (service (serveur)))

attention (serveur) $\leftarrow 0$

dbq (tournemond (classe (service (serveur))))

goto scheduler

Reveille (pointdeservice, retour)

garçon $\leftarrow$ valfile (fileserveurlibre (pointdeservice))

dort (garçon) $\leftarrow 0$ *veille garçon*

trie (fileserveurlibre (pointdeservice)) *align*

dbq (salledattente (service (serveur)))

attention (serveur) $\leftarrow 0$ *pas de pénétration*

blq (salledegarde (classe (pointdeservice)))

ajoute (filgarçon (classe (pointdeservice)), garçon)

dbq (salledegarde (classe (pointdeservice)))

dbq (tournemond (classe (pointdeservice)))

attention (serveur) $\leftarrow 1$ *ne pas être la*

goto retour *danger*

Le Regulateur est alors

(152)

Regulateur (classe de hardware)

blq (tournevis (classe de hardware))

blq (salle des fardes (classe de hardware))

meilleux garçon ← choix serveur (fil garçon (classe de hardware))

moins bon hardware ← choix hardware (fil hard (classe de hardware))

if digne (meilleux garçon, moins bon hardware)

then sacifié ← emprunteur (moins bon hardware)

blq (attention (sacifié))

STOP (moins bon hardware)

dblq (attention (sacifié))

adr (sacifié) ← compteur d'adresse (moins bon hardware)

compteur d'adresse (moins bon hardware) ←

adr (meilleur garçon)
emprunteur (moins bon hardware) ← meilleur garçon

enleve (fil garçon (classe de hardware), meilleur garçon)

ajoute (fil garçon (classe de hardware), sacifié)

GETTABLE (moins bon hardware)

dblq (salle des fardes (classe de hardware))

goto R

un seul réglage (tournevis)

avoir de
- mieux
- serveur
- on stop

change

affectation
de la file

lancer

stop
le compt
de la file

le compteur d'adresse
est décalé de 1

Analyse de Régulateur (classe de hardware)

Nous allons suivre un hardware -
 lorsque Régulateur est débloqué
 (tournée ← o sur intervention
 de endort ou reveille), il bloque
 la table des fardes et jace à deux
 algorithmes qui définissent sa politique
 il choisit parmi [les serveurs qui attendent
 un hardware dans file de passage celui
 qui lui paraît le meilleur] (meilleur gap
 et il choisit parmi tous les hardware
 celui qui lui semble le moins bon
 (par exemple parce qu'il travaille avec
 un serveur d'un point de service peu
 important ou pour tout autre raison
 dépendant de sa politique) -

Il regarde ensuite si le meilleur gap
 est digne du moins bon hardware (c'est
 encore une question de politique) - Dans
 l'affirmative, il va sacrifier l'importateur

de moins bon hardware, c'est à dire (dit)
qu'il va ~~se remettre~~ ^{le remettre} dans file de priorité -
Pour cela il arrête le moins bon hardware.

lui change son serveur et le redémarre -
Enfin Regulateur débloque la salle
des gardes puis boucle sur R -

Le Regulateur comme nous l'avons
dit peut être (indirectement) activé
soit par endort soit par réveil.

Lorsqu'il est activé par endort il
est évident que le régulateur aura
Tendance à choisir comme moins
bon hardware celui qui justement a
pour ~~serveur~~ impromptu le serveur que
l'on vient d'endormir - c'est pourquoi
on ne va pas tourner sur ~~Ally~~ (dort
(serveur)) dans Scheduler -

Lorsqu'il est activé par réveil, le
régulateur remet éventuellement en
cause l'affectation des hardware aux
serveurs en tenant compte du nouveau
serveur qui vient d'entrer dans file de priorité.

On remarque que dans le Régulateur⁽¹⁵⁾ on ne stoppe le Hardware sacrifié que lorsque l'opération blq (attention (sacrifié)) le permet. Ceci est nécessaire afin d'éviter que on ne mette dans filed'attente un serveur qui a bloqué une salle d'attente (point de service) car alors ce point de service serait inutilisable jusqu'à ce que le serveur sacrifié ~~se~~ repointe de nouveau un hardware -

Il faut remarquer que ce blocage ~~sur~~ sur attention (sacrifié) ainsi que le blocage sur salle des fardes n'entraîne pas d'événement fatal entre le Régulateur et un hardware exécutant Revil car dans cette dernière opération on a pris soin de faire attention (serveur) ~~à~~ o (à qui débloque éventuellement le Régulateur) juste avant de faire blq (salle des fardes) -

Il faut noter la très grande analogie ⁴⁵⁶
 qui existe entre les éléments du tableau
 ci dessous

Point de service	classe de hardware
serveur	hardware
client	serveur
demandeur et file d'attente	file de parq. 57
de serveur libre + autres serveurs du point de service	file hard
appl	recuil
sortie	endort
salle d'attente	salle des fardes
Scheduler	Regulateur
choix client	choix serveur

Nous allons enfin introduire une dernière
 opération qui permet à un serveur
 de changer de classe de hardware

change (serveur, ~~hardware~~^{classe 1}, classe 2, adresse)

(152)

imposeurs (hard (serveur)) ← 0

blq (salle des gardes (classe 2))

ajoute (file de garçons (classe 2), serveur)

adr (serveur) ← adresse

dblq (salle des gardes (classe 2))

dblq (tourne rond (classe 2))

dblq (tourne rond (classe 1))

goto 0

0 : goto 0

On suppose que serveur travaille
avec un hardware de 1 : classe 2 et
que pour une raison ^(quelconque) il ne puisse plus
continuer à travailler avec lui

(par exemple ce hardware n'est
pas assez puissant : il n'a plus
le repertoire d'instruction suffisant
pour que serveur puisse continuer à

travailler de façon satisfaisante (158)
pour son client -) . Le serveur va
alors changer de classe de hardware
pour emporter un exemplaire de la
classe 2 - Les opérations de change-
ments terminées les deux régulateurs
sont invoqués - Le Hardware de la
classe 1 qui a perdu son serveur
n'a alors plus rien d'autre à faire
qu'à boucler - On se doute bien
que le Régulateur de ~~la classe~~ ^{sa classe} ~~pourrait~~
va le désigner ~~comme~~ tout de
suite comme moins bon hardware
et qu'il va rapidement faire un
travail plus profitable à la commu-
nauté des serveurs et des clients -
~~donc~~ Nous en fin qu'il peut être possible
à un organe extérieur aux hardware
et aux Régulateurs de faire de temps
en temps l'opération ^(toujours) ~~obligatoire~~ de ~~classe de hardware~~
pour l'une quelconque des classes de hardware.

Le Régulateur ainsi délégué peut remettre en cause les affectations de ses ~~processus~~ hardwares aux serveurs. Lorsque l'organe externe lui exécute ces délégations pour une certaine classe de hardware est une horloge travaillant à intervalle régulier on dit que les hardwares de cette classe travaillent en time slicing (tranche de temps)

Dans la pratique des systèmes existant actuellement on trouve généralement 2 classes de hardwares :

- les Unités Centrales
- les Canaux

a) Unité Centrale

Nous allons d'abord considérer le cas d'une installation (3^e génération) n'ayant qu'une seule Unité Centrale -
Le problème est compliqué par le fait

- que une Unité Centrale cache en (160) fait 2 hardwares distincts (ou plus) :
 - l'Unité Centrale en mode normal
 - l'Unité Centrale en mode maître
- que les Régulateurs de chacune de ces 2 classes de hardware (mode normal, mode maître) s'exécutent sur l'Unité Centrale elle-même et non dans un organe incluant d'autre comme nous l'avons décrit dans le système "théorique".

L'opération :

change (servant, mode normal, mode maître, adresse)

s'appelle un appel superviseur (SVC)

L'opération :

change (servant, mode maître, mode normal, adresse, ^{d'ancien})

s'appelle un rechargement d'état programme ~~PROGRAM~~

En résumé on a ⁴ ~~théoriques~~ organes distincts

puissent partager l'usage de l'Unité Centrale

- le Régulateur (mode-maitre)
- le Régulateur (mode-normal)
- l'Uniq. hard (mode-maitre)
- l'Uniq. hard (mode-normal)

Bien entendu ces organes sont en compétition pour utiliser l'Unité Centrale : il existe donc une hiérarchie entre ces organes qui est la suivante

- ① Régulateur (mode-maitre)
- ② hard (mode-maitre)
- ③ Régulateur (mode-normal)
- ④ hard (mode-normal)

On s'aperçoit que l'on a ici une compétition assez arbitraire qui tend d'ailleurs à disparaître

Dans le cas d'un système avec plusieurs Unités Centrales il n'est plus possible de faire exécuter les Régulateurs sur une

Unité Centrale (laquelle ?) - On (162)
voit alors apparaître un ~~bandeau~~
organe indépendant qui est le
Régulateur - C'est le cas par exemple
du 360/67 biprocesseur ou de l'IRIS-86
(CII)

⑥ Canal

L'"histoire" du canal depuis la 1^{ère}
génération jusqu'à nos jours ... serait
à rapprocher de l'histoire d'un organe
dans une espèce vivante sous la pression
de l'évolution sélective -

On peut voir comment le canal
très ~~complexe~~ simplifié et totalement
incorporé à l'Unité Centrale lors de
la 1^{ère} génération a toujours tendu
à prendre de plus en plus son indé-
pendance -

Si nous observons où en est l'évolution
du canal dans un système 3^e génération

classique (1344/360) nous voyons
qu'il est un organe hybride qui
n'a pas encore réussi à se séparer
totalement de l'Unité Centrale et
qui doit encore s'en servir :-

Autrement dit nous ne pouvons pas
~~aller~~ (malheureusement) identifier un
canal à un point de service c'est
pourquoi nous sommes obligé de prati-
quer des changements de hardware
comme serveur

change (mode maître, canal, adresse)

Ceci s'appelle finalement un

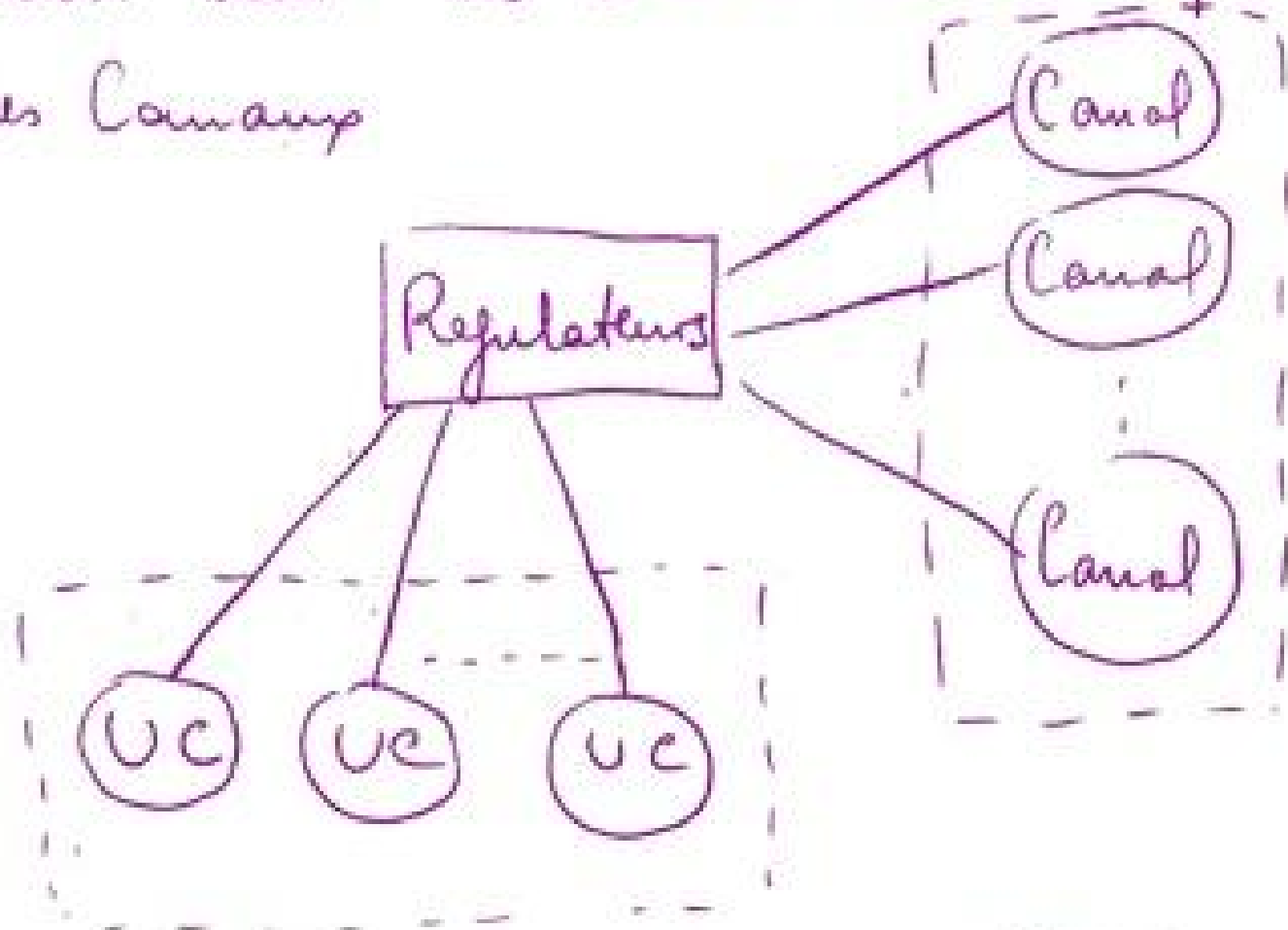
START I/O (SIO)

et comme

change (serveur, canal, mode maître, adresse)

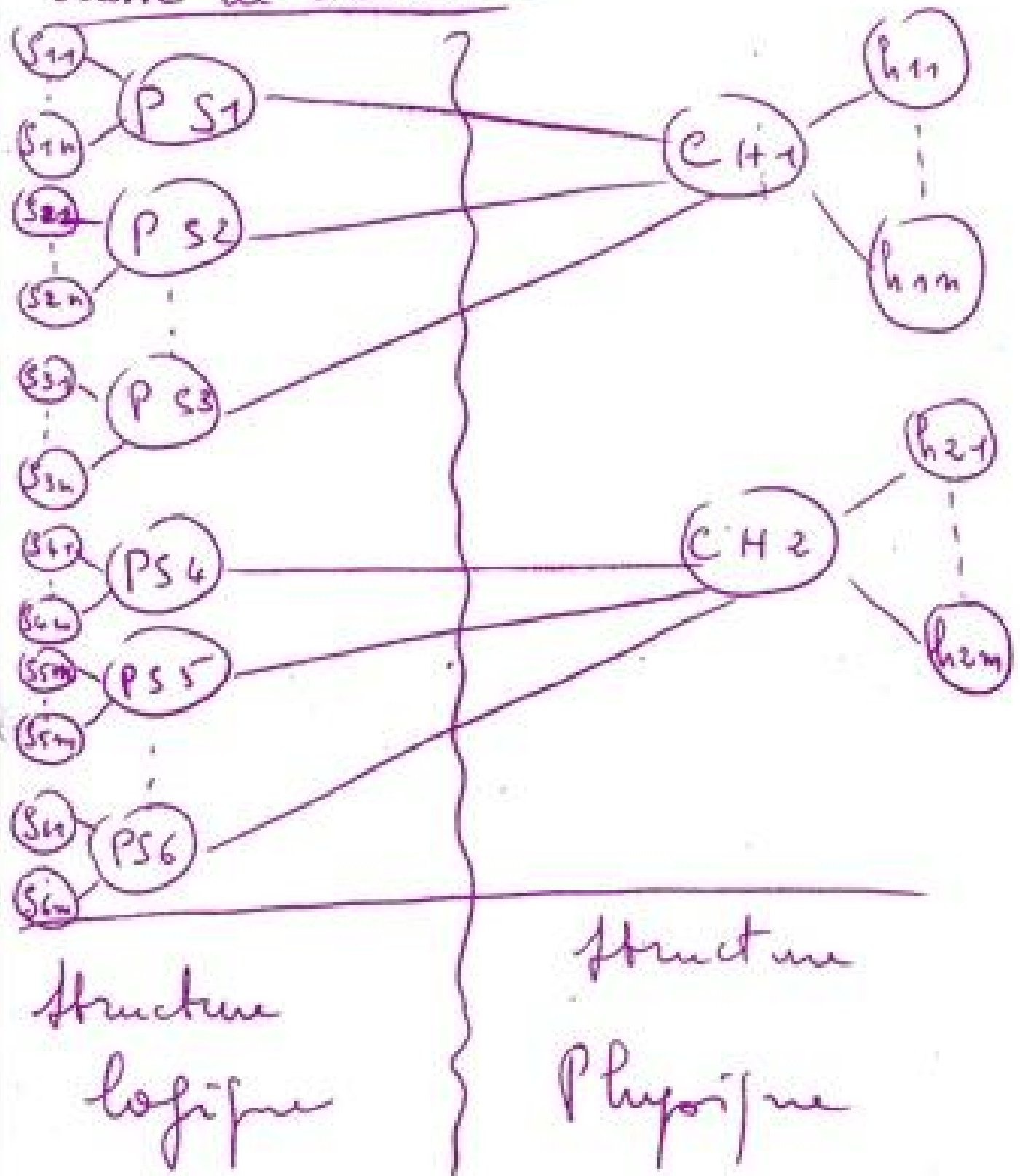
Ceci s'appelle finalement une interrup-
tion Canal

Mais l'évolution qui va dans le (966)
 sens d'une identification des ~~points~~
 Canaux à ~~plus~~^{un} point de service est irréver-
 sible en même temps que les travaux
de Régulation se regroupent tous dans
un même organe chargé de réguler
 aussi bien les Unités Centrales que
 les Canaux



Avec la tendance vers ce schéma
 disparaît ~~donc~~ de plus en plus le besoin des
 opérations change (server, h1, h2, adress).
 on tend donc à avoir des systèmes

fini fonctionnement sans interruption (165)
 Dans de tels systèmes à un point de
 service il ne correspond pas une seule
classe de hardware



le nombre de serveurs d'un point de service (degré de multitraitement) est égal au nombre de hardware de la classe unique correspondant à ce point de service (Voir schéma : m, n)

Lorsque un client (on appelle parfois dans la littérature un client un processus)

passé d'un point de service PS_i à un point de service PS_j ou effectue éventuellement l'opération réveille (PS_j , retour)

(erreur de réveil, ^{le serveur qui a le client en file att de PS_i})
(Voir cette opération page 151) - On voit qu'au cours de cette opération on ~~ajoute~~ place un serveur endormi dans la file d'attente de la classe de hardware de PS_j (opération soulagée puis on débloque le régulateur de cette classe de hardware. Si le serveur réveillé est digne d'emprunter un hardware de la classe i P va effectuer

Scheduler (puisqu'il dormait il 167)
pseudo-tournant sur blq (dont (server))
se debloquer (puisqu'il est réveillé)
et éventuellement (choisissent dans
Scheduler page 142) choisir ce
client - (dans durée fixe chois + tard)

On voit donc que le changement
de hardware pour un ~~processus~~ ^{client}
se fait "en douceur" à l'occasion d'un
changement de serveur et non pas une
opération change beaucoup plus brutale -
à noter que dans ce cas on peut
appeler un serveur un hardware
(ou processeur) virtuel -

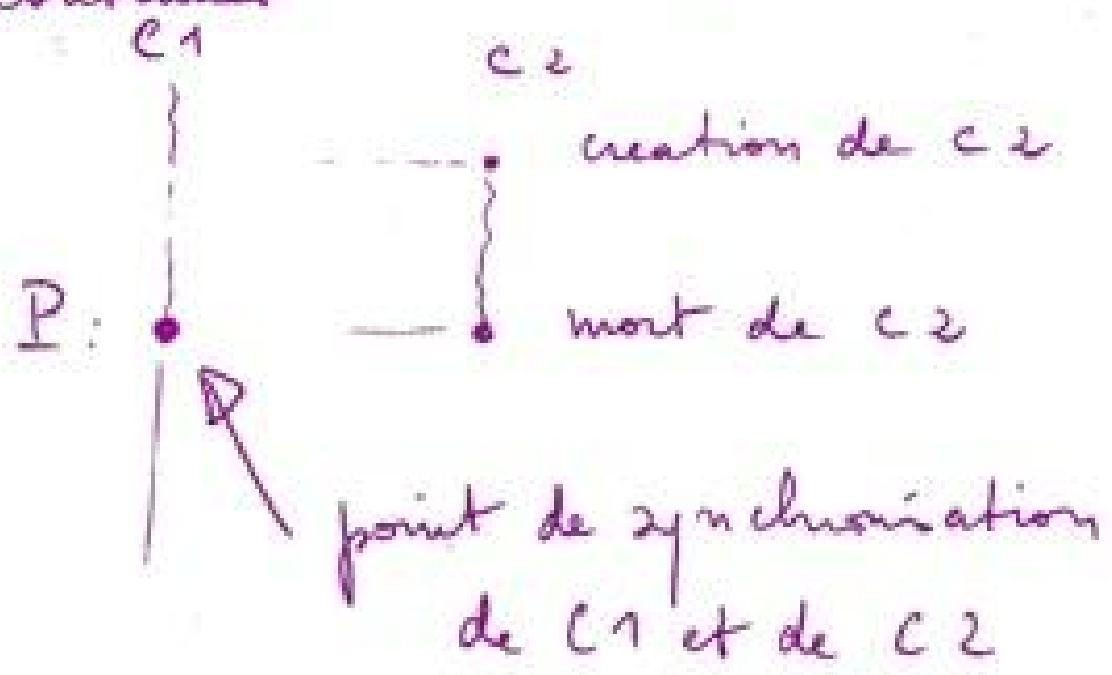
Problème de synchronisation

Nous allons enfin considérer un dernier
problème : celui de la synchronisation
de 2 clients


autre client 192 193 194 195

Supposons qu'un client C_1 ~~soit~~ soit
servi à un point de service PS_1 et qu'il
crée un ~~autre~~ autre client C_2 (comme
dans un appel sans retour) pour être servi
à un 2^e point de service PS_2 -

On veut alors pouvoir synchroniser les
deux clients C_1 et C_2 c'est à dire
qu'une fois que C_2 est créé on veut
que C_1 puisse continuer jusqu'à
un certain point P à partir duquel il
est nécessaire pour lui que C_2 soit
mort (qu'il ~~soit~~ ait fait sortir) pour
continuer



Deux cas sont alors possible

(189)

(a) C1 arrive au point de synchronisation alors que C2 est déjà mort. C1 peut donc continuer sans plus attendre.

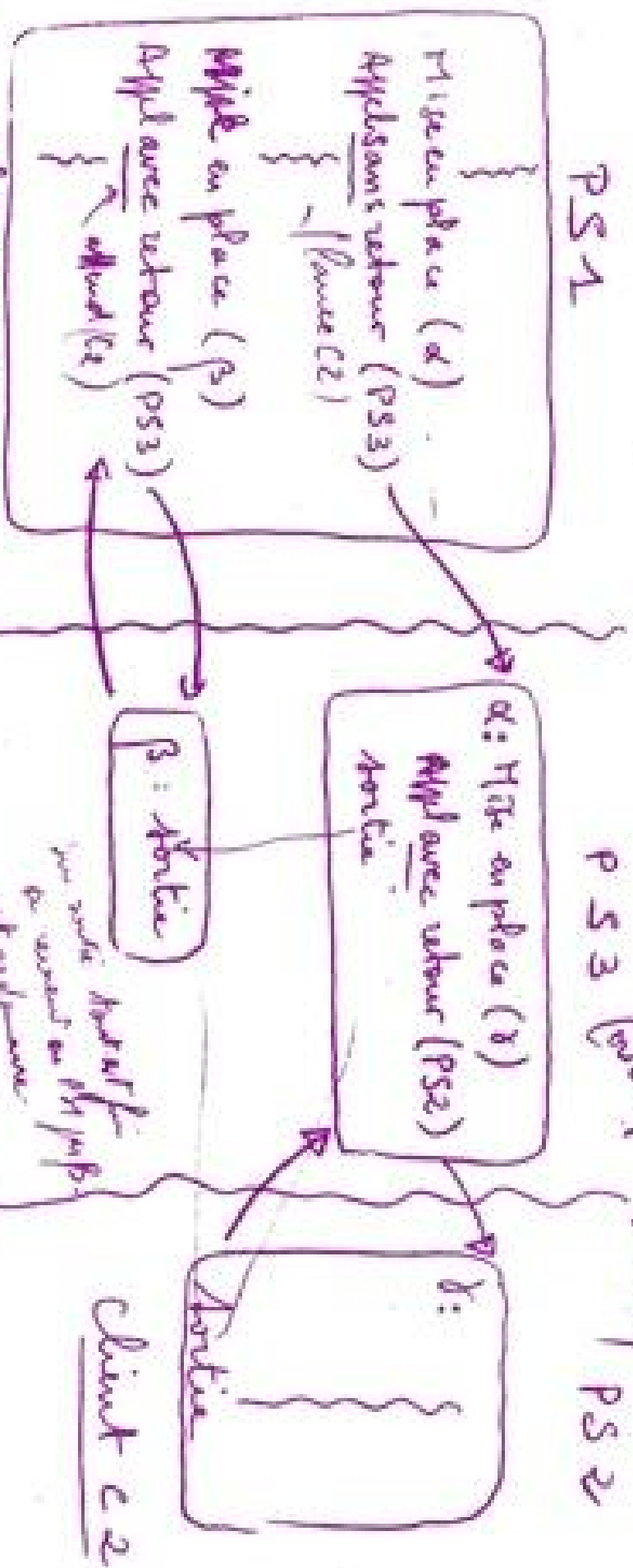
(b) C1 arrive au point de synchronisation alors que C2 n'est pas encore mort - C1 doit attendre.

Dans le cas (b), l'attente de C1 doit avoir lieu comme ~~tous~~ les phénomènes d'attente dans une file, or toutes les files d'attente de clients ~~existent~~ ^{n'existent} ~~seulement~~ que par rapport à des points de service.

On voit donc la nécessité ~~de~~ ~~se~~ d'avoir un point de service supplémentaire PS3 qui sert à synchroniser C1 et C2.

Ce point de service PS3 doit avoir un degré de multiprogrammation maximum égal à 1 : $DM(PS3) = 1$

⑩ du "programme" puis il faut écrire ce qui suit :



Client C1 l'opération Place en place (adresse) consiste à insérer dans la pile du client l'adresse d'un programme - (cf § 2.1)

On voit que cet exemple peut être vu comme une pile de service PS3 ne dépendant pas du client C1 et non

du programme exécuté par le serveur (171)
de C1 -

Ceci amène deux opérations supplémentaires :

- Créer ps

- Supprimer ps (point de service)

qui permettent de créer et supprimer
dynamiquement des points de services -

2.4 Allocation des listes linéaires (1)

Nous allons décrire différentes méthodes d'allocation de listes linéaires ~~en~~ en faisant différentes hypothèses quant à la structure de la mémoire :

Nous supposons d'abord que nous avons une mémoire homogène (un seul algorithme de sélection) adressable au moyen d'atomes de mémoire que nous appellerons des mots (sans préciser leur taille) - Nous supposons donc que les listes que nous voulons allouer sont des suites de mots : un élément de la liste (élément log) correspond à un atome de mémoire (élément physique)

Dans la 2^e hypothèse nous 173
supposons que la mémoire n'est plus
homogène elle comporte 2
niveaux ; ~~le~~ premier est analogue
au précédent et le second plus
grand n'est pas directement interpré-
table par le hardware mais
doit au préalable transférer ^{partie de} son
contenu dans le premier.

A noter que ici n'est pas toujours
le cas dans les hiérarchies de
mémoires : les "extended core memo-
ries" ~~elles~~ constituent des mémoires
d'un niveau inférieur à celui des
mémoires core et pourtant leur
contenu est directement interpré-
table par le hardware - le système
time-sharing ACME de l'hôpital

universitaire de Stanford utilise (174)
une telle mémoire étendue en finc
de mémoire secondaire rapide - (8 p. sec
de temps d'accès)

2.4.1 lien mémoire homogène

Nous allons étudier 7 méthodes
différentes et les comparer sur 6
critères différents qui sont les suivants

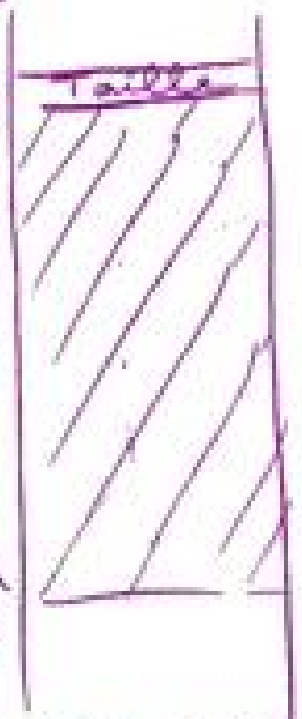
- facilité d'allocation
- facilité de deallocation
- indexation / Recherche
- ajout ^{d'éléments} au milieu de la liste
- place occupée
- possibilité d'extension en mémoire

secondaire (étudier plus en détail au
paragraphe ^{2.4.4} ~~suivant~~)

① Méthode purement séquentielle

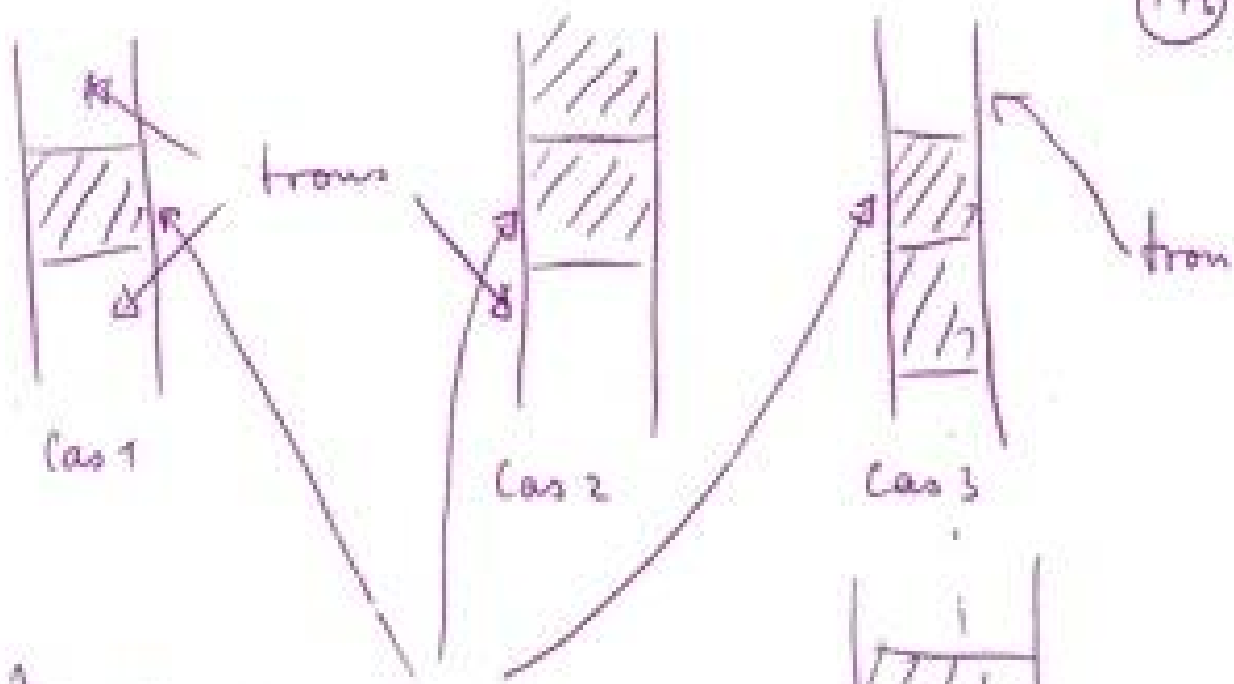
La liste est allouée séquentiellement
en mémoire

L'allocation ou la croissance
font des opérations éventuel-
lement complexes car il est
nécessaire de trouver un tron-
çon de taille suffisante ou alors



de faire des opérations de compactage
(très coûteuse) entre les différentes listes
linéaires de façon à concaténer entre
eux les petits trons pour en former
de plus gros -

La deallocation par contre est simple :
on doit seulement vérifier que la
deallocation d'une liste rentre dans
l'un des 4 cas suivants et reformater
l'espace en conséquence :



liste à supprimer

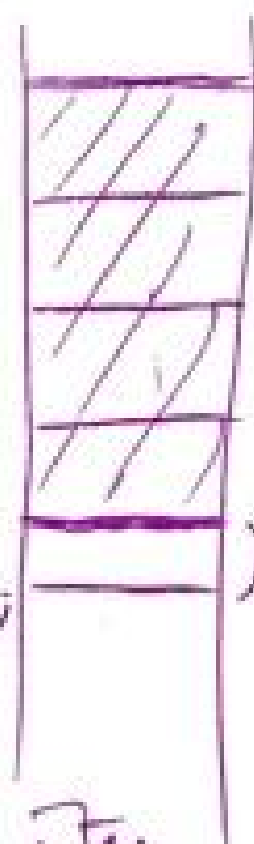
L'indirection est
indépendamment immédiate
ou contre l'aport ou
la suppression d'un

élément au milieu peut être
assez complexe (éventuellement obligation
de compactage)

La place occupée est minimale et
les possibilités d'extension en mémoire
secondaire sont complexes

② Méthode purement séquentielle (177) avec allocation par groupes de mots de taille fixe

Cette méthode est une
simple variante de la
méthode précédente
de façon à améliorer
la croissance éventuelle;



) perte

Ceci se paie par une
perte de place qui peut être
plus ou moins importante suivant
le rapport du reste de la division
de la taille de la liste par
la taille du pavé, à la taille
du pavé $\text{mod}(t, p) / p$.

3) Allocation par groupe de mots de taille variable

(178)

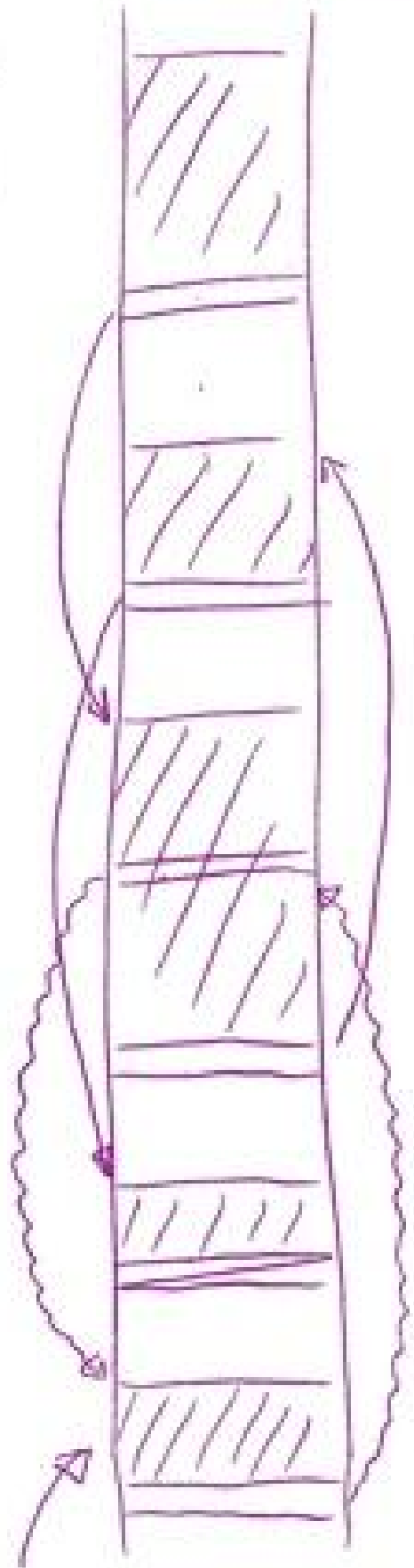
De façon à améliorer les ^{problèmes} ~~difficultés~~ d'allocation des méthodes précédentes on peut fractionner une liste en différentes parties séquentielles de taille variable, chaque partie étant reliée à la suivante par un pointeur -

Cette méthode se caractérise par le fait que l'on peut de temps en temps réorganiser les listes de façon à les rendre ponctuellement séquentielles -

Les caractéristiques d'une telle méthode sont donc variables.

Son principal défaut consiste en l'impossibilité d'indexer en général

par contre il est facile
 d'ajouter un (ou mieux
 plusieurs) éléments au
 milieu d'une liste
 organisée de cette façon :
 comme s'indique le
 schéma ci contre, on
 tronçonne ^(en 2) la partie
 concernée et l'on "plante"
 un pointeur (avec copie
 du mot écrasé par le
 pointeur dans la partie
 ajoutée) vers le nouveau
 groupe fin l'on chaîne
 à la 2^e partie ~~de~~ de
 l'ensemble coupé -



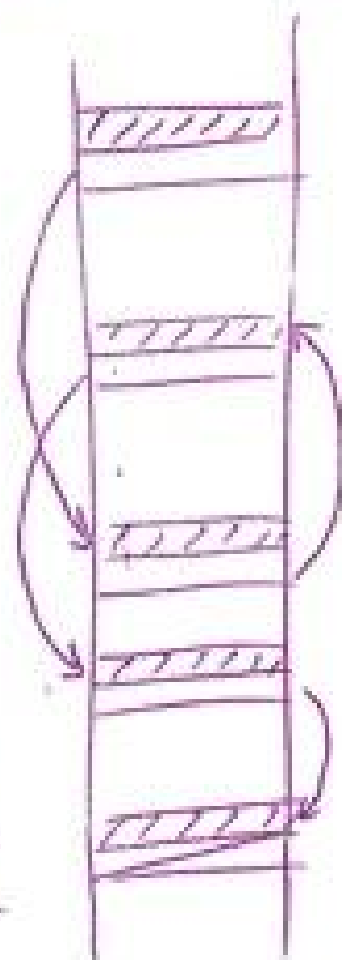
groupes de mots fin
 l'on ajoute "au milieu"

④ Méthode purement chaînée

(190)

Très pratique pour ajouter des éléments au milieu, cette méthode interdit toute indexation et prend beaucoup de place (50% en pointeurs)

⑤ Méthode chaînée par éléments de plus d'un mot de taille fixe



C'est une variante de la méthode précédente dans la quelle on chaîne entre eux non plus chaque mot $\neq$ mais des groupes de taille constante de mots -

Prend moins de place que la méthode précédente et permet une certaine indexation -

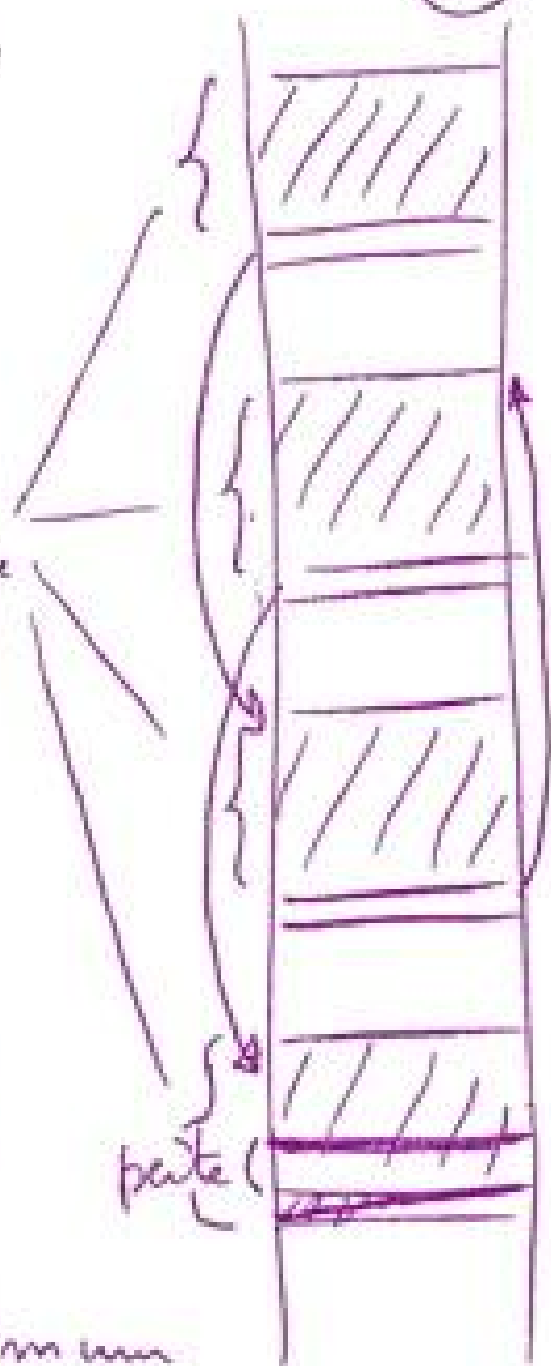
si p_i est la taille
de chaque groupe
• t est la taille
de la liste

la perte "perdue" est :

$$\left[\frac{t}{p} \right] + \frac{p}{2}$$

↑
pour le pointeur

↑
valeur moyenne de la
perte en fin de liste



Il existe donc un optimum
de la taille p approximativement
pour : $p = \sqrt{2t}$

Avec cette méthode on peut commencer
à pratiquer une extension en mémoire
secondaire au niveau de chaque groupe

⑥ Méthode du dictionnaire (pagination)

(182)

Cette méthode apparaît comme
une extension
de la méthode
précédente :

on lie de chaînes
entre elle les différentes
pages comme dans

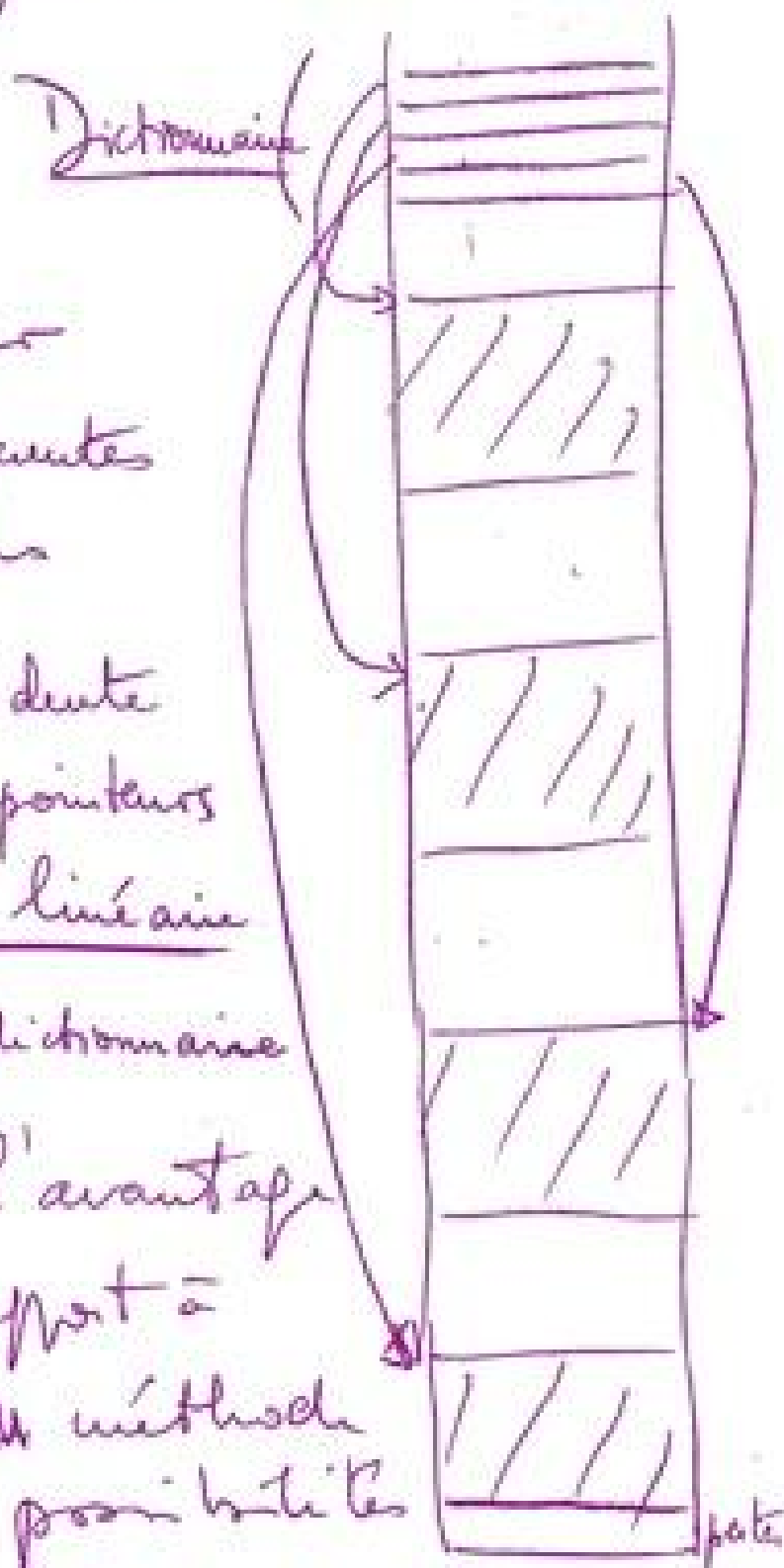
la méthode précédente
on regroupe ces pointeurs
dans une liste linéaire

généralisée : le dictionnaire

Elle présente l'avantage
maigre par rapport à

la méthode précédente
elle donne des possibilités

d'indexation (en fait double indexation)



Les possibilités peuvent d'ailleurs (183)
être facilement implémentées par
hardware -

Par contre on retombe sur une nouvelle
série de difficultés du fait que le
dictionnaire doit être une liste séquen-
cielle (Méthode ①) . On peut donc
être amené pour conserver les possibi-
lités d'indexation du dictionnaire
à le paginer lui-même et donc
à envisager un dictionnaire du diction-
naire etc.

Cette méthode peut être assez contentueu-
se place par contre son extension
à la mémoire secondaire est simple

2) Méthode du Système SOCRATE

Cette méthode vise à avoir les
avantages de la précédente (indexation)

sans avoir les inconvénients de la (196)
méthode 1 (risque d'avoir à faire
des compactage) ni les inconvénients
de la méthode précédente (allocation
d'une pagination de dictionnaire)

On suppose que l'espace mémoire est
extrêmement grand (par exemple 2^{31}
mots $\gg 2$ milliards de mots ----) -
Les diverses listes sont allouées statique
ment dans cet espace - leur nombre
et leur taille sont donc fixes (main
ten constants) - Cet espace s'appelle
l'espace virtuel - Le véritable espace
mémoire s'appelle l'espace réel -

La fonction qui permet de passer
de l'adresse virtuelle d'un mot à
son adresse réelle est une fonction
de repliement puis de projection de

l'espace virtuel sur l'espace réel

(185)

comme θ définit le espace

hémi ci contour -

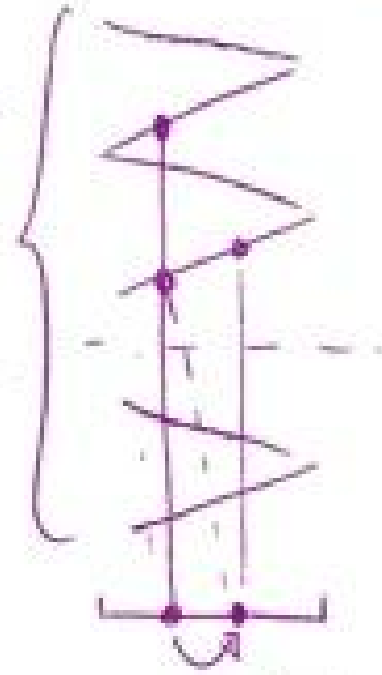
cette fonction

entraîne bien entendu reflé

plusieurs mots virtuels

à avoir la même

projection réelle -



espace réel

Pour résoudre ce problème on cherche à placer dans l'espace réel les

mots qui sont ~~associés~~ issus

de mots virtuels ayant la même

projection - De façon à optimiser

on peut comme dans la méthode 3

(par rapport à la méthode 2) grouper

plusieurs mots virtuels consécutifs

~~de~~ de taille fixe de façon

à avoir un pointeur de chargement

non plus pour un mot mais pour (186)
plusieurs - Appelons page ce groupe
de mots -

L'espace réel est donc formé
d'une suite de pages qui peuvent
être dans l'un des états suivants

libre : cette page réelle ne
correspond à aucune page virtuelle

occupée et bien placée : cette page
réelle correspond à une page virtuelle
dont elle est la projection.

occupée et mal placée : cette page
~~virtuelle~~^{réelle} correspond à une page virtuelle
mais sa position n'est pas la projec-
tion de cette page virtuelle : on l'appelle
un squatter -

Lorsque l'on lit un mot d'une ~~page~~
page virtuelle trois cas peuvent se
présenter :

(187)
- la page projection est libre :
la page virtuelle n'existe pas

- la page projection est occupée
et bien placée : on lit directement le
mot

- la page projection est occupée
et mal placée : on se déplace, on
le cherche pour trouver (éventuelle-
ment) la bonne page virtuelle -

Lorsque l'on ^{veut} écrit ^(un mot d') une page
virtuelle on peut se retrouver
dans l'un de ces 3 cas : dans la
troisième on cherche ~~au hasard~~
une page ^(dans l'espace réel) libre ^{que l'on} compare
aux autres squatters -

Nous allons maintenant analyser
ces diverses méthodes dans ~~des~~ les
diagrammes comparatifs finissant.

facilité
possibilités

d'extension
en valeur loca-
daire

place occupée

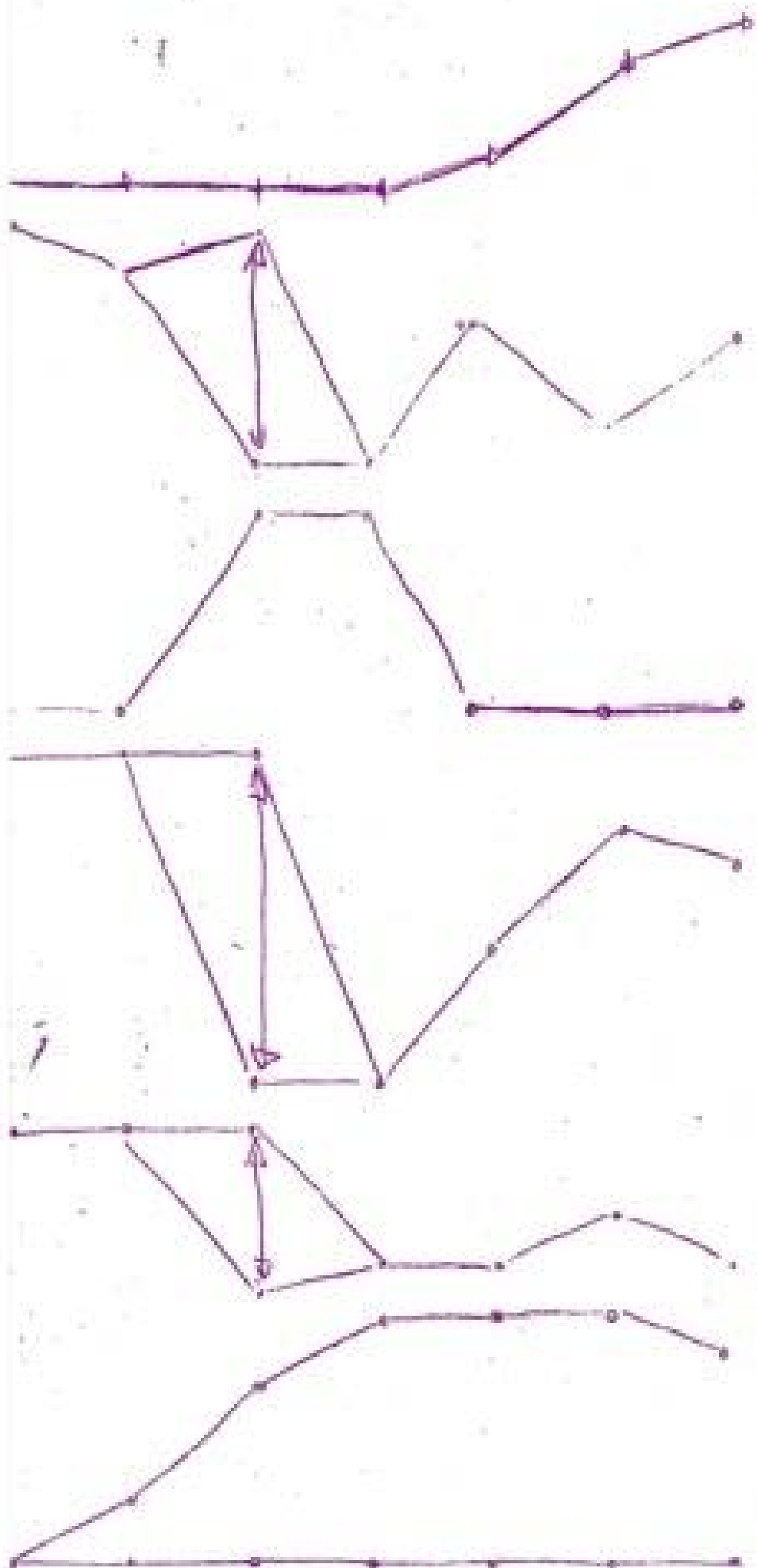
marque

apport "au milieu"

indication/Recherche

d'allocation/décroissance
allocation/croissance

Différentes
Méthodes



1. 2. 3. 4. 5. 6. 7. 8. 9. 10. 11. 12. 13. 14. 15. 16. 17. 18. 19. 20. 21. 22. 23. 24. 25. 26. 27. 28. 29. 30. 31. 32. 33. 34. 35. 36. 37. 38. 39. 40. 41. 42. 43. 44. 45. 46. 47. 48. 49. 50. 51. 52. 53. 54. 55. 56. 57. 58. 59. 60. 61. 62. 63. 64. 65. 66. 67. 68. 69. 70. 71. 72. 73. 74. 75. 76. 77. 78. 79. 80. 81. 82. 83. 84. 85. 86. 87. 88. 89. 90. 91. 92. 93. 94. 95. 96. 97. 98. 99. 100.

Sur la base de ces diagrammes (quelque peu subjectifs il est vrai) nous allons étudier les cas pratiques d'emploi de ces différentes méthodes.

Pour la méthode 3 on peut avoir une infinité de points suivants l'état de la fragmentation. On a représenté les extrêmes sur certains diagrammes.

Lorsque l'on a peu ou pas d'allocation dynamique, le peu de places et les facilités d'indirection procurées par ces méthodes les rendent intéressantes.

Les compilateurs allouent de cette façon les données statiques en mémoire centrale. A noter aussi que les machines FREEMAN et FREEMAN (1964) travaillent de cette façon. On les emploie aussi parfois dans les systèmes en temps réel pour

(de mémoire centrale)
gère des listes éventuellement dynamiques⁽¹⁹⁰⁾ (on emploie surtout alors la méthode 2) - De façon à ne pas perdre de temps en compactage lorsque l'on veut agrandir ou créer une nouvelle liste on utilise les temps morts pour opérer des compactages même si cela n'est pas nécessaire. En effet les systèmes temps réel sont caractérisés par une courbe de charge très variable donc de nombreux temps morts -

Méthodes 3 et 4

L'intérêt essentiel de ces méthodes réside dans le fait de pouvoir agrandir ou diminuer des listes par le milieu. Si par ailleurs

les besoins en indirection ne sont ⁽¹⁹⁾
pas fondamentaux par rapport
aux besoins en extension dynami-
ques on choisira ces méthodes
de préférence aux autres -

C'est le cas des langages de listes.
La méthode 4 tend maintenant
à supplanter la méthode 3 à cause
de l'économie de place qui peut
être très importante

Voir à ce sujet les articles
suivants :

- Compact list representation :
definition, garbage collection
and system implementation

par W.J. Hansen CACM Vol 12 N 9 sept 69

- A non recursive list compacting
algorithm par E.J. Cheney CACM V.13 N 11
Nov 70

La méthode 5 est intéressante car elle permet une allocation / désallocation simple (on chaîne un nouveau bloc ou on en détache un), une certaine indirection et n'utilise pas trop de place -

Cette méthode est en particulier très efficace pour manipuler des
files / ou des files d'attente en mémoire centrale
ou plus généralement des listes linéaires dont on ^{peut} indexer par les mots voisins de l'extrémité. Il suffit d'employer un registre qui pointe sur le dernier bloc chaîné -

C'est la méthode qui sert à implémenter l'allocation

dynamique de PL/1 (Voir en (193)
ce sujet le paragraphe intitulé
"Run time stacks" page 132 du
document IBM System/360

Operating System PL/1 (F) : Program
meas finale)

Cette méthode est ^{aussi} ~~généralisée~~ ^{généralisée} pour stocker
les contextes d'exécution (les registres)

des programmes écrits en assembleur.

Voir en particulier le paragraphe
intitulé "Providing a save area"
page 11 du document IBM System/360

Operating System : Super-IO ~~and~~ and
Data Management Services —

Méthode 6 (pagination)

des facilités conjuguées d'indexation et d'allocation constituent le mérite majeur de cette méthode. Les possibilités d'extension à une mémoire secondaire et sa réalisation facile par Hardware l'ont rendue très populaire. Elle est implantée sur de très nombreuses machines : CJC 10070 et IRIS 70, IBM 360/67 et au dessus, GE 645 etc...

Méthode 7

Cette méthode présente un bon compromis, elle est plus économique en place que la méthode 6 mais son implémentation par hardware est problématique — d'allocation

(195)

statique dans l'espace virtuel
présente un très grand intérêt
pour la simplification des profonds
ves; en particulier l'adresse virtuelle
l'une liste est son nom et
le catalogue peut se réduire à une
chaîne de bits. Nous reviendrons
sur ces points dans le paragraphe
consacré aux catalogues — (§ 2.4.3)

2.4.2 Représentation des tables d'occupation

Le terme "Table d'occupation" & rappelle
le sert à désigner un système
& donnée en soi qui représente
l'état d'occupation d'un support
(ici homogène)

Il y a deux cas à distinguer suivant
que les éléments (libres ou occupés)
sont de taille variable (méthodes

1, 2 et 3) ou de tailles fixe (96)
(méthodes 4, 5, 6 et 7)

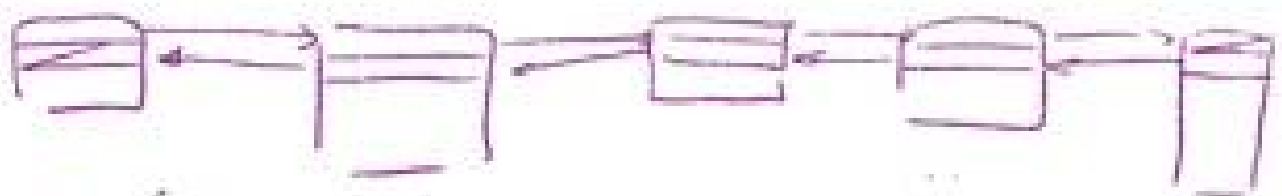
Dans tous les cas la "table
d'occupation" ne représente
finalement que les éléments
libres (les trous)

1^{er} cas : éléments de longueur variable
chaque trou contient sa propre taille -
les ~~éléments~~ ^{trous} sont chaînés entre eux.

on peut le faire de 3 façons
différentes

① par ordre chronologique

dès qu'une suite de mots ~~devient~~
libres on les ^(éventuellement doublement) chaîne à la fin
de la liste des trous.



cette méthode est rapide mais
 ne faut pas la concaténation
 des trous ~~de la zone~~ lorsque cela
 est nécessaire comme dans



la libération de la zone
 hachurée ci-contre car il n'y a

aucune raison pour les trous adjacents
~~de~~

le soit dans la liste des trous -

De façon à résoudre ce problème

on peut mettre un code spécial
 (par exemple 0) à la fin et au

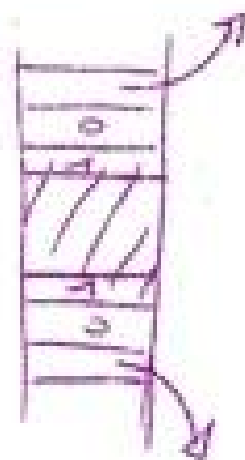
debut d'un trou - ~~différent~~

(Mais alors il faut mettre

aussi un code spécial : par

exemple 1, à la fin et

au debut d'une zone pleine)



Lorsque deux trous se fondent en un seul on est donc amené à reorganiser la chaîne des trous c'est pourquoi il est nécessaire de disposer d'un double chaînage.

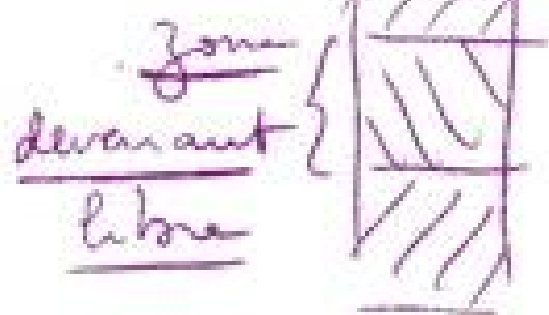
① par ordre d'adresse croissante

De façon à simplifier les problèmes précédents on peut chaîner les trous par ordre d'adresse croissante -

Cela ^{amène} ~~représente~~ une perte de temps supplémentaire lorsque l'on

est dans le cas

du schéma ci-contre devant libra



car il faut faire une

recherche dans la chaîne des trous pour insérer le nouveau trou à la bonne place.

③ par taille croissante

199

Lorsque l'on a besoin d'une zone d'une certaine taille on peut employer l'une ou l'autre des 2 méthodes

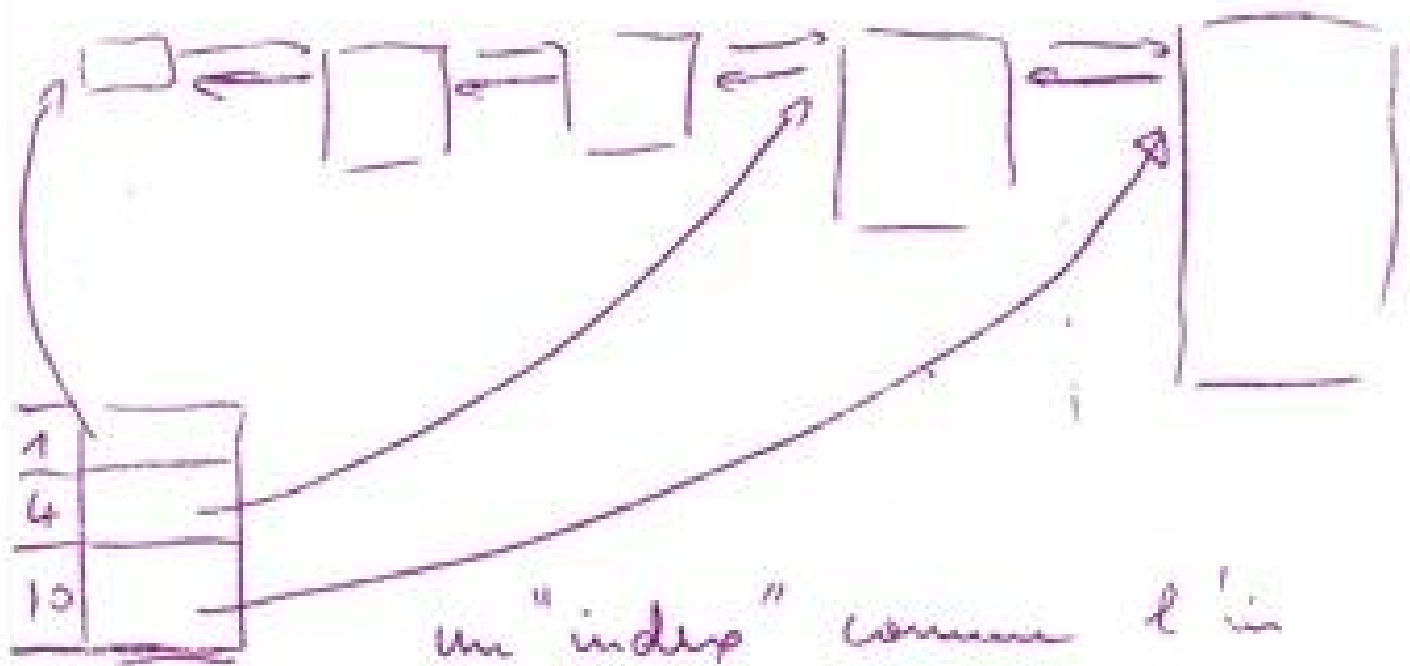
- "first fit" : on cherche dans la liste des trous le premier d'entre eux qui a une taille au moins égale à celle désirée - cette méthode

favorise la fragmentation et on peut lui préférer la suivante

- "best fit" : on cherche dans la liste des trous le meilleur d'entre eux c'est à dire le plus petit parmi ceux qui conviennent -

Cette dernière méthode peut être longue car il faut chercher un minimum c'est pour quoi on

l'assemblage des trans par feuille (200)
croissante peut être intéressant



d'après le schéma ci-dessus peut
encore améliorer cette recherche -

Une étude poussée des méthodes
"best fit" et "first fit" et
de diverses implémentations de la
dernière peut être trouvée dans
l'ouvrage: "The Art of Computer
Programming" de D. Knuth (Vol 1)

Addison-Wesley) au paragraphe (201)

5^e page 435 (Dynamique storage allocation) - Voir aussi le paragraphe "Allocation and release of storage in an Area" page 145 document IBM System/360 operating system PL/1 (P) Programming guide

2. Cas : éléments de longueur fixe

Dans ce cas la plupart des problèmes rencontrés précédemment n'existent plus ; les pages libres forment une liste libre simplement chaînée -

De façon à grouper les allocations couvrant plus d'une page dans une même zone (intéressant surtout lorsque l'on emploie une mémoire secondaire)

on peut employer la méthode de la "bit map" c'est à dire

une chaîne de bits contenant (202)
autant de bits qu'il y a de
pages : un bit égal à 1 indique
que la page correspondante est
occupée ; on trouve zéro dans
le cas contraire -

On peut alors rechercher facile-
ment un certain nombre de
pages consécutives libres en "promenant"
un masque sur la chaîne de bits.

2.4.3 Représentation des Catalogues

Le catalogue est un système
de donnée en soi qui contient
pour chaque donnée (au minimum) deux valeurs :
son nom, et un pointeur vers
son lieu -

Rappelons que le catalogue sert 203
à deux "moments" :

- lorsqu'on baptise une donnée pour vérifier que le nom que l'on donne n'est pas déjà celui d'une autre donnée déjà existante
- lorsque l'on veut faire un accès logique à une donnée pour pouvoir trouver son ~~chemin~~ ^{lien} à partir de son nom

Dans les 2 cas la vertu principale d'un catalogue doit être sa rapidité de consultation.

Implémenter un catalogue sous la forme d'une liste linéaire séquentielle ^(ordinaire) est souvent ^(une méthode) encore trop lente car : il faut parcourir

toute la liste.

(20)

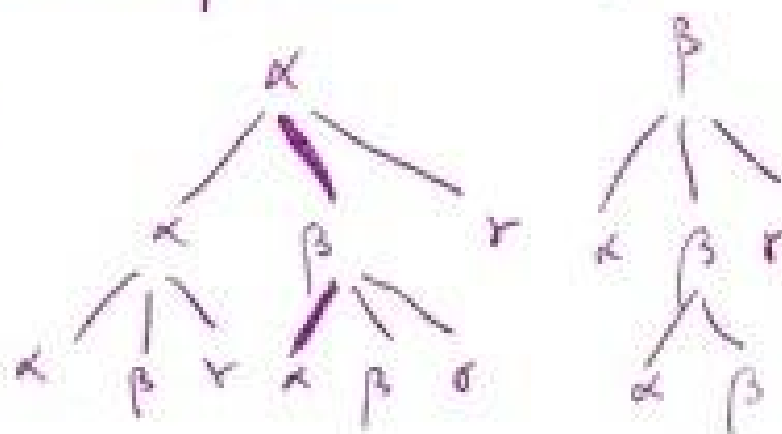
c'est pourquoi les catalogues sont

a) des listes linéaires séquentielles
classés par ~~une~~ ordre ~~alphabétique~~ des
noms (pour pouvoir faire des
recherches dichotomiques) ;

b) des structures en arbres
(ou mieux en forêt) : un nom est
alors en fait une suite de chaînes
alphabétiques. Les nœuds
fils d'un nœud donné sont classés
par ordre - c'est le cas
de la plupart des catalogues de
fichiers dans les systèmes d'exploit
ation classés

Exemple

les noms peuvent être
 α, β, α



Les "Index" des clés des enregistrements ⁽¹⁰⁾
des fichiers indexés séquentiels sont
organisés aussi de cette façon

③ des listes linéaires à support
séquentiel mais remplies de façon
pseudo-aléatoire par hachcoding

Cette méthode est rapide mais
elle prend beaucoup de place -
De plus lorsque le ~~catalogue~~
catalogue hachcodé est saturé on
ne peut plus l'étendre dynamiquement
sans en faire une reorganisa-
tion complète

C'est pourquoi cette méthode est
intéressante lorsque l'on peut se
fixer avec certitude (ou imposer
aux utilisateurs) un nombre maxi

nom de nom à entrer dans le (206)
catalogue et lorsque d'autres
ont le catalogue (qui prend
beaucoup de place) a une durée
de vie limitée -

C'est typiquement le cas pour
les catalogues qui contiennent
les noms génériques des données
utilisés dans les programmes
~~des~~ assemblés ou compilés : c'est
à dire les Table des symboles
les compilateurs ou assembleurs

④ Cas particuliers

Dans le système SOCRATE, l'adresse
virtuelle d'une donnée est son nom
et le mécanisme de passage de
l'adresse virtuelle à l'adresse réelle

se fait par programme : 1^o 2^o (204)
fonction du catalogue est donc
inutile.

La première fonction est réalisée
par une chaîne de bits.

Exemple : si l'on veut réaliser
un système pouvant supporter 1024
listes linéaires on devra utiliser
une chaîne de 128 octets. Chaque
bit indique la liste correspon-
dante existe ou non.

Si l'on desire par ailleurs donner
un nom externe alphabétique
à une donnée on peut réaliser
un catalogue bas codé contenant
pour chaque donnée son nom externe
et son nom interne (son adresse
virtuelle). Cependant on minimise
la place occupée par ce 2^e catalogue.

car il est lui-même "virtualisé" (208)
et les trous qu'il comporte n'ont
pas leur contrepartie dans l'espace
réel -

2.4.4 Milieu mémoire non homo-

gène

Comme nous l'avons précisé plus
haut nous allons considérer le
cas de 2 milieux :

- l'un le milieu tore ayant
un temps d'accès de l'ordre de 1 μ sec

- l'autre le disque ayant un
temps d'accès de l'ordre de 100 ms

(on pourrait aussi considérer le

cas du Tambour ou du disque à
tête fixe ayant un temps d'accès
de l'ordre de 10 ms)

Toute la difficulté de l'allocation 239
dans ce cas vient du
rapport 10000 (ou 100000) entre
les temps d'accès -

Le milieu Tore ne sert que pour
les traitements, les listes sont
normalement stockées sur le
Disque -

On peut considérer plusieurs cas :

① la taille du tore est suffisante
pour contenir entièrement la
liste linéaire stockée de façon
séquentielle sur le disque
On entre alors toute la liste
en mémoire centrale lorsque l'on
en a besoin - Si l'on désire accéder

a une autre liste ou la remettre ^(2^{es})
en mémoire centrale s'il y a encore
de la place pour contenir entièrement
cette nouvelle liste ou bien dans
le cas contraire on remet sur
disque l'ancienne pour ~~suffire~~ faire de
la place - Cette méthode est
connue sous le nom de swapping -
Elle a été utilisée dans les premiers
systèmes fonctionnant en multiprogram-
mation (les listes linéaires en
fonction etant alors des programmes)
plus particulièrement pour la Burroughs
B 5000 et les suivantes -

⑤ La taille du tou n'est
pas suffisante pour contenir la
liste sur laquelle on veut

travailler mais on est capable (211)
d'indiquer a priori la manière
dont on va employer la liste
(ou plus largement la structure
de données) de façon à faire
des entrées partielles en mémoire
centrale qui vont se recouvrir
les unes les autres - on place
ces entrées partielles dans une
zone tampon (buffer) dont la
taille est la plus grande taille
parmi celles des structures partielles
sur l'entrée - Dans le cas
où la structure en question est
un programme on dit sur l'on

fait de l'"overlay" - Dans le (212)
cas où la structure en question
est une donnée on a à faire à un
fichier dont les différents enregistre-
ments se ~~retrouvent~~ en mémoire
centrale sur le buffer - Dans ce
dernier cas, de façon à ~~opti-~~
miser l'accès des données, on
peut (surtout lorsque l'on doit
pratiquer un parcours séquentiel sur
les enregistrements) rentrer plusieurs
enregistrements par le buffer en une seule opé-
ration d'accès : le nombre d'enregis-
trements que l'on rentre simultané-
ment s'appelle le facteur de
blockage -

(c) La mémoire tore ne peut (213)
contenir la totalité de chaque liste
et de plus il est impossible de
prévoir à l'avance le comportement
du programme de traitement à
l'égard de la liste. On pratique
alors une homogénéisation du tore et
du disque par pagination -

Le disque est divisé en un certain
nombre de pages - Le tore est
divisé de même en un certain nombre
de cadres de la même taille que
les pages -

Supposons pour fixer les idées une
zone dynamique en mémoire centrale
de 150 K ^(octets) et que l'on pagine ~~à~~ 100
Millions d'octets sur disque, on aura

alors si la taille d'une page est de 2 K octets :

75 cadres en memoire centrale

50.000 pages sur disque

les cadres sont gérés par une table d'occupation comprenant pour chacune des 75 entrées l'adresse disque de la page correspondante ainsi que des renseignements sur la date de sa consultation et le fait qu'elle a été mise à jour ou non -

| | | | |
|----|---------|------|-----|
| 1 | | | |
| | | | |
| 75 | Adresse | date | maj |
| | | | |
| | | | |
| | | | |

Table
d'occupation
des
Cadres

le fonctionnement est très simple : (2/5)

lorsque l'on veut consulter une page
on cherche dans la table d'occupation
des ~~don~~ cadres si elle est là

- 1° oui on la consulte

- 2° non on choisit un cadre

pour servir de place à la page que
l'on va amener depuis le disque -

De nouveau 3 cas

- le cadre choisi est inoccupé

- le cadre choisi est occupé par

une page non modifiée : on la recopie

- le cadre choisi est occupé par

une page modifiée : on la recrée sur

le disque avant d'utiliser le cadre

d'algorithme que nous venons de

défini s'appelle un algorithme (216)
le "demand paging" par opposi-
tion aux algorithmes où l'on a
toujours "un coup d'avance" de façon
à avoir systématiquement un cadre
libre en mémoire centrale lorsque
l'on a besoin de rentrer une nou-
velle page -

L'algorithme de choix des ~~à~~ cadre
à libérer ou "page removal algorithm"
peut faire appel à différents
critères : le but ^à atteindre est
évidemment ~~nécessaire~~ d'utiliser
le cadre qui contient une
page qui ne risque pas d'être réutilisée

dans un proche avenir (le
cas d'un cadre ne contenant
pas de page ne se rencontre ~~plus~~
pas en dehors de la période tout
à fait initiale de fonctionnement)

Une critère classique consiste à
choisir comme cadre celui qui
contient la page la plus ancien
nement consultée. De façon

à ce fin ~~plus~~ le cadre soit
déterminé rapidement on ^{doublement} peut vérifier

entre ^{les} ~~plus~~ les cadres par ordre de
consultation (chaque fois que l'on
consulte un cadre on ~~peut~~ le place
en queue de la file ^à : celui qui
est le premier de la file est le
plus anciennement consulté)

La disposition d'un milieu
mémoire intermédiaire entre le
Tore et le disque améliore
fondamentalement le comportement
général de tels systèmes :

on peut placer des milieux
intermédiaires soit "pres" du tore
("extended core memories" ou les plus
rares) soit pres du disque (Tambour
ou disque à têtes fixes) - Pour
des raisons de prix on trouve
généralement des Tambours ou
des disques à têtes fixes -

Pour une étude plus complète
de toutes ces techniques on pourra
se reporter avec profit à l'article

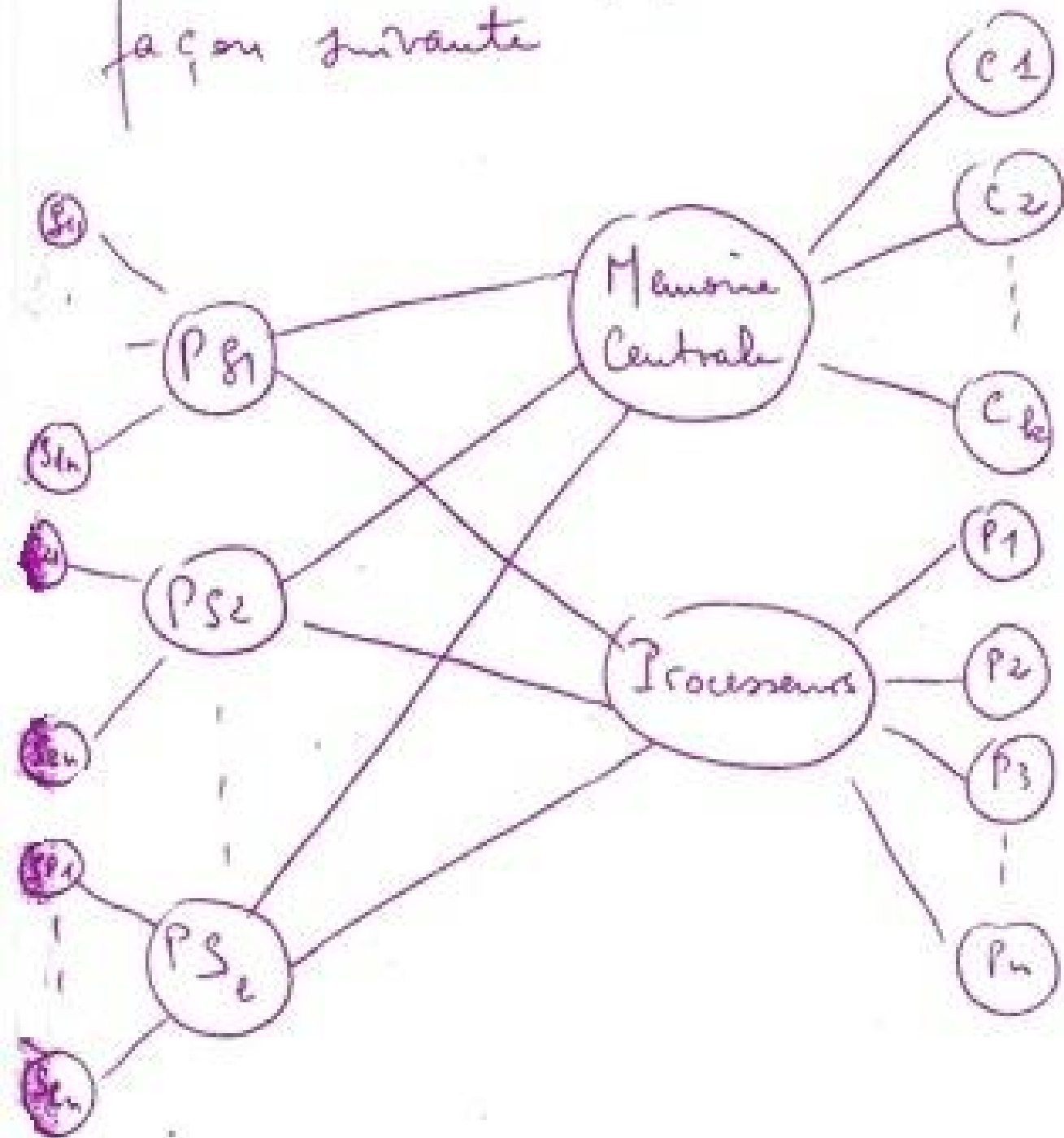
suivant: "Virtual memory" (219)
par P. J. Denning Computer
surveys Vol 2 N.3 Sept 70 -

J'espère que le lecteur qui
réfléchit n'aura pas manqué
de remarquer l'analogie
très grande qui existe entre
le "page removal algorithm"
et le Regulateur de circuit page

(152) - En fait il s'agit
~~de~~ du même problème : on
est en présence dans les 2 cas
d'une ressource quantifiée :
la classe de hardware dont les
quantas sont les différents hardware

de la classe et la mémoire
centrale dont les fonctions sont
les différents cadres.

Il convient donc d'améliorer le
schéma de la page (165) de la
façon suivante



Chaque serveur ne se trouve plus
dans une seule file d'attente :

1 file de garçon (processeurs) mais
il se trouve aussi dans la file
d'attente file de garçon (mémoire centrale)

Les deux Régulateurs :

Régulateur (processeurs) et

Régulateur (mémoire centrale) ne
doivent donc pas travailler indépendan-
ment mais au contraire se coupler
de façon à donner de façon
harmonieuse un processeur et un
ou plusieurs) cadric(s) au même ^{serveur} ~~systeme~~
et ceci le plus simultanément
possible de façon à ce qu'un

ne trouve pas dans la situation (222)
(cela fréquente dans la pratique)
d'un client fin² un processeur et
pas de mémoire ou l'inverse -

Ceci est un des problèmes les plus
difficile à résoudre actuellement
dans l'écriture des systèmes -

De la même façon que nous
avons suggéré la réalisation
d'un organe ^{unique} chargé de réguler
toutes les classes de hardware
on pourrait ~~attribuer~~ adjoindre
à cet organe le rôle supplémentaire
de réguler la mémoire centrale -
Ceci faciliterait d'autant plus
les complaisances entre régulation -