

Retournement de l'aimantation des nanoaimants

Edgar Bonet

Version 2.0

Remerciements

Un grand merci tout d'abord à ceux qui m'ont accueilli dans ce laboratoire. Je veux nommer en particulier Dominique Givord, qui a dirigé le laboratoire Louis Néel pendant toute la durée de ma thèse, Bernard Barbara, qui a dirigé cette thèse, ainsi que Wolfgang Wernsdorfer, qui m'a encadré au jour le jour pendant ces trois années. Merci aussi à ceux qui ont accepté de participer à mon jury, et dont certains ont dû venir de Paris : outre Bernard et Wolfgang, qui sont ainsi doublement remerciés, il y a aussi Albert Fert, Jacques Ferré, André Thiaville, Jean-Louis Porteseil et Alain Benoît. André Thiaville doit être remercié aussi pour les nombreuses discussions que nous avons eues ensemble, et qui m'ont fait aimer la géométrie des surfaces critiques.

Je remercie aussi tous les autres membres du laboratoire, puisque c'est grâce à eux que le laboratoire Louis Néel est un endroit agréable, et en particulier Christophe Thirion et Irinel Chiorescu, qui ont beaucoup participé à la bonne humeur de notre petite équipe.

Pendant le début de ma thèse, j'ai passé quelque temps dans le laboratoire voisin : le CRTBT. Je remercie donc une deuxième fois Alain Benoît, qui m'a un peu encadré là-bas, en m'initiant aux plaisirs de l'électronique numérique et des circuits d'acquisition. Merci aussi à Bartosz, Armelle, Sarah et Franck pour les bons moments que nous avons passés ensemble au CRTBT.

Je ne remercierai probablement jamais assez Stéphanie, qui m'a fait voir une autre lumière, et tout en m'encourageant me fait comprendre que la physique n'est pas la seule chose qui compte dans la vie. Merci aussi à mon cher papa, qui était un fan-club à lui tout seul. Merci enfin à ceux qui, étant mes cothurnes, ont dû me supporter à la maison : Michel et Mélanie, mais aussi, même s'ils le sont restés moins longtemps, Petro et Benoît.

Table des matières

Introduction	9
Notations	13
I Théorie du retournement uniforme de l'aimantation	15
1 Théorie quasi-statique	17
1.1 Retournement uniforme ou non uniforme?	17
1.1.1 Les énergies en compétition	18
1.1.2 Configurations possibles	19
1.1.3 Près du retournement	22
1.1.4 Calculs plus précis	23
1.2 Théorie du retournement uniforme	24
1.3 Modèle de Stoner-Wohlfarth	25
1.4 Anisotropie arbitraire	27
1.4.1 Notation complexe	27
1.4.2 Points d'équilibre	27
1.4.3 Sauts d'aimantation	28
1.4.4 Autre description	28
1.5 Bifurcations	29
1.6 En trois dimensions	32
2 Théorie dynamique	33
2.1 Modèle de Néel-Brown	33
2.1.1 Modèle des orientations discrètes	34
2.1.2 Modèle de Néel	35
2.1.3 Modèle de Brown	36
2.1.4 Approximation de la barrière haute	38
2.2 Reformulation de Kurkijärvi	40

3	Calculs numériques	43
3.1	Position du problème	43
3.2	Plongement dans l'espace 3D	44
3.3	Sauts d'aimantation	45
3.4	Forme du potentiel au voisinage du saut	46
3.4.1	Hauteur de la barrière d'énergie	47
3.4.2	Courbures du potentiel	48
3.4.3	Temps caractéristique τ_0	48
II	Techniques expérimentales	51
4	Présentation de l'expérience	53
4.1	Cryogénie	54
5	Chaîne de mesure	57
5.1	Micro-SQUID	57
5.1.1	Une théorie simpliste	57
5.1.2	Comportement observé	59
5.2	Circuit de détection	60
5.3	Interface de contrôle	61
5.4	Carte d'acquisition	61
5.5	Logiciel	62
5.5.1	Routine temps réel	63
5.5.2	API de la routine temps réel	63
5.5.3	Interface utilisateur	63
6	Protocole de mesure	65
6.1	Commande du champ	65
6.1.1	Schéma de principe	65
6.1.2	Comment obtenir le champ qu'on veut	66
6.1.3	Effet de la discrétisation du temps	72
6.2	Modes de mesure	73
6.2.1	Cycle d'hystérésis	73
6.2.2	Mode froid	74
6.2.3	Mode aveugle	74
7	Problème de l'oracle qui a bu	77
7.1	Posons le problème	77
7.2	Éléments de réponse	78
7.2.1	Reposer la question pour confirmer la réponse	78

7.2.2	Demander n confirmations	79
7.2.3	Maximiser l'information soutirée à l'oracle	81
7.2.4	Autres pistes	83
III	Résultats	85
8	Mesures quasi-statiques	87
8.1	En deux dimensions	87
8.1.1	Particule de Cobalt	87
8.1.2	Particule de fer-cuivre	91
8.2	En trois dimensions	93
8.2.1	Particules de cobalt	93
8.2.2	Particule de ferrite de baryum	95
9	Mesures dynamiques	99
9.1	Champs de retournement	99
9.2	Dépendance par rapport au point de la surface critique	100
9.3	Bruit télégraphique	101
10	Discussion	105
10.1	Interactions avec le milieu	105
10.2	Hypothèse de l'aimantation uniforme	106
10.3	Spins nucléaires et dissipation	107
	Bibliographie	109
	Conclusion	113
	Annexes	119
A	Compléments mathématiques	119
A.1	Quelques anisotropies	119
A.1.1	Anisotropie quadratique	119
A.1.2	Anisotropie cubique	119
A.1.3	Uniaxial d'ordre 4	120
A.1.4	Uniaxial d'ordre 6	120
A.2	Changements de repère orthonormé	120
A.2.1	Formes linéaires, quadratiques et cubiques	121
A.2.2	Base propre d'une forme quadratique	122
A.2.3	Angles d'Euler	122

B	Mode d'emploi de la carte d'acquisition	125
B.1	Description de la carte	125
B.1.1	Schéma général	125
B.1.2	Aspect physique	126
B.1.3	Connecteur externe	127
B.1.4	Modes de fonctionnement	128
B.1.5	Registre <i>flags</i>	129
B.2	Interface C	130
B.2.1	Description des fonctions	130
B.2.2	Utilisation de plusieurs cartes	131
B.2.3	Exemples de code	132

Introduction

Motivations de ce travail

On a vu depuis quelques années le terme *nano-science* fleurir dans des titres de colloques et congrès. Cet engouement vient du fait que la physique, qui s'est intéressée au cours de la dernière décennie à l'étude de systèmes de plus en plus petits, se concentre maintenant dans l'étude des systèmes de taille nanoscopique. Le magnétisme trouve un très grand intérêt dans cette mouvance. En effet, en magnétisme l'échelle nanométrique s'est révélée être l'échelle mésoscopique, c'est à dire celle où le comportement passe d'un comportement macroscopique (comme dans un matériau massif) à un comportement microscopique (comme dans des atomes isolés). Cette échelle s'est révélée d'une grande richesse en phénomènes nouveaux. Ainsi, en magnétisme macroscopique les longueurs caractéristiques sont la longueur d'échange et l'épaisseur d'une paroi de Bloch. Ceci suffit à expliquer que les particules magnétiques dans cette gamme de taille présentent un intérêt singulier [1]. Les théories qui traitent de ces particules sont utilisées depuis un demi siècle pour l'interprétation d'expériences portant sur des grandes assemblées de particules. Dans de tels systèmes, le comportement individuel de chaque particule est caché derrière des distributions de taille, d'orientation et de forme. Pour valider et raffiner les modèles, il est donc nécessaire de les confronter à des expériences réalisées sur des particules individuelles.

Ces particules sont aussi intéressantes pour les applications. Dans le développement d'aimants permanents on s'intéresse à des matériaux nanostructurés dont les composants les plus coercitifs sont comparables aux nanoparticules qui nous intéressent. Pourtant, le principal intérêt de ces systèmes semble se situer dans l'enregistrement magnétique à haute densité. Dans ces dernières années, la densité de stockage sur des supports magnétiques a augmenté de 60 % par an. Aujourd'hui, des densités de l'ordre de cinq mégabits par millimètre carré (5 Mb/mm²) sont déjà utilisées commercialement. Si cette évolution se poursuit, il faudra très rapidement se tourner vers des particules de taille nanométrique comme support de l'information. Actuellement la limite superparamagnétique (qu'on situe vers 60 Mb/mm²) semble être la première difficulté. Cependant, grâce à des nouveaux matériaux à très forte

anisotropie, cette limite devrait pouvoir être efficacement repoussée. Il est alors vraisemblable que dans le futur on utilisera des particules dans des gammes de taille qui aujourd'hui sont au delà de cette limite. L'évolution de l'enregistrement magnétique se faisant aussi vers des vitesses d'enregistrement de plus en plus élevées, l'étude de la dynamique du renversement de l'aimantation dans les nanoparticules promet d'être un sujet d'intérêt industriel dans un futur proche.

Des nouvelles techniques expérimentales ont permis très récemment de faire des mesures magnétiques sur des nanoparticules isolées. Citons par exemple la microscopie à force magnétique [2, 3], la magnétorésistance [4, 5] les micro-sondes de Hall [6, 7, 8], la microscopie de Lorentz [9], la magnétométrie à micro-SQUID [10] et la spectroscopie tunnel à travers une nanoconstriction [11]. Ces techniques évoluent continuellement vers une sensibilité de plus en plus grande. Alors que les premières nanoparticules étudiées isolément avaient des tailles de plusieurs centaines de nanomètres, nous arrivons actuellement dans la gamme de trois à cinq nanomètres.

Mon travail de thèse a consisté à développer et à utiliser la technique du magnétomètre à micro-SQUID. En travaillant dans la gamme de 15 à 25 nm, j'ai réalisé des mesures montrant la nécessité d'étendre les modèles théoriques actuels afin de tenir compte des trois dimensions de l'espace et d'une anisotropie magnétique arbitraire. Ces extensions ayant été faites par André Thiaville, j'ai mis au point une méthode de mesure des champs de retournement en trois dimensions qui ont donné un support expérimental à cette nouvelle théorie. Des mesures sur la dynamique confirment que les particules ont un comportement classique bien décrit par le modèle de Néel-Brown au moins pour des températures au dessus de 1 K, les résultats n'étant pas très concluants en dessous. Toutes ces mesures ont été réalisées à des températures comprises entre 50 mK à 4 K et pour des champs appliqués allant jusqu'au demi tesla environ.

Présentation du manuscrit

La première partie de ce manuscrit s'intéresse à la théorie de ces nanoparticules. Nous verrons que quand la taille de la particule est suffisamment petite, l'échange est le terme dominant de l'énergie, ce qui conduit à des configurations en monodomaine. L'aimantation étant uniforme et ayant sa valeur à saturation, le moment magnétique de la particule est constant en module mais a la liberté de changer d'orientation. Plusieurs phénomènes sont responsables d'une anisotropie magnétique qui s'exprime comme un terme de l'énergie dépendant de cette orientation. On notera ici $V(\mathbf{r})$ ce potentiel, où $\mathbf{r}$ est l'orientation de l'aimantation. À température nulle, l'aimantation se trouve en permanence dans l'un des minima de V . L'application d'un champ magnétique extérieur modifie le potentiel V et permet donc de modifier l'orientation $\mathbf{r}$ de l'aimantation. Nous verrons que l'irréversibilité de ces changements est liée à

l'existence de champs critiques pour lesquels l'aimantation fait des sauts discontinus. Ces champs sont ceux pour lesquels une position d'équilibre stable devient instable. Autrement dit, ce sont les champs pour lesquels une barrière d'énergie autour d'un état métastable disparaît. Le lieu des champs critiques est étroitement lié à l'anisotropie de la particule. Ainsi, en mesurant les premiers il est possible de remonter à la seconde par un ajustement.

Quand la particule est à température non nulle, l'activation thermique permet à l'aimantation de franchir une barrière de potentiel qui n'a pas été complètement annulée par le champ appliqué. Le temps qui caractérise ce franchissement dépend fortement de la hauteur de cette barrière, de la température, mais aussi de choses plus difficiles à connaître comme l'importance de la dissipation dans la dynamique de l'aimantation ou les courbures du potentiel V . Il est intéressant de noter que ces courbures peuvent être calculées si on arrive à déterminer l'anisotropie à partir de la mesure des champs critiques.

La deuxième partie concerne les méthodes expérimentales utilisées. Notre dispositif expérimental comporte beaucoup d'éléments maison qu'il a fallu développer ou adapter. La chaîne de mesure commence par le micro-SQUID. Celui-ci se présente comme une boucle supraconductrice sensible au flux magnétique qui la traverse. En posant les nanoparticules directement sur le SQUID, on obtient un couplage direct très fort qui permet d'avoir une sensibilité très élevée en moment magnétique. Le micro-SQUID est contrôlé par un circuit analogique secondé par une logique programmable. Ce dispositif avait été mis au point au CRTBT pour la mesure de courants permanents et il a dû être modifié par reprogrammation de la logique afin de l'adapter à nos expériences. Toute la chaîne de mesure, ainsi que les alimentations de courant des bobines de champ, sont contrôlées par un microordinateur via une carte de communication et logiciel développés en interne. Les différents modes de fonctionnement de l'expérience sont implémentés dans le logiciel pour s'adapter aux divers types de mesures qu'on veut faire. Un soin particulier a été apporté à la commande du champ magnétique pour tenir compte du temps de retard de l'ensemble alimentation + bobines. Le dernier chapitre de cette partie a un caractère un peu plus récréatif et présente un problème algorithmique qu'il a fallu traiter pour avoir une bonne fiabilité dans les mesures.

La troisième partie présente les résultats obtenus avec ce magnétomètre. Les mesures effectuées en deux dimensions sur une particule de fer-cuivre ont montré la nécessité de développer une théorie du retournement uniforme qui s'adapte à une anisotropie arbitraire et en trois dimensions. Des résultats qualitatifs importants de cette nouvelle théorie ont été mis en évidence expérimentalement dans ces mesures. Nous présentons ensuite les premières mesures de surfaces critiques en trois dimensions. Ces mesures ont permis d'obtenir les premières déterminations expérimentales de l'anisotropie effective de nanoparticules individuelles.

Les mesures dynamiques confirment dans leur ensemble que le modèle de Néel-

Brown donne une description correcte de la dynamique du retournement de l'aimantation dès que les particules étudiées ne présentent pas de désordre magnétique en surface. On note cependant quelques déviations par rapport à ce modèle aux températures les plus basses, en particulier pour les particules de ferrite de baryum. Ces déviations ont pu être interprétées comme la signature d'un effet tunnel de l'aimantation. Cette interprétation est toutefois sujette à caution. L'observation d'un effet tunnel résonnant serait une preuve beaucoup plus convainquante de l'existence de cet effet tunnel. Les derniers développements de notre expérience ayant permis d'atteindre des sensibilités bien plus élevées, l'espoir d'obtenir une telle preuve sur des systèmes plus petits est permis.

Notations

Nous utiliserons, tout au long de ce manuscrit, le système international de poids et mesures : le système SI. Voici une liste des notations qui seront couramment employées. Entre parenthèses est notée l'unité SI correspondant à chacune de ces grandeurs.

$\mathbf{H}$	champ magnétique appliqué (A/m) ;
$H = \mathbf{H} $	module de $\mathbf{H}$ (A/m) ;
$\mathbf{H}_c$	champ critique de retournement de la particule (A/m) ;
$\delta\mathbf{H} = \mathbf{H} - \mathbf{H}_c$	écart au champ critique (A/m) ;
$\delta H = \delta\mathbf{H} $	module de $\delta\mathbf{H}$;
$\mathbf{M}$	aimantation de la particule (A/m) ;
$M_s = \mathbf{M} $	aimantation à saturation (A/m) ;
$\mathbf{r} = \mathbf{M}/M_s$	direction de $\mathbf{M}$ ($ \mathbf{r} = 1$) ;
v	volume de la particule (m ³) ;
$v\mathbf{M}$	moment magnétique (Am ²) ;
γ	rapport gyromagnétique (T ⁻¹ s ⁻¹) ;
$\mathbf{h} = \mu_0 v M_s \mathbf{H}$	champ réduit (J) ;
$h = \mathbf{h} $	module de $\mathbf{h}$ (J) ;
K	constante d'anisotropie magnétocristalline (J/m ³) ;
J	constante d'échange (J/m ³) ;
$V_0(\mathbf{r})$	énergie d'anisotropie magnétique (J) ;
$V = V_0 - \mu_0 v \mathbf{M} \cdot \mathbf{H}$	énergie magnétique totale (J) ;
k_B	constante de Boltzmann (J/K).

On utilisera bien entendu les multiples et sous multiples des unités SI. Dans la pratique, et par respect pour l'usage établi, nous n'exprimerons pas les champs $\mathbf{H}$ en A/m mais toujours $\mu_0 \mathbf{H}$ en T (ou plutôt en mT). De même, plutôt que d'exprimer V en J, on exprimera V/k_B en K.

Une entorse aux règles sera faite au niveau du traitement théorique de l'équation

$$V = V_0 - \mathbf{r} \cdot \mathbf{h}.$$

En principe, les trois termes de cette équation sont des énergies et s'expriment en joules. Cependant, lorsque nous traiterons des exemples purement théoriques, nous considérerons les trois termes comme adimensionnés.

Première partie

Théorie du retournement uniforme
de l'aimantation

Chapitre 1

Théorie quasi-statique

Les trois chapitres qui constituent cette partie traitent de la théorie du comportement magnétique de nanoparticules individuelles. Dans ces trois chapitres, le système considéré est une particule ferromagnétique monocristalline, isolée de son milieu extérieur sauf en ce qui concerne le champ magnétique appliqué par l'expérimentateur et un éventuel thermostat. C'est le comportement de son aimantation qui sera l'objet de notre intérêt.

Dans ce chapitre nous allons nous intéresser à la limite $T \rightarrow 0$ et nous allons négliger tous les effets quantiques ainsi que ceux liés à la dynamique du système. Cette dernière hypothèse revient à considérer que la dynamique propre du système est très rapide comparée aux temps accessibles expérimentalement. Dans ces conditions, le système qu'on considère va se trouver à chaque instant dans un état d'équilibre qui correspond à un minimum local de son énergie. Lorsqu'on fait varier le champ appliqué, ces variations seront suffisamment lentes pour qu'on puisse considérer l'évolution du système comme une succession d'états d'équilibre.

En considérant la limite $T \rightarrow 0$, on interdit au système de franchir toute barrière de potentiel. Ainsi, lorsqu'il se trouve dans un état qui minimise localement son énergie, il y reste tant que cet état demeure un état d'équilibre stable.

1.1 Retournement uniforme ou non uniforme ?

Une particule ferromagnétique présente en champ nul au moins deux configurations magnétiques stables. En effet, toutes les contributions à l'énergie de la particule sont invariantes par la transformation qui change en tout point l'aimantation $\mathbf{M}$ en $-\mathbf{M}$. Cette transformation permet donc, dès qu'on connaît une configuration stable, d'en déduire une deuxième.

Nous pouvons essayer, par l'application d'un champ magnétique extérieur, de faire passer le système d'une de ces configurations à une autre. En faisant varier

continuellement ce champ appliqué, nous allons modifier les configurations qui correspondent aux minima locaux de l'énergie. Cette modification sera aussi continue, sauf pour certains champs appelés *champs critiques* qui correspondent à l'apparition ou à la disparition d'un de ces minima. Ce sont ces champs critiques qui vont nous occuper tout au long de ce chapitre.

Supposons que la particule se trouve dans un certain état d'équilibre qui minimise localement son énergie. Supposons aussi que qu'on atteigne le champ critique pour lequel ce minimum disparaît. Pour ce champ particulier, l'état d'équilibre considéré reste toujours un état d'équilibre, mais il a perdu sa stabilité et est devenu un état d'*équilibre instable*. Dans cette situation, une perturbation infinitésimale peut s'amplifier jusqu'à emmener le système très loin de son équilibre de départ et, *in fine*, jusqu'à un autre état stable. C'est ce phénomène qu'on appelle *retournement de l'aimantation*. La perturbation infinitésimale qui est susceptible de s'amplifier ainsi est en fait une solution des équations de la dynamique de l'aimantation, linéarisées autour de l'équilibre instable. On l'appelle *mode de retournement* et on peut la voir comme la description de la façon dont le retournement *commence*. Connaître la façon dont le retournement se poursuit est quelque chose de plus difficile qui met en jeu des équations non linéaires complexes.

Il est temps maintenant de préciser ce qu'on entend par *retournement uniforme*. On parlera de retournement uniforme lorsque deux conditions sont réunies. D'une part, l'état d'équilibre qu'on va déstabiliser est un état à aimantation uniforme, c'est à dire un état dans lequel l'aimantation est constante dans tout l'échantillon. D'autre part, le mode de retournement est un mode uniforme, ce qui veut dire que, au début du retournement, l'aimantation varie de manière identique en tous les points de l'échantillon. Cela *ne* veut *pas* dire que l'aimantation demeure uniforme tout au long du processus de retournement, mais seulement au tout début.

Dans la suite de cette section nous allons essayer de déterminer les conditions dans lesquelles on peut s'attendre à avoir des états d'équilibre à aimantation uniforme et, dans une certaine mesure, des modes de retournement uniformes.

1.1.1 Les énergies en compétition

La nature des états d'équilibre est déterminée par la compétition entre différentes contributions à l'énergie totale de la particule. Nous allons essayer d'évaluer ces différentes contributions afin de déterminer les conditions sous lesquelles les états d'équilibre sont à aimantation uniforme. Nous n'utiliserons aucun facteur numérique dans cette évaluation. Il ne s'agit donc que d'une sorte de calcul « aux dimensions » qui cherche à mettre en évidence les paramètres pertinents mais qui ne sera valable qu'à des facteurs numériques près.

Les énergies mises en jeu dans ce problème sont :

- l'énergie d'échange $E_e = J \int (\nabla \cdot \mathbf{M})^2 / M_s^2$;

- l'anisotropie $E_a = K \int 1 - (\mathbf{M} \cdot \mathbf{u})^2 / M_s^2$;
- l'énergie du champ démagnétisant $E_d = \int \mu_0 \mathbf{M} \cdot \mathbf{H}_d$;

où J désigne la constante d'échange, K la constante d'anisotropie (dans le cas d'une simple anisotropie uniaxiale), $\mathbf{u}$ un vecteur unitaire le long de l'axe de facile aimantation et $\mathbf{H}_d$ le champ démagnétisant (champ produit par $\mathbf{M}$ lui même). Les intégrales s'étendent sur tout le volume de la particule.

Supposons que la particule est de taille L et que les variations d'aimantation au sein de cette particule se font sur des longueurs de l'ordre de l . On suppose que là où l'aimantation est uniforme, elle est parallèle à une direction de facile aimantation. L'énergie d'anisotropie est donc nulle à cet endroit. Dans les énergies d'échange et d'anisotropie l'intégration se fera alors dans un volume de l'ordre de $L^2 l$. Pour l'énergie du champ démagnétisant ce sera sur le volume L^3 de la particule. Les énergies considérées ci-dessus vont donc varier comme

- $E_e \sim JL^2 l^{-1}$;
- $E_a \sim KL^2 l$;
- $E_d \sim \mu_0 M_s^2 L^3$.

1.1.2 Configurations possibles

Nous allons considérer trois types de configurations magnétiques et comparer leurs énergies afin de déterminer, en fonction des paramètres de la particule, laquelle est la plus stable. Les configurations considérées sont les suivantes :

Aimantation uniforme

Dans cette configuration, l'aimantation est constante sur tout l'échantillon. Elle est alors parallèle à l'une des directions de facile aimantation. Les énergies d'échange et d'anisotropie sont donc nulles, mais il reste l'énergie du champ démagnétisant. L'énergie de cette configuration est donc

$$E_1 \sim \mu_0 M_s^2 L^3.$$

Aimantation non homogène

Ici on considère une configuration dans laquelle l'aimantation varie sur *tout* l'échantillon. Un exemple d'une telle configuration pourrait être un vortex. La longueur caractéristique sur laquelle varie l'aimantation est donc la taille de l'échantillon ($l = L$). On suppose qu'en adoptant une telle configuration, l'aimantation s'arrange pour annuler (ou presque) les charges magnétiques et donc aussi l'énergie dipolaire par la même occasion. En revanche, il restera dans cette configuration aussi bien de l'énergie d'échange que de l'anisotropie. En considérant que la somme de ces deux

contributions est approximativement égale à la plus grande d'entre-elles, l'énergie de cette configuration est alors

$$E_2 \sim \sup(JL, KL^3).$$

Paroi de domaine

Pour cette troisième configuration, on va considérer que l'aimantation est uniforme sur toute la particule excepté dans l'épaisseur d'une ou plusieurs parois de domaine. On suppose que la formation de domaines permet d'annuler (ou presque) les charges magnétiques et l'énergie dipolaire. Ce cas est similaire au précédent à ceci près que la longueur caractéristique de variation de l'aimantation est $l = \delta$, où δ désigne l'épaisseur d'une paroi de domaine. L'énergie de cette configuration est alors $JL^2/\delta + KL^2\delta$. En minimisant ceci par rapport à δ , on trouve que l'épaisseur d'une paroi de domaine est $\delta = \sqrt{J/K}$. L'énergie de cette configuration vaut alors, en négligeant le facteur deux,

$$E_3 \sim \sqrt{JK}L^2.$$

Il faut toutefois noter que cette configuration n'est possible que si la particule est suffisamment grande pour abriter en son sein une paroi de domaine. Autrement dit il faut que

$$L > \delta = \sqrt{J/K}.$$

Configuration la plus stable

La figure 1.1 montre la dépendance de ces trois énergies par rapport à la taille L de la particule. Sur cette figure on voit apparaître trois longueurs caractéristiques :

- $\delta = \sqrt{J/K}$, épaisseur d'une paroi de domaine ;
- $\lambda = \sqrt{J/\mu_0 M_s^2}$ longueur d'échange ;
- $l_c = \lambda^2/\delta$ taille critique pour l'apparition d'une paroi.

Le tableau 1.1 montre les valeurs de δ et de λ pour quelques métaux. Ces longueurs caractéristiques permettent de définir un paramètre sans dimension et spécifique au matériau qu'on appelle *facteur de qualité* et qu'on définit comme

$$Q = \frac{K}{\mu_0 M_s^2} = \left(\frac{\lambda}{\delta}\right)^2 = \left(\frac{l_c}{\lambda}\right)^2.$$

À partir des graphiques précédents, on peut déterminer la configuration la plus stable de la particule en fonction des deux paramètres sans dimensions L/λ et Q . La figure 1.2 est une sorte de « diagramme de phase » montrant la configuration la plus stable en fonction de ces deux paramètres. Le domaine qui nous intéressera dorénavant est celui où la configuration la plus stable est celle en monodomaine.

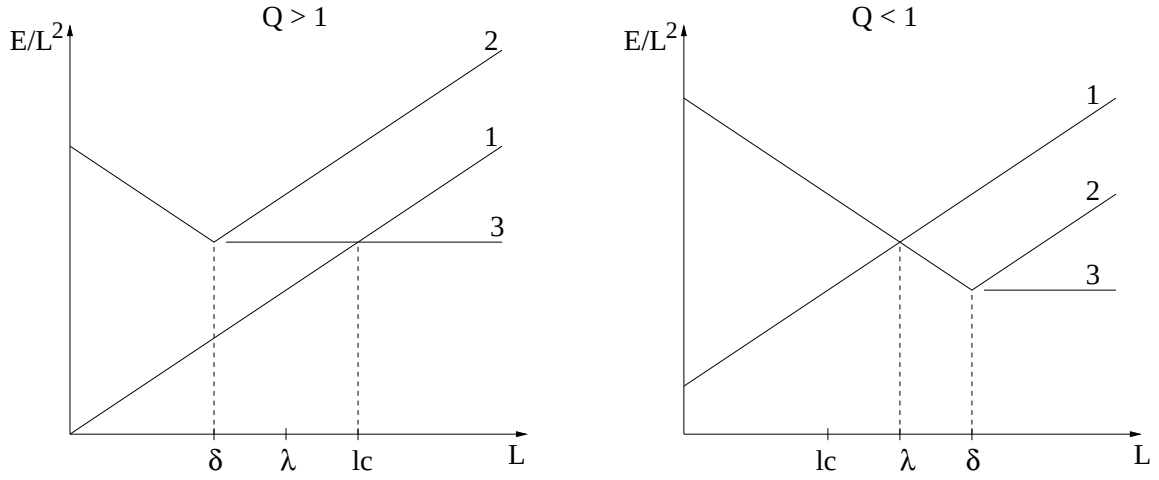


FIG. 1.1 – Énergie des trois configurations considérées en fonction de la taille L de la particule. 1 : configuration uniforme, 2 : configuration non homogène, 3 : paroi de domaine. Deux cas doivent être distingués suivant que $Q > 1$ ou $Q < 1$. Les axes sont en échelle logarithmique.

élément	δ (nm)	λ (nm)
Fe	15	6
Co	5	7
Ni	50	20

TAB. 1.1 – Longueurs magnétiques caractéristiques du fer, du cobalt et du nickel [12]. Les constantes d'échange J étant mal connues, ces valeurs sont entachées d'une incertitude importante.

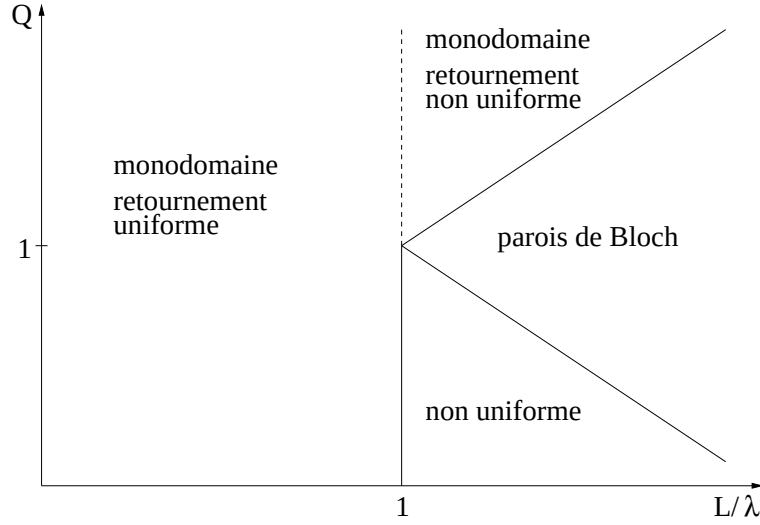


FIG. 1.2 – Diagramme de phase montrant la configuration la plus stable d'une nanoparticule magnétique. Q est le facteur de qualité, λ la longueur d'échange et L la taille de l'échantillon. Les axes sont en échelle logarithmique.

1.1.3 Près du retournement

Supposons donc que nous sommes dans des conditions pour lesquelles la configuration la plus stable en champ nul est celle en monodomaine. Nous allons déstabiliser cette configuration par l'application d'un champ extérieur et nous voudrions savoir si on a des chances de voir un mode de retournement uniforme. On peut imposer par l'imagination à l'aimantation de rester uniforme en toutes circonstances. En appliquant un champ suffisamment intense dans une direction appropriée, il y aura alors une valeur de champ pour laquelle l'orientation globale de l'aimantation devient instable, au profit d'une autre orientation. Cette instabilité est celle du retournement uniforme.

Pour avoir des chances d'observer cette instabilité, il faut qu'il n'y ait aucune autre instabilité correspondant à un mode non uniforme pour des champs inférieurs à celle du mode uniforme. Autrement dit, il faut que la configuration à aimantation uniforme soit encore, juste avant le champ de retournement du mode uniforme, plus stable que toute configuration non uniforme. Nous sommes donc ramenés au problème précédant : on a à comparer les stabilités relatives des différentes configurations, mais cette fois-ci pour un champ appliqué proche de l'instabilité de rotation uniforme.

Toute la discussion précédente considère un échantillon en champ nul. Cependant, il est très facile de la généraliser pour tenir compte d'un champ appliqué. Le

champ extérieur $\mathbf{H}$ apporte à l'énergie une contribution dite *terme Zeeman* qui vaut

$$E_z = \int \mu_0 \mathbf{M} \cdot \mathbf{H}.$$

Ce terme est un peu analogue au terme d'anisotropie en ce sens que l'intégrand dépend uniquement de l'orientation locale de $\mathbf{M}$ et non de ses dérivées (comme c'est le cas de l'échange) ou de l'orientation de $\mathbf{M}$ en d'autres points de l'échantillon (comme c'est le cas du terme démagnétisant). Nous pouvons donc sans scrupules l'intégrer dans notre raisonnement en l'incluant purement et simplement dans une anisotropie *effective*. Évidemment, de ce fait l'anisotropie perd sa forme simple ainsi que sa symétrie centrale. Du coup la constante d'anisotropie K devient, en première approximation, la dérivée seconde de cette nouvelle anisotropie par rapport à la direction de $\mathbf{M}$. L'effet du champ appliqué est alors tout simplement de modifier la valeur de cette nouvelle constante d'anisotropie effective. Ceci correspond à une modification de Q à λ constant et donc à un déplacement vertical dans le diagramme de la figure 1.2.

Nous avons supposé que l'état le plus stable en champ nul est celui en monodomaine et nous voulons savoir s'il le reste juste avant l'instabilité du retournement uniforme. Dans le cas où $L < \lambda$ c'est évidemment le cas puisque cet état est le plus stable indépendamment de Q . Il reste à voir le cas où $L > \lambda$ et $Q > (L/\lambda)^2$ correspondant au triangle en haut à droite du diagramme. Puisque dans cette région $Q > 1$, l'anisotropie domine le terme démagnétisant et la configuration d'équilibre est celle qui minimise l'anisotropie. Le champ appliqué devant diminuer la stabilité de cette configuration, il ne peut que s'opposer à l'effet de l'anisotropie en diminuant la valeur de l'anisotropie effective. Par conséquent, quand on s'approche du champ de retournement uniforme, la valeur du Q effectif diminue et on se déplace *vers le bas* dans le diagramme de la figure 1.2. On voit alors que, tout en restant dans la région $Q > 1$ (hypothèse utilisée ci-dessus), on va passer dans une région où la configuration multi-domaines devient plus favorable. À ce stade on peut s'attendre à voir la nucléation d'une paroi et donc un mode de retournement non uniforme.

Ceci n'est pas une preuve que le mode non uniforme aura effectivement lieu dans cette région puisqu'il n'est pas exclu qu'il existe une barrière de potentiel entre l'état uniforme et celui non uniforme. La présence d'une telle barrière nous semble pourtant très improbable puisque le système dispose d'une infinité de degrés de liberté lui permettant d'aller vers un état non uniforme. C'est pour ça qu'on a indiqué un retournement non uniforme pour cette région de la figure 1.2.

1.1.4 Calculs plus précis

Tous les calculs qui précèdent ne sont valables que de façon qualitative, à des constantes numériques près. Établir un diagramme de phases exact équivaut à

celui de la figure 1.2 n'est pas chose aisée car il dépend de la forme précise de la particule considérée ainsi que de l'expression exacte de son anisotropie cristalline. Il est toutefois possible d'établir un tel diagramme par des simulation numériques. Par exemple, dans [13], W. Rave a établi un tel diagramme pour un cube avec anisotropie uniaxiale. Ses résultats sont comparables à ceux exposés ci-dessus mais la taille limite en dessous de laquelle la particule est toujours monodomaine n'est pas λ mais 25λ .

Les particules utilisées dans nos expériences ont généralement des tailles de quelques λ . Leur comportement statique et dynamique est celui de particules monodomaine. Ceci semble cohérent avec les résultats de Rave qui indiquent que la limite trouvée ici pour le retournement uniforme est très pessimiste.

1.2 Théorie du retournement uniforme

Dans la théorie du retournement uniforme de l'aimantation, on pose comme postulat que l'aimantation de la particule est tout le temps uniforme. Comme on ne s'intéresse qu'aux états d'équilibre et à leurs instabilités, à l'exclusion de ce qui peut se passer *au cours* du retournement, il suffit en fait pour pouvoir appliquer cette théorie d'être dans une situation de retournement uniforme, comme décrit précédemment. La température de la particule étant très en dessous de la température de Curie, on peut aussi considérer que son aimantation $\mathbf{M}$ est à saturation. On a donc $|\mathbf{M}| = M_s$. Puisque $\mathbf{M}$ est uniforme, le moment magnétique total de la particule vaut vM_s , si v est son volume.

Il ne reste alors comme degrés de liberté à considérer que l'orientation du vecteur aimantation qu'on note $\mathbf{r} = \mathbf{M}/M_s$. Le système est alors entièrement décrit par une fonction $V(\mathbf{r})$ qui est le potentiel associé à l'aimantation $M_s\mathbf{r}$. Pour tenir compte du champ extérieur $\mathbf{H}$ appliqué sur la particule, on écrit $V(\mathbf{r}) = V_0(\mathbf{r}) - \mu_0 v \mathbf{H} \cdot \mathbf{M}$ où $V_0(\mathbf{r})$ est le potentiel en champ nul. Par commodité on notera

$$\boxed{V = V_0 - \mathbf{r} \cdot \mathbf{h}}$$

où $\mathbf{h} = \mu_0 v M_s \mathbf{H}$ est l'aimantation réduite (homogène, comme V , à une énergie). Les variations de V_0 par rapport à la direction de $\mathbf{M}$ caractérisent l'anisotropie de la particule. *A priori* il y a au moins deux sources d'anisotropie : l'anisotropie cristalline qui est celle considérée dans la section précédente et l'anisotropie de forme qui est due à la différence d'intensité du champ démagnétisant en fonction de la direction de l'aimantation. À celles-ci on peut rajouter l'anisotropie de surface qui dépend aussi de la forme de la particule. Le potentiel étant invariant par réflexion plane, V est inchangé lorsqu'on change $\mathbf{H}$ en $-\mathbf{H}$ et $\mathbf{M}$ en $-\mathbf{M}$. On a donc $V_0(\mathbf{r}) = V_0(-\mathbf{r})$.

Pour la théorie quasi-statique, on s'intéresse à la limite $T \rightarrow 0$ et on néglige tous les effets quantiques [14]. Le système choisit un état d'équilibre qui est un minimum

local de V . Lorsqu'on fait varier $\mathbf{H}$, le potentiel V se déforme continuellement et le minimum dans lequel se trouve $\mathbf{M}$ se déplace donnant lieu à une rotation continue de l'aimantation. Pour certaines valeurs de $\mathbf{H}$ il se produit une bifurcation : le minimum de V disparaît et l'aimantation varie brusquement pour se stabiliser dans un autre minimum. Ce sont ces « sauts d'aimantation » qui sont responsables du comportement hystérétique de la particule.

Souvent, par souci de simplification, on ne considère qu'un seul degré de liberté de rotation θ . Ceci est bien justifié dans deux cas particuliers :

- D'une part, lorsque l'anisotropie présente un axe difficile suffisamment fort, l'aimantation est contrainte à rester dans un plan. Dans ce cas, l'angle θ est mesuré dans ce plan. On a d'ailleurs le choix de l'origine des angles ;
- D'autre part, lorsque le potentiel V_0 admet un axe de symétrie, on va noter θ l'angle que fait le vecteur aimantation avec cet axe. Dans ce cas on n'a plus le choix de l'origine des angles, mais on sait que V_0 respecte la symétrie $\theta \leftrightarrow -\theta$.

De toutes façons, pour des raisons techniques, le champ $\mathbf{H}$ est souvent contraint à évoluer dans un plan. Dans les deux cas on note Ox la direction $\theta = 0$.

1.3 Modèle de Stoner-Wohlfarth

C'est le modèle le plus simple qu'on puisse construire sur ces considérations [15, 16]. Le potentiel V_0 ne dépend que de la variable θ et est périodique de période π . On prend alors $V_0 = K_1 \sin^2(\theta - \theta_0)$. Dans le cas où V_0 est à symétrie de révolution, la symétrie $\theta \leftrightarrow -\theta$ impose $\theta_0 = 0$. Dans le cas où $\mathbf{M}$ est contraint dans un plan, on peut imposer la même condition par choix de l'origine des angles. Donc $V_0 = K_1 \sin^2 \theta$, d'où

$$\begin{aligned} V &= K_1 \sin^2 \theta - \mathbf{r} \cdot \mathbf{h} \\ &= K_1 \sin^2 \theta - h \cos(\theta - \varphi) \end{aligned}$$

si φ désigne l'angle entre le vecteur $\mathbf{h}$ et l'axe Ox et h le module de $\mathbf{h}$.

Les états d'équilibre sont donnés par

$$\begin{aligned} \frac{\partial V}{\partial \theta} &= 0 \\ 2K_1 \sin \theta \cos \theta + h \sin(\theta - \varphi) &= 0 \\ \sin 2\theta &= -\frac{h}{K_1} \sin(\theta - \varphi). \end{aligned} \tag{1.1}$$

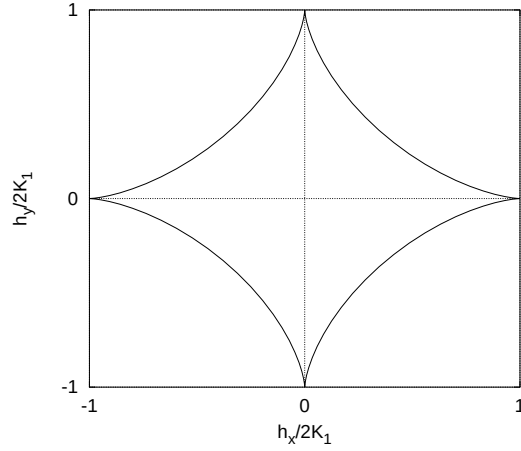


FIG. 1.3 – Courbe d'instabilité pour le modèle de Stoner-Wohlfarth.

Les bifurcations correspondent aux états d'équilibre marginal où

$$\begin{aligned} \frac{\partial^2 V}{\partial \theta^2} &= 0 \\ 2K_1 \cos 2\theta + h \cos(\theta - \varphi) &= 0 \\ \cos 2\theta &= -\frac{h}{2K_1} \cos(\theta - \varphi). \end{aligned} \quad (1.2)$$

Le système constitué des deux équations précédentes permet de trouver les sauts d'aimantation : en divisant (1.1) par (1.2) on obtient

$$\begin{aligned} \tan 2\theta &= 2 \tan(\theta - \varphi) \\ \frac{2 \tan \theta}{1 - \tan^2 \theta} &= 2 \frac{\tan \theta - \tan \varphi}{1 + \tan \theta \tan \varphi} \\ \tan \theta + \tan^2 \theta \tan \varphi &= \tan \theta - \tan \varphi - \tan^3 \theta + \tan^2 \theta \tan \varphi \\ \tan \theta &= -\tan^{1/3} \varphi. \end{aligned}$$

Ce qui permet d'éliminer θ . Ensuite, $(1.1)^2 + (1.2)^2$ donne le champ de retournement en fonction de l'angle du champ appliqué φ :

$$\frac{h}{2K_1} = \frac{1}{(\sin^{2/3} \varphi + \cos^{2/3} \varphi)^{3/2}}.$$

La figure 1.3 montre la courbe d'instabilité en variables adimensionnées, c'est à dire le lieu de $\frac{h}{2K_1}$ lorsque φ parcourt $[0, 2\pi[$.

Ces équations permettent aussi de calculer la forme du cycle d'hystérésis de la particule. Cependant, ces résultats sont moins intéressants pour nous puisque seul le champ de retournement peut être mesuré de manière précise dans notre expérience.

1.4 Anisotropie arbitraire

Le potentiel de Stoner-Wohlfarth est en fait le premier terme d'un développement en série de Fourier de $V(\theta)$. Lorsque ses prédictions ne correspondent pas complètement aux résultats mesurés, on peut envisager de rajouter des termes au développement [17, 18, 19].

On choisit ici d'écrire V sous la forme d'un développement de Fourier plutôt qu'un développement en puissances de $\cos \theta$ et $\sin \theta$ (ce qui est équivalent). En développant à l'ordre 4 on a :

$$V = -\frac{1}{2}A_2 \cos 2\theta - \frac{1}{4}A_4 \cos 4(\theta - \delta) - h \cos(\theta - \varphi).$$

L'absence de « déphasage » dans le premier terme du développement se justifie de même que précédemment. Notons que dans le cas où V est à symétrie de révolution, on a en plus la condition $\delta = 0$.

1.4.1 Notation complexe

On note $V = \Re \mathcal{V}$ avec

$$-\mathcal{V} = h e^{i(\theta - \varphi)} + \frac{1}{2}A_2 e^{2i\theta} + \frac{1}{4}A_4 e^{4i(\theta - \delta)}.$$

Notons $g = h e^{i\varphi}$, $a_2 = A_2$, $a_4 = A_4 e^{4i\delta}$, $z = e^{i\theta}$. Alors

$$-\mathcal{V} = \bar{g}z + \frac{1}{2}a_2 z^2 + \frac{1}{4}\bar{a}_4 z^4.$$

1.4.2 Points d'équilibre

Ils sont donnés par

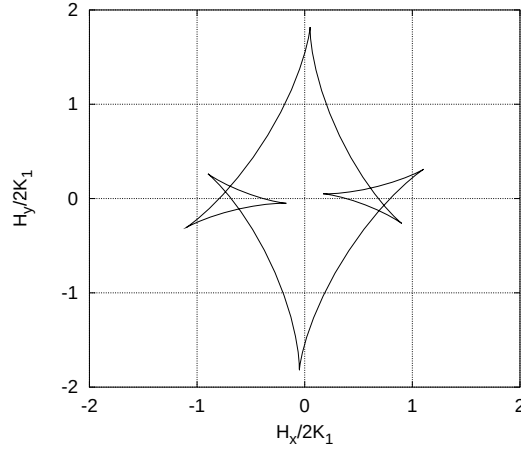
$$\begin{aligned} \frac{dV}{d\theta} &= 0 \\ \Re \frac{d\mathcal{V}}{d\theta} &= 0 \\ \frac{d\mathcal{V}}{d\theta} + \frac{d\bar{\mathcal{V}}}{d\theta} &= 0, \end{aligned}$$

or $\frac{d}{d\theta} = iz \frac{d}{dz}$, donc

$$-\frac{d\mathcal{V}}{d\theta} = i\bar{g}z + ia_2 z^2 + i\bar{a}_4 z^4.$$

La condition d'équilibre est alors :

$$\bar{g}z - g\bar{z} + a_2 z^2 - a_2 \bar{z}^2 + \bar{a}_4 z^4 - a_4 \bar{z}^4 = 0. \quad (1.3)$$

FIG. 1.4 – Courbe d'instabilité pour $a_2 = 1$ et $a_4 = -0,4 + 0,1i$.

1.4.3 Sauts d'aimantation

Ils se produisent lorsqu'on a $\frac{dV}{d\theta} = 0$ et $\frac{d^2V}{d\theta^2} = 0$. Puisque

$$\frac{d^2\mathcal{V}}{d\theta^2} = \bar{g}z + 2a_2z^2 + 4\bar{a}_4z^4,$$

l'équation des points d'inflexion s'écrit :

$$\bar{g}z + g\bar{z} + 2a_2z^2 + 2a_2\bar{z}^2 + 4\bar{a}_4z^4 + 4a_4\bar{z}^4 = 0. \quad (1.4)$$

(1.4) – (1.3) donne

$$2g\bar{z} + a_2z^2 + 3a_2\bar{z}^2 + 3\bar{a}_4z^4 + 5a_4\bar{z}^4 = 0.$$

En remarquant que $\bar{z} = z^{-1}$ et en divisant par $2\bar{z}$ on a :

$$-2g = a_2z^3 + 3a_2z^{-1} + 3\bar{a}_4z^5 + 5a_4z^{-3}.$$

L'équation précédente est une description paramétrée (par z) de la courbe d'instabilité. Elle est directement comparable aux données expérimentales. Elle dépend des seuls paramètres complexe a_2 et a_4 et redonne Stoner-Wohlfarth pour $a_4 = 0$. La figure 1.4 montre ce que devient la courbe d'instabilité pour $a_2 = 1$ et $a_4 = -0,4 + 0,1i$.

1.4.4 Autre description

On aurait pu écrire de façon équivalente

$$V = K_1 \sin^2 \theta' + K_2 \sin^4(\theta' - \delta') - h \cos(\theta' - \varphi').$$

La relation entre les deux descriptions est donnée par :

$$\begin{aligned} A_2^2 &= (K_1 + K_2 \cos 2\delta')^2 + (K_2 \sin 2\delta')^2 \\ A_4 &= -\frac{K_2}{2} \\ \theta &= \theta' - \theta'_0 \\ \delta &= \delta' - \theta'_0 \\ \varphi &= \varphi' - \theta'_0, \end{aligned}$$

où θ'_0 est l'écart entre les deux systèmes de coordonnées et est donné par

$$\begin{aligned} \sin 2\theta'_0 &= \frac{K_2 \sin 2\delta'}{A_2} \\ \cos 2\theta'_0 &= \frac{K_1 + K_2 \cos 2\delta'}{A_2}. \end{aligned}$$

1.5 Bifurcations

Dans le cas de l'anisotropie de Stoner-Wohlfarth, la signification de la courbe critique semble relativement claire. C'est le tracé, en coordonnées polaires, de l'intensité du champ qu'il faut appliquer pour avoir un retournement d'aimantation, et ceci en fonction de la direction de ce champ appliqué. Lorsqu'on rajoute des termes supplémentaires à l'anisotropie, on peut obtenir une courbe qui se recoupe elle-même. Celle-ci indique donc plusieurs valeurs de champ le long de certaines directions. Nous allons, dans cette section, discuter de la signification de la courbe critique dans cette situation.

Pour cette discussion, nous allons considérer la courbe de la figure 1.5. Celle-ci correspond à l'anisotropie

$$V_0 = -\cos^2 \theta - \frac{3}{4} \cos^4 \theta.$$

Supposons qu'on parte d'un état initial où l'aimantation est orientée vers la gauche. Le champ que nous allons appliquer pour retourner l'aimantation va suivre l'un des trajets représentés sur la figure 1.6. Chacun de ces trajets coupe la courbe critique en trois points. Est-ce que cela signifie qu'il y aura trois sauts d'aimantation ? Ou alors, s'il y a un seul saut, laquelle des intersections correspondra au champ du saut ?

Pour répondre à ces questions, nous allons considérer la variation du potentiel $V(\mathbf{r})$ le long de ces trajets. La figure 1.7 montre les potentiels en question. Lorsque le champ est en A , V a deux minima et $\mathbf{r}$ est dans le minimum métastable. Quand le champ augmente, ce minimum devient de moins en moins stable. La figure 1.7 montre le potentiel V au voisinage du minimum métastable pour différentes valeurs

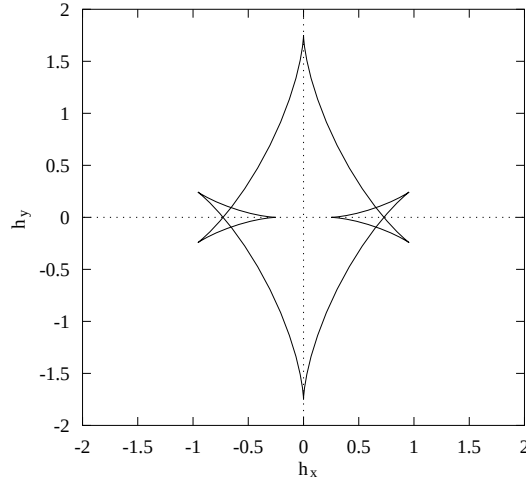


FIG. 1.5 – Courbe critique obtenue pour $V_0 = -\cos^2 \theta - \frac{3}{4} \cos^4 \theta$.

de $\mathbf{h}$ le long des chemins considérés. Cette figure montre aussi les états successifs de $\mathbf{r}$ le long de ces mêmes chemins.

Comme on peut le voir sur la figure 1.7, dans les deux cas $\mathbf{r}$ subira un seul saut. Pour le premier chemin le saut aura lieu en F et pour le second il aura lieu en D . La courbe critique délimite autour de C une région dans laquelle le puits métastable est dédoublé. Pour savoir à quel endroit aura lieu le saut, il faut savoir dans lequel des deux sous-puits se trouve le système. Ceci dépend du chemin emprunté par $\mathbf{h}$ pour aller de A à C . Ça dépend en fait du fait que $\mathbf{h}$ atteigne la région de C en passant au dessus ou en dessous du point I . Dans ce point particulier les deux puits métastables se fondent en un seul de manière symétrique. Ce type de situation est connu sous le nom de *bifurcation supercritique* [20].

Si on lit trop rapidement les figures 1.6 et 1.7, on a l'impression qu'il n'y a jamais de saut d'aimantation en B_1 ou en B_2 . En fait ces sauts sont bel et bien possibles. En suivant le chemin $A \rightarrow B_2 \rightarrow C \rightarrow B_1 \rightarrow A$, on peut voir qu'il y aura un petit saut en B_1 . En fait il est possible d'avoir un saut d'aimantation en tout point de la courbe critique, pourvu que le trajet du champ soit choisi convenablement.

Dans certains cas, le saut d'aimantation est très petit et difficile à mesurer. Dans la région autour du point C , lorsqu'on se rapproche de I les deux puits métastable deviennent très proches. Un saut d'aimantation entre ces deux puits donnera lieu à une variation d'aimantation $\Delta \mathbf{M}$ très faible qui à son tour induira une très faible variation $\Delta \Phi$ du flux à travers le SQUID. La méthode indirecte qui a été mise au point pour mesurer ces tout petits sauts d'aimantation est décrite dans le chapitre 6.

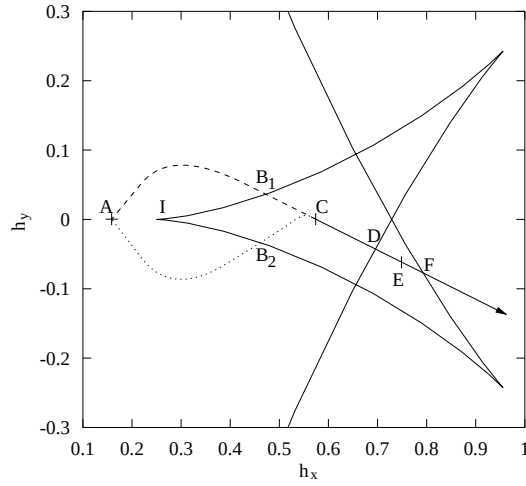


FIG. 1.6 – Agrandissement de la partie droite de la figure 1.5. Les deux lignes partant de A sont les trajets du champ discutés dans le texte.

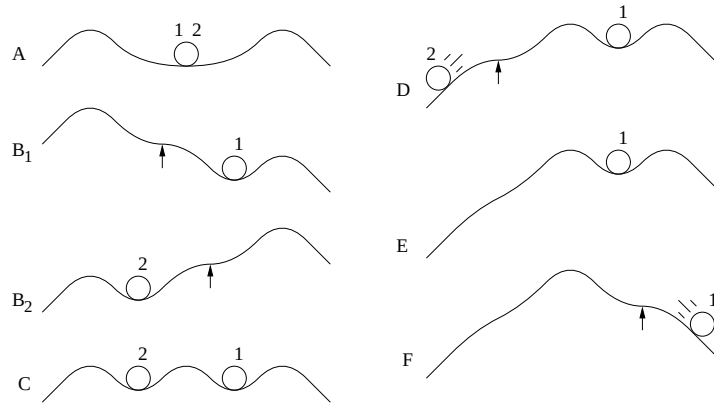


FIG. 1.7 – Allure du potentiel $V(\mathbf{r})$ en plusieurs points le long des chemins de la figure 1.6. Sur ces figures l'abscisse représente la direction de $\mathbf{r}$, les balles sont les états par lesquels passe $\mathbf{r}$ et les flèches indiquent les disparitions d'un minimum local du potentiel.

1.6 En trois dimensions

Cette théorie quasi-statique a été généralisée à trois dimensions [21]. La principale différence avec le cas bidimensionnel est que les sauts d'aimantation ne surviennent pas quand un minimum du potentiel rencontre un maximum mais quand il rencontre un col.

Comme la direction $\mathbf{r}$ de l'aimantation a deux degrés de liberté, la condition d'équilibre est l'annulation des *deux* dérivées premières du potentiel par rapport à des coordonnées curvilignes sur la sphère unité. Dans [21] les coordonnées utilisées sont les coordonnées sphériques θ et φ . Dans le chapitre 3 nous montrons une méthode plus symétrique basée sur un système de coordonnées local.

Avec deux degrés de liberté, un point d'équilibre peut être un maximum, un minimum ou un col du potentiel. Ces trois cas se distinguent par la forme quadratique des dérivées secondes. Cette forme a ses deux valeurs propres positives dans un minimum et négatives dans un maximum. Dans un col les deux valeurs propres ont des signes opposés. Pour écrire la condition d'équilibre marginal, il faut trouver le moment où un minimum rejoint un col. À ce moment la forme quadratique a une valeur propre positive et une valeur propre nulle. En écrivant l'annulation du déterminant on a une solution de trop : le moment où un maximum rejoint un col. Il faut donc écrire aussi que la trace doit être positive.

Le lieu des champs critiques est alors une surface dans l'espace des champs. Cette surface est constituée de différentes nappes qui se rejoignent sur des fronces, lieux des points super-critiques. La surface critique présente une certaine analogie avec le diagramme de phases magnétique du système : ce dernier est constitué de nappes correspondant aux transitions du premier ordre qui se terminent sur des lignes correspondant à des transitions du second ordre [22]. Ces lignes critiques du diagramme de phases sont les mêmes que les fronces de notre surface critique.

Chapitre 2

Théorie dynamique

Dans le chapitre précédent nous nous intéressions à la limite $T \rightarrow 0$. Les expériences ayant lieu à température non nulle, il convient de voir quel peut être l'effet de la température sur le retournement de l'aimantation.

Lorsque le champ appliqué est inférieur au champ critique H_c , il existe un état métastable séparé de l'état le plus stable par une barrière d'énergie. Le champ H_c est celui qui annule la hauteur de cette barrière et qui permet donc le retournement de l'aimantation à température nulle. Si le système se trouve à température non nulle, il est capable de sauter par activation thermique une barrière d'énergie de l'ordre de quelques $k_B T$. Le retournement de l'aimantation aura donc lieu pour un champ inférieur à H_c , lorsque la hauteur de la barrière sera d'un ordre de grandeur compatible avec l'activation thermique.

Une autre différence importante est que le retournement devient un événement stochastique. Pour chaque hauteur de la barrière d'énergie, le système a une certaine probabilité par unité de temps de se retourner. On va montrer dans ce chapitre comment on peut calculer ces lois de probabilité. Une bonne partie de ce qui est ici est une relecture d'un article de synthèse de Coffey [23]. Pour une présentation générale de la théorie de l'activation thermique, il convient de lire [24].

2.1 Modèle de Néel-Brown

Nous considérons ici une particule dont le potentiel présente deux puits séparés par une barrière d'énergie. Ce modèle s'intéresse à l'activation thermique qui permet au système de franchir la barrière entre les deux puits. Celle-ci est décrite à l'aide d'un temps caractéristique τ appelé *temps de relaxation*. En général on suppose que la hauteur de la barrière est grande devant $k_B T$ de sorte que le système est presque tout le temps au fond de l'un des puits et il passe très peu de temps au voisinage du sommet de la barrière. Ceci est connu sous le nom d'*approximation de la barrière*

haute. Deux cas se présentent suivant la forme de cette barrière.

Barrière peu ou pas asymétrique : Le franchissement de la barrière est possible dans les deux sens. On considère alors un grand nombre de copies identiques du système et on s'intéresse aux populations dans chacun des deux puits. Le temps de relaxation τ est alors le temps qu'il faut à ces populations pour atteindre une valeur d'équilibre.

Barrière très asymétrique : Dans ce cas, le franchissement de la barrière n'est possible que dans un sens : du puits le moins profond (puits métastable) au plus profond (puits stable). Le franchissement dans l'autre sens demanderait un temps généralement supérieur à l'âge de l'univers et en tout cas toujours incompatible avec la patience d'un expérimentateur. Le temps de relaxation est alors le temps moyen qu'il faut au système pour franchir de façon irréversible la barrière. Plus précisément, si on sait que le système est dans le puits métastable à l'origine des temps, la probabilité pour qu'il y soit encore à la date t est

$$P(t) = e^{-t/\tau}.$$

On suppose ou on montre que le temps de relaxation suit une loi d'Arrhénius :

$$\tau = \tau_0 e^{E/k_B T},$$

où τ_0 est une grandeur qui varie peu et E est la hauteur de la barrière d'énergie qu'on franchit. En fait, comme on verra par la suite, ce cas peut être considéré comme un cas particulier du précédent.

2.1.1 Modèle des orientations discrètes

En faisant l'approximation de la barrière haute, on peut supposer que $\mathbf{r}$ suit un chemin d'orientations stables le long des minima locaux du potentiel. On a alors à traiter un processus Markovien où la distribution W est remplacée par n_i : le nombre de particules d'orientation i . Si on a un grand nombre de particules $n = \sum n_i$ qui n'interagissent pas, on a l'équation maîtresse :

$$\dot{n}_i = \sum_{j \neq i} (\nu_{j,i} n_j - \nu_{i,j} n_i), \quad (2.1)$$

où $\nu_{i,j}$ est la probabilité par unité de temps qu'une particule d'orientation i transite à l'orientation j . $\nu_{i,j}$ dépend de la température T , de la constante d'anisotropie K , et du champ appliqué H .

Si on se place dans le cas uniaxial, on a deux orientations possibles et l'équation 2.1 devient ¹ :

$$\dot{n}_1 = -\dot{n}_2 = \nu_{2,1}n_2 - \nu_{1,2}n_1.$$

En remplaçant n_2 par $n - n_1$ on trouve l'équation d'évolution pour n_1 :

$$\dot{n}_1 = \nu_{2,1}n - (\nu_{1,2} + \nu_{2,1})n_1.$$

On aboutit finalement au temps de relaxation de Néel :

$$\tau = (\nu_{1,2} + \nu_{2,1})^{-1}.$$

On pose généralement $\nu_{i,j} = \nu_{i,j}^0 e^{(V_0 - V_i)/k_B T}$, où V_0 et V_i sont les potentiels au sommet de la barrière et au fond des puits et où $\nu_{i,j}^0$ varie lentement avec la température et peut donc être considéré en première approximation comme constant. On a alors, quand les hauteurs des barrières sont égales :

$$\tau = C e^{(V_0 - V_1)/k_B T},$$

où C est une constante.

2.1.2 Modèle de Néel

Ce modèle a été développé par Louis Néel en 1949 pour étudier la relaxation de l'aimantation des poudres ferromagnétiques en grains fins [25]. Il suppose que chaque grain est une particule monodomaine avec anisotropie uniaxiale. Une barrière d'énergie de hauteur Kv , où v est le volume de la particule, s'oppose au retournement de son aimantation. La température est supposée être suffisamment faible devant Kv/k_B pour que la densité de probabilité de présence de l'aimantation soit concentrée au voisinage des directions faciles. Ceci correspond à l'approximation des orientations discrètes où la relaxation est décrite par les probabilités P et $1 - P$ d'avoir l'aimantation au voisinage de chacune de ces deux directions.

Néel a supposé que l'aimantation se trouve initialement dans une direction connue correspondant à la saturation de l'échantillon ($P(t = 0) = 1$). À $t = 0$, le champ responsable de cette saturation est coupé et cette probabilité évolue exponentiellement suivant

$$P(t) = \frac{1 + e^{-t/\tau}}{2}.$$

Le temps de relaxation τ a été estimé en considérant que, lorsque l'aimantation atteint par activation thermique le plan de difficile aimantation, l'anisotropie n'exerce

¹On a choisi arbitrairement 1 comme orientation positive.

plus aucun couple sur elle. Le seul couple Γ qui s'exerce est alors dû aux vibrations thermiques de la particule, via le couplage magnétoélastique. Ce couple fait précesser l'aimantation suivant la loi

$$\frac{d}{dt} \frac{v\mathbf{M}}{\gamma} = \Gamma,$$

où γ est le rapport gyromagnétique (rapport de la charge de l'électron à sa masse) et $v\mathbf{M}/\gamma$ le moment cinétique associé à l'aimantation de la particule.

En ne considérant que les modes de vibration uniformes comme réservoir thermique, Néel a estimé la moyenne quadratique de ce couple :

$$\sqrt{\langle \Gamma^2 \rangle} = 3\lambda_s \sqrt{Gvk_BT},$$

où λ_s est la constante magnétoélastique et G le module de cisaillement. Il en a déduit que le temps de relaxation vaut

$$\tau = \tau_N e^{Kv/k_BT}$$

où le préfacteur τ_N , appelé temps de relaxation de Néel, est donné par :

$$\tau_N = \frac{M_s}{6K\gamma|\lambda_s|} \sqrt{\frac{\pi k_BT}{2Gv}}.$$

τ_N a une valeur typique de 10^{-10} s.

2.1.3 Modèle de Brown

Brown a étendu ce modèle avec un traitement plus rigoureux de la densité de probabilité de présence de l'aimantation. Le modèle de Brown [26, 27, 28, 29] est plus général car indépendant de la nature physique du couplage entre le bain thermique et le moment magnétique de la particule. Ce couplage peut en effet faire intervenir des courants de Foucault dans un conducteur et non seulement l'interaction magnétoélastique. Son modèle est aussi généralisé à n'importe quelle forme de l'anisotropie effective.

Dans un monodomaine, l'anisotropie magnétique V_0 et le champ appliqué $\mathbf{H}$ se résument à un champ effectif

$$\mathbf{H}_{ef} = \mathbf{H} - \frac{1}{\mu_0 v M_s} \frac{dV_0}{d\mathbf{r}},$$

où $\mathbf{r} = \mathbf{M}/M_s$ est la direction de l'aimantation. En absence d'amortissement, ce champ ferait précesser $\mathbf{r}$ suivant l'équation gyromagnétique

$$\dot{\mathbf{r}} = \mu_0 \gamma \mathbf{r} \times \mathbf{H}_{ef}.$$

Pour tenir compte des phénomènes dissipatifs tels que les courants de Foucault et le couplage magnétoélastique avec le réseau, Gilbert a introduit un terme d'atténuation effective ayant pour effet d'aligner $\mathbf{r}$ parallèlement à $\mathbf{H}_{ef}$. L'équation ci dessus devient

$$\dot{\mathbf{r}} = \mu_0 \gamma \mathbf{r} \times (\mathbf{H}_{ef} - \eta M_s \dot{\mathbf{r}}), \quad (2.2)$$

où η est une constante d'atténuation phénoménologique. Après quelques manipulations vectorielles, on trouve la forme explicite de l'équation de Gilbert :

$$\dot{\mathbf{r}} = a M_s \mathbf{r} \times \mathbf{H}_{ef} + b M_s (\mathbf{r} \times \mathbf{H}_{ef}) \times \mathbf{r},$$

où

$$a = \frac{\mu_0 \gamma}{(1 + \alpha^2) M_s}, \quad b = \frac{\alpha \mu_0 \gamma}{(1 + \alpha^2) M_s} \quad \text{et} \quad \alpha = \eta \mu_0 \gamma M_s = \frac{b}{a}.$$

Si on suppose $\alpha^2 = 0$, l'équation 2.2 se simplifie et on retrouve l'équation de Landau-Lifshitz. Le terme contenant a est appelé terme gyromagnétique. Celui contenant b , terme d'alignement. En introduisant le potentiel

$$V = V_0 - \mu_0 v \mathbf{M} \cdot \mathbf{H},$$

le champ effectif s'exprime par

$$\mathbf{H}_{ef} = -\frac{1}{v \mu_0 M_s} \frac{\partial V}{\partial \mathbf{r}}.$$

On remplace $\mathbf{H}_{ef}$ dans 2.2 et comme $\frac{\partial V}{\partial \mathbf{r}} \perp \mathbf{r} \Rightarrow \mathbf{r} \cdot \frac{\partial V}{\partial \mathbf{r}} = 0$, on trouve :

$$\dot{\mathbf{r}} = -a \mathbf{r} \times \frac{\partial V}{\partial \mathbf{r}} - b \frac{\partial V}{\partial \mathbf{r}}.$$

Brown a considéré les orientations des aimantations individuelles du système de particules magnétiques comme un courant de points représentatifs se déplaçant sur la sphère unité avec une densité $W(\mathbf{r}, t)$ et un courant de densité $\mathbf{J}(\mathbf{r}, t)$. Comme il y a conservation du nombre de ces points représentatifs, on a l'équation de continuité

$$\dot{W} = -\frac{\partial}{\partial \mathbf{r}} \cdot \mathbf{J}.$$

Pour tenir compte de la force thermique aléatoire qui disperse les concentrations de ces points représentatifs autour des points d'orientations d'aimantation stable, il a postulé l'existence d'un terme de diffusion de la forme $-k' \frac{\partial W}{\partial \mathbf{r}}$, où $k' > 0$ et est constant pour une température donnée. On a alors

$$\mathbf{J} = W \dot{\mathbf{r}} - k' \frac{\partial W}{\partial \mathbf{r}}.$$

On peut alors achever le calcul en substituant $\dot{\mathbf{r}}$ par son expression, et en développant les doubles produits vectoriels (on trouvera le détail du calcul dans [30]). On obtient alors l'équation de Fokker-Planck pour la distribution de l'orientation des aimantations :

$$\dot{W} = a\mathbf{r} \cdot \left(\frac{\partial V}{\partial \mathbf{r}} \times \frac{\partial W}{\partial \mathbf{r}} \right) + b \frac{\partial}{\partial \mathbf{r}} \cdot \left(W \frac{\partial V}{\partial \mathbf{r}} \right) + k' \frac{\partial^2 W}{\partial \mathbf{r}^2}.$$

A l'équilibre thermique, on a $\dot{W} = 0$, et W est la distribution d'équilibre $W_0 = Ae^{-V/k_B T}$. Si on substitue W par son expression, on a alors

$$-\frac{a}{k_B T} e^{-V/k_B T} \mathbf{r} \cdot \left(\frac{\partial V}{\partial \mathbf{r}} \times \frac{\partial V}{\partial \mathbf{r}} \right) + (b - k'/k_B T) \frac{\partial}{\partial \mathbf{r}} \cdot \left(e^{-V/k_B T} \frac{\partial V}{\partial \mathbf{r}} \right) = 0.$$

Le terme gyroscopique disparaît, et comme le dernier terme est non nul, on a $k' = bk_B T$, donc

$$\dot{W} = a\mathbf{r} \cdot \left(\frac{\partial V}{\partial \mathbf{r}} \times \frac{\partial W}{\partial \mathbf{r}} \right) + b \frac{\partial}{\partial \mathbf{r}} \cdot \left(W \frac{\partial V}{\partial \mathbf{r}} \right) + bk_B T \frac{\partial^2 W}{\partial \mathbf{r}^2}.$$

2.1.4 Approximation de la barrière haute

Si on a un puits de potentiel asymétrique avec deux minima en $\mathbf{n}_1$ et $\mathbf{n}_2$ et un point de selle en $\mathbf{n}_0$ les séparant, et dans le cas où $(V(\mathbf{n}_0) - V(\mathbf{n}_i))/k_B T \gg 1$, la relaxation de l'aimantation vers le puits stable suit une loi de probabilité exponentielle. On a reporté dans la littérature des lois de relaxation non exponentielles, mais elles correspondent soient à une situation où la barrière est basse [31], soit à l'existence d'un très grand nombre de barrières séparant l'état métastable de l'état stable [32], soit à des systèmes avec anisotropie infinie (modèle d'Ising) [33, 34]. Dans le cas qui nous intéresse, l'évolution de l'équation de Fokker-Planck est dominée par le temps de relaxation le plus long, c'est à dire par la plus petite valeur propre non nulle de l'opérateur de Fokker-Planck. Coffey *et al.* ont montré que ce temps de relaxation peut se calculer en fonction de différentes courbures du paysage de potentiel.

Amortissement moyen et fort

Dans le cas de l'amortissement critique défini par $\alpha(V(\mathbf{n}_0) - V(\mathbf{n}_i))/k_B T > 1$ (voir [35]), le temps de relaxation est donné par

$$\frac{1}{\tau} = \frac{\Omega_0}{2\pi\omega_0} \left(\omega_1 e^{-(V(\mathbf{n}_0)-V(\mathbf{n}_1))/k_B T} + \omega_2 e^{-(V(\mathbf{n}_0)-V(\mathbf{n}_2))/k_B T} \right),$$

où ω_0 et Ω_0 sont les fréquences angulaires au point de selle et au point de selle atténué :

$$\omega_0 = \frac{\gamma}{vM_s} \sqrt{-c_1^{(0)} c_2^{(0)}},$$

$$\Omega_0 = \frac{\gamma}{2vM_s} \frac{\alpha}{1 + \alpha^2} \left(-c_1^{(0)} - c_2^{(0)} + \sqrt{(c_2^{(0)} - c_1^{(0)})^2 - 4\alpha^{-2} c_1^{(0)} c_2^{(0)}} \right),$$

avec ω_i la pulsation :

$$\omega_i = \frac{\gamma}{vM_s} \sqrt{c_1^{(i)} c_2^{(i)}}.$$

$c_1^{(i)}$ et $c_2^{(i)}$ sont les courbures du potentiel aux points $\mathbf{n}_i$, et peuvent être déterminée par la mesure tridimensionnelle de la surface critique du champ de retournement en appliquant le calcul de Thiaville [36, 21].

Dans l'expression de τ on peut distinguer deux termes qui correspondent au franchissement de la barrière dans un sens et dans l'autre. Dans nos expériences, nous avons toujours une barrière très asymétrique de sorte que l'un des temps de passage peut être considéré comme infini. Nous retrouvons alors la loi d'Arrhénus

$$\frac{1}{\tau} = \frac{1}{\tau_0} e^{-E/k_B T},$$

où $E = V(\mathbf{n}_0) - V(\mathbf{n}_1)$ est la hauteur de la barrière d'énergie et

$$\frac{1}{\tau_0} = \frac{1}{2\pi} \frac{\Omega_0 \omega_1}{\omega_0}.$$

Soit, en remplaçant Ω_0 et ω_i par leur valeur,

$$\frac{1}{\tau_0} = \frac{\gamma}{4\pi v M_s} \frac{\alpha}{1 + \alpha^2} \frac{\sqrt{c_1^{(1)} c_2^{(1)}}}{\sqrt{-c_1^{(0)} c_2^{(0)}}} \left(-c_1^{(0)} - c_2^{(0)} + \sqrt{(c_2^{(0)} - c_1^{(0)})^2 - 4\alpha^{-2} c_1^{(0)} c_2^{(0)}} \right). \quad (2.3)$$

Amortissement faible

Si nous sommes dans un régime d'amortissement sous-critique, c'est à dire dans le cas où $\alpha(V(\mathbf{n}_0) - V(\mathbf{n}_i))/k_B T < 1$, l'expression du temps de relaxation est aussi une loi d'Arrhénus avec

$$\frac{1}{\tau_0} = \frac{\alpha}{2\pi} \omega_1 \frac{E}{k_B T}.$$

Soit, en remplaçant ω_1 par sa valeur,

$$\frac{1}{\tau_0} = \frac{\alpha \gamma}{2\pi v M_s} \frac{E}{k_B T} \sqrt{c_1^{(1)} c_2^{(1)}}. \quad (2.4)$$

2.2 Reformulation de Kurkijärvi

Nous avons vu que, dans le cas d'une barrière très asymétrique, la probabilité pour que le système soit dans l'état métastable à la date t sachant qu'il y est à l'origine des temps est

$$P(t) = e^{-t/\tau}, \quad (2.5)$$

où τ suit la loi d'Arrhénus

$$\tau = \tau_0 e^{E/k_B T} \quad (2.6)$$

et on peut souvent négliger les variations de τ_0 . Il convient cependant de noter que la loi de probabilité ci-dessus suppose que τ_0 , et donc la hauteur E de la barrière, sont constants. Cette description n'est donc pas adaptée à l'interprétation de nos mesures de champs de retournement puisque celles-ci se font avec un champ, et donc une hauteur de barrière, qui varient continuellement. Il est cependant possible, comme l'a montré Kurkijärvi [37], d'exprimer cette loi de probabilité dans le cas d'un champ appliqué variant à vitesse constante. Plus que la loi de probabilité, c'est en fait la position de son maximum qui nous intéresse, c'est à dire la valeur du champ de retournement le plus probable. C'est le calcul de ce champ qui est exposé ci-dessous.

La première chose à faire est de remarquer que l'équation 2.5 est en fait la forme intégrée d'une loi plus générale. Cette loi dit que, dans un système *sans mémoire*, la probabilité pour que la barrière soit franchie entre les dates t et $t+dt$, sachant qu'elle n'a pas été franchie à la date t , ne dépend que de dt . Cette probabilité conditionnelle s'écrit donc

$$-\frac{dP}{P} = \frac{dt}{\tau}. \quad (2.7)$$

On voit alors que l'équation 2.5 n'est que la forme intégrée de celle-ci pour τ constant.

Il faut ensuite exprimer la hauteur E de la barrière d'énergie en fonction du champ appliqué. En développant cette dépendance autour du champ critique, on trouve

$$E = E_0(\delta H)^{\alpha'}, \quad (2.8)$$

où $\delta H = |\mathbf{H} - \mathbf{H}_c|$ est l'écart au champ critique. L'exposant α' , qui correspond au premier terme non nul du développement, vaut toujours 3/2 sauf dans quelques cas pathologiques [38, 39].

Dans une expérience de mesure de champ de retournement, on part d'un état initial à $t = 0$ où l'aimantation est dans son état métastable et on balaye le champ jusqu'à ce que l'aimantation se renverse. Dans ce cas, δH , E et τ sont toutes les

trois des fonctions décroissantes du temps. La densité de probabilité de retournement $-\frac{dP}{dt}$ est maximale quand

$$\frac{d^2P}{dt^2} = 0. \quad (2.9)$$

Or, en dérivant 2.7 on trouve

$$\frac{d^2P}{dt^2} = \frac{P}{\tau^2} \left(1 + \frac{d\tau}{dt} \right).$$

L'équation 2.9 s'écrit donc

$$\frac{d\tau}{dt} = -1.$$

Si on remplace τ par son expression dans (2.6) et qu'on néglige les variations de τ_0 par rapport au champ et à la température, l'équation ci-dessus devient

$$\frac{dE}{dt} = -k_B T \frac{1}{\tau_0} e^{-E/k_B T}.$$

En utilisant l'expression de E donnée par l'équation 2.8, on trouve

$$(\delta H)^{\alpha'} = \frac{k_B T}{E_0} \ln \frac{k_B T}{E_0 (\delta H)^{\alpha'-1} \alpha' v \tau_0},$$

où $v = -\frac{d(\delta H)}{dt}$ est la vitesse de balayage du champ.

Cette relation a été généralisée par Garg [40] pour tenir compte des variations de τ_0 en fonction du champ appliqué. Cependant, ces variations étant bien en dessous de notre précision de mesure, nous n'utiliserons pas son développement.

Chapitre 3

Calculs numériques

Afin de comprendre les effets de l'anisotropie sur le comportement physique des nanoparticules, j'ai développé le programme *astroide* qui permet de calculer des grandeurs d'intérêt physique à partir de l'expression du potentiel d'anisotropie. Ce programme permet donc d'exploiter les ajustements du potentiel. Les grandeurs calculées sont le lieu des champs de retournement (la surface entière en 3D ou une coupe plane) et les préfacteurs géométriques qui interviennent dans le calcul du temps caractéristique τ_0 .

Ce chapitre se propose de décrire les méthodes mathématiques mises en œuvre dans le programme. La méthode générale est celle décrite par Thiaville dans [21] mais l'approche technique est légèrement différente. Alors que dans [21] il est largement fait usage des coordonnées sphériques, ici nous n'utilisons que des coordonnées cartésiennes. Ce chapitre peut être considéré, en quelque sorte, comme un guide de lecture des sources, ou comme les commentaires qu'il aurait été trop difficile d'y inclure. C'est un peu dans cet état d'esprit qu'il a été écrit. Les menus détails algébriques ont été mis dans l'annexe A.

3.1 Position du problème

Une particule monodomaine d'aimantation $M_s \mathbf{r}$, soumise à un champ extérieur $\mathbf{H}$ sent le potentiel

$$V(\mathbf{r}, \mathbf{h}) = V_0(\mathbf{r}) - \mathbf{r} \cdot \mathbf{h} \quad (3.1)$$

où $\mathbf{h} = v\mu_0 M_s \mathbf{H}$ est un champ réduit et v est le volume de la particule. Le retournement de l'aimantation à température nulle se produit quand un minimum de $V(\mathbf{r})$ rejoint un col. Nous voulons déterminer le champ de retournement pour ce cas limite. À température très basse mais non nulle, le retournement a lieu très près de ce champ limite. Ce retournement est un phénomène stochastique qui dépend de la hauteur et la forme de la barrière d'énergie.

3.2 Plongement dans l'espace 3D

Le potentiel $V(\mathbf{r}, \mathbf{h})$ n'est défini que pour $\mathbf{r}$ appartenant à la sphère unité. Cependant, pour les calculs il est plus pratique d'utiliser des dérivées par rapport aux coordonnées cartésiennes de l'espace plutôt que par rapport à un système de coordonnées curvilignes sur la sphère. Nous allons donc exprimer les dérivées par rapport aux coordonnées curvilignes en fonction de celles par rapport aux coordonnées cartésiennes.

On choisit un point $\mathbf{r}_0$ sur la sphère unité et on regarde s'il existe un $\mathbf{h}$ pour lequel il y a retournement d'aimantation à température nulle. On choisit un repère orthonormé (x, y, z) tel que le point $\mathbf{r}_0$ choisi ait pour coordonnées $(0, 0, 1)$.

Au voisinage de $\mathbf{r}_0$, on repère les points de la sphère par les coordonnées (u, v) définies par

$$\begin{aligned} x &= u \\ y &= v \\ z &= \sqrt{1 - u^2 - v^2} = 1 - \frac{u^2 + v^2}{2} + O((u^2 + v^2)^2) \end{aligned}$$

Les potentiels V_0 et V définis sur la sphère sont étendus à tout l'espace. À partir des dérivées sur tout l'espace on retrouve les dérivées sur la sphère par composition des dérivées. Les dérivées premières valent

$$\begin{aligned} \frac{\partial}{\partial u} &= \frac{\partial}{\partial x} - x \frac{\partial}{\partial z} + O((u^2 + v^2)^{3/2}) \\ \frac{\partial}{\partial v} &= \frac{\partial}{\partial y} - y \frac{\partial}{\partial z} + O((u^2 + v^2)^{3/2}) \end{aligned}$$

les dérivées secondes,

$$\begin{aligned} \frac{\partial^2}{\partial u^2} &= \frac{\partial^2}{\partial x^2} - \frac{\partial}{\partial z} - 2x \frac{\partial^2}{\partial x \partial z} + O(u^2 + v^2) \\ \frac{\partial^2}{\partial u \partial v} &= \frac{\partial^2}{\partial x \partial y} - y \frac{\partial^2}{\partial x \partial z} - x \frac{\partial^2}{\partial y \partial z} + O(u^2 + v^2) \\ \frac{\partial^2}{\partial v^2} &= \frac{\partial^2}{\partial y^2} - \frac{\partial}{\partial z} - 2y \frac{\partial^2}{\partial y \partial z} + O(u^2 + v^2) \end{aligned}$$

et les dérivées troisièmes

$$\begin{aligned}\frac{\partial^3}{\partial u^3} &= \frac{\partial^3}{\partial x^3} - 3\frac{\partial^2}{\partial x\partial z} + O((u^2 + v^2)^{1/2}) \\ \frac{\partial^3}{\partial u^2\partial v} &= \frac{\partial^3}{\partial x^2\partial y} - \frac{\partial^2}{\partial y\partial z} + O((u^2 + v^2)^{1/2}) \\ \frac{\partial^3}{\partial u\partial v^2} &= \frac{\partial^3}{\partial x\partial y^2} - \frac{\partial^2}{\partial x\partial z} + O((u^2 + v^2)^{1/2}) \\ \frac{\partial^3}{\partial v^3} &= \frac{\partial^3}{\partial y^3} - 3\frac{\partial^2}{\partial y\partial z} + O((u^2 + v^2)^{1/2})\end{aligned}$$

On note $g(\mathbf{r})$, $q(\mathbf{r})$ et $c(\mathbf{r})$ les termes d'ordre 1, 2 et 3 du développement en série de $V_0(x, y, z)$ autour de $\mathbf{r}_0$:

$$g_i = \frac{\partial V_0}{\partial i}(\mathbf{r}_0) \quad q_{ij} = \frac{1}{2!} \frac{\partial^2 V_0}{\partial i \partial j}(\mathbf{r}_0) \quad c_{ijk} = \frac{1}{3!} \frac{\partial^3 V_0}{\partial i \partial j \partial k}(\mathbf{r}_0)$$

pour $i, j \in \{x, y, z\}$. D'après l'équation 3.1, les dérivées de V en $\mathbf{r}_0$ ($x = y = 0$) sont :

$$\begin{aligned}\frac{\partial V}{\partial u}(\mathbf{r}_0) &= g_x - h_x \\ \frac{\partial V}{\partial v}(\mathbf{r}_0) &= g_y - h_y \\ \frac{\partial^2 V}{\partial u^2}(\mathbf{r}_0) &= 2q_{xx} - g_z + h_z \\ \frac{\partial^2 V}{\partial u\partial v}(\mathbf{r}_0) &= 2q_{xy} \\ \frac{\partial^2 V}{\partial v^2}(\mathbf{r}_0) &= 2q_{yy} - g_z + h_z \\ \frac{\partial^3 V}{\partial u^3}(\mathbf{r}_0) &= 6c_{xxx} - 6q_{xz} \\ \frac{\partial^3 V}{\partial u^2\partial v}(\mathbf{r}_0) &= 6c_{xxy} - 2q_{yz} \\ \frac{\partial^3 V}{\partial u\partial v^2}(\mathbf{r}_0) &= 6c_{xyy} - 2q_{xz} \\ \frac{\partial^3 V}{\partial v^3}(\mathbf{r}_0) &= 6c_{yyy} - 6q_{yz}\end{aligned}$$

3.3 Sauts d'aimantation

Les sauts d'aimantation à température nulle ont lieu lorsqu'un minimum du potentiel fusionne avec un col. Le point commun entre un minimum et un col est

que en ces points les dérivées premières sont nulles. Ce qui les différencie est la forme quadratique correspondant au terme d'ordre deux du développement limité en ce point. À un minimum, les deux valeurs propres de cette forme quadratique sont positives. À un point col l'une est positive et l'autre est négative. Quand les deux se rejoignent on a une valeur propre nulle et l'autre positive. Autrement dit, le déterminant est nul et la trace est positive. Nos conditions s'écrivent donc

$$\begin{aligned} \frac{\partial V}{\partial u} &= 0 \\ \frac{\partial V}{\partial v} &= 0 \\ \begin{vmatrix} \frac{\partial^2 V}{\partial u^2} & \frac{\partial^2 V}{\partial u \partial v} \\ \frac{\partial^2 V}{\partial v \partial u} & \frac{\partial^2 V}{\partial v^2} \end{vmatrix} &= 0 \\ \frac{\partial^2 V}{\partial u^2} + \frac{\partial^2 V}{\partial v^2} &\geq 0. \end{aligned}$$

En utilisant les notations g et q pour la partie linéaire et quadratique de V_0 , on a

$$\begin{aligned} g_x - h_x &= 0 \\ g_y - h_y &= 0 \\ (2q_{xx} - g_z + h_z)(2q_{yy} - g_z + h_z) - 4q_{xy}^2 &= 0 \\ 2q_{xx} + 2q_{yy} - 2g_z + 2h_z &\geq 0, \end{aligned}$$

ce qui nous donne comme solution

$$\begin{aligned} h_x &= g_x \\ h_y &= g_y \\ h_z &= -(q_{xx} + q_{yy} - g_z) + \sqrt{(q_{xx} - q_{yy})^2 + 4q_{xy}^2}. \end{aligned}$$

3.4 Forme du potentiel au voisinage du saut

Nous nous plaçons au voisinage du champ $\mathbf{h}$ qui produit un saut d'aimantation en $\mathbf{r}_0$ à température nulle. Il s'agit de déterminer la forme locale du potentiel. Celui-ci se présente sous la forme d'un minimum proche d'un col. Pour préciser cette forme, on développe $V(\mathbf{r})$ par rapport aux coordonnées (u', v') sur la sphère :

$$V(\mathbf{r}, \mathbf{h}) = a_{02}v'^2 + a_{30}u'^3 + a_{21}u'^2v' + a_{12}u'v'^2 + a_{03}v'^3 + O((u'^2 + v'^2)^2)$$

où $a_{02} > 0$ et $a_{30} > 0$. Le choix des coordonnées u' et v' est analogue au choix de u et v à ceci près que les directions des axes ont été choisies de sorte que

- la valeur propre nulle de la forme quadratique de V est suivant la direction de u' ;
- le saut d'aimantation se fait suivant les v' décroissants.

Nous ajoutons maintenant au champ $\mathbf{h}$ un petit champ $\delta\mathbf{h}$ perpendiculaire à l'astroïde et dirigé vers l'intérieur, afin de créer une petite barrière de potentiel. Cette direction choisie pour le champ est celle qui augmente le plus vite la hauteur de la barrière de potentiel. C'est donc la direction de l'axe u' dans le sens de u' croissants. Le nouveau potentiel s'écrit alors

$$V(\mathbf{r}, \mathbf{h} + \delta\mathbf{h}) = V(\mathbf{r}, \mathbf{h}) - \delta h u'$$

3.4.1 Hauteur de la barrière d'énergie

La barrière d'énergie se mesure entre le minimum et le col. Ces deux points correspondent à l'annulation de la partie linéaire de $V(\mathbf{r})$. On résout donc

$$\frac{\partial V}{\partial u'} = \frac{\partial V}{\partial v'} = 0$$

soit

$$\begin{aligned} 2a_{02}v' + a_{21}u'^2 + 2a_{12}u'v' + 3a_{03}v'^2 + O((u'^2 + v'^2)^{3/2}) &= 0 \\ 3a_{30}u'^2 + 2a_{21}u'v' + a_{12}v'^2 - \delta h + O((u'^2 + v'^2)^{3/2}) &= 0 \end{aligned}$$

La première de ces deux équations nous montre que

$$v' = -\frac{a_{21}}{2a_{02}}u'^2 + O(u'^3)$$

ce qui veut dire que le minimum et le col sont proches de l'axe $v' = 0$ et donc u' et v' ne sont pas du même ordre en δh . On en déduit la position du minimum et du col à l'ordre 1/2 en δh :

$$\begin{aligned} u' &= \pm \left(\frac{\delta h}{3a_{30}} \right)^{1/2} + O(\delta h^{3/4}) \\ v' &= O(\delta h) \end{aligned}$$

Pour ces deux points, un développement à l'ordre 3/2 en δh donne

$$\begin{aligned} V(\mathbf{r}) &= a_{30}u'^3 - \delta h u' + O(\delta h^2) \\ &= \pm \frac{2}{3\sqrt{3}} \frac{\delta h^{3/2}}{a_{30}^{1/2}} + O(\delta h^2) \end{aligned}$$

D'où la hauteur de la barrière d'énergie

$$\boxed{\Delta V = \frac{4}{3\sqrt{3}} \frac{\delta h^{3/2}}{a_{30}^{1/2}} + O(\delta h^2)}$$

3.4.2 Courbures du potentiel

Dans l'étude de la dynamique de l'aimantation à température non nulle, il faut connaître les courbures du potentiel au minimum et au col. Celles-ci sont données par les dérivées secondes de V qui sont

$$\begin{aligned}\frac{\partial^2 V}{\partial u'^2} &= 6a_{30}u' + 2a_{21}v' + O(u'^2) \\ \frac{\partial^2 V}{\partial u'\partial v'} &= 2a_{21}u' + 2a_{12}v' + O(u'^2) \\ \frac{\partial^2 V}{\partial v'^2} &= 2a_{02} + 2a_{12}u' + 6a_{03}v' + O(u'^2)\end{aligned}$$

Les courbures qui nous intéressent correspondent aux valeurs propres de la matrice formée par ces dérivées. Elles se calculent à partir de la trace et du déterminant de cette matrice. Les termes non diagonaux de cette matrice n'interviennent que dans le déterminant, mais à un ordre plus grand que les termes diagonaux. Par conséquent, au premier ordre les valeurs propres sont les mêmes si on annule les termes non diagonaux. Ce qui revient à dire que les courbures principales sont

$$\begin{aligned}\frac{1}{2} \frac{\partial^2 V}{\partial u'^2} &= \pm \sqrt{3a_{30}} \delta h^{1/2} + O(\delta h) \\ \frac{1}{2} \frac{\partial^2 V}{\partial v'^2} &= a_{02} + O(\delta h^{1/2})\end{aligned}$$

On remarque que la courbure dans la direction v' est, au premier ordre, indépendante de δh et du point considéré. Celle dans la direction u' est d'ordre $\delta h^{1/2}$ et elle a des valeurs opposées au minimum et au col.

3.4.3 Temps caractéristique τ_0

Nous avons vu dans la section 2.1.4 l'expression générale de ce temps caractéristique en fonction de certaines courbures du potentiel au minimum et au col. On se contente maintenant de reprendre ces expressions avec le changement de notations suivant :

$$\begin{aligned}c_i^{(0)} &\rightarrow c_i \\ c_i^{(1)} &\rightarrow c'_i.\end{aligned}$$

Amortissement moyen et fort

L'équation 2.3 s'écrit maintenant :

$$\frac{1}{\tau_0} = \frac{\gamma}{4\pi v M_s} \frac{\alpha}{1 + \alpha^2} \frac{\sqrt{c'_1 c'_2}}{\sqrt{-c_1 c_2}} \left(-c_1 - c_2 + \sqrt{(c_2 - c_1)^2 - 4\alpha^{-2} c_1 c_2} \right).$$

Sachant que

$$\begin{aligned} c_1 &= a_{02} + O(\delta h^{1/2}) \\ c_2 &= -\sqrt{3a_{30}}\delta h^{1/2} + O(\delta h) \\ c'_1 &= c_1 \\ c'_2 &= -c_2, \end{aligned}$$

et en tenant compte que c_2 et c'_2 sont très petits par rapport à c_1 et c'_1 , cette expression de τ_0 se développe en

$$\begin{aligned} \frac{1}{\tau_0} &\approx \frac{\gamma}{4\pi v M_s} \frac{\alpha}{1 + \alpha^2} (1 + \alpha^{-2}) c'_2 \\ \boxed{\frac{1}{\tau_0} &\approx \frac{\gamma}{4\pi \alpha v M_s} \sqrt{3a_{30}} \delta h^{1/2}.} \end{aligned}$$

Amortissement faible

En reprenant de la même manière l'équation 2.4, on a

$$\begin{aligned} \frac{1}{\tau_0} &= \frac{\alpha \gamma}{2\pi v M_s} \frac{E}{k_B T} \sqrt{c'_1 c'_2} \\ \boxed{\frac{1}{\tau_0} &\approx \frac{\alpha \gamma}{2\pi v M_s} \frac{E}{k_B T} \sqrt{a_{02} \sqrt{3a_{30}} \delta h^{1/4}}.} \end{aligned}$$

Deuxième partie

Techniques expérimentales

Chapitre 4

Présentation de l'expérience

Le but de l'expérience est de mesurer l'aimantation d'une particule magnétique en fonction du champ $\mathbf{H}$ appliqué et de la température T . On utilise pour cela un SQUID (*Superconducting QUantum Interference Device* = interféromètre quantique supraconducteur) qui est un détecteur sensible au flux magnétique qui le traverse [41]. La particule étant de très petite taille, il est nécessaire de la poser directement sur le SQUID pour avoir un signal mesurable. Cette technique consistant à obtenir un couplage direct entre le SQUID et le système à mesurer a été développée pour la mesure des courants permanents dans des conducteurs mésoscopiques [42]. La figure 4.1 montre qu'une telle disposition permet au champ engendré par la particule de traverser la boucle du SQUID.

Cette méthode impose essentiellement deux contraintes. D'une part, une seule composante de l'aimantation peut être mesurée. D'autre part, le champ appliqué doit se trouver dans le plan du SQUID pour ne pas être mesuré par ce dernier.

La figure 4.2 montre le schéma général de l'expérience. Au centre du dispositif se trouve le SQUID avec la particule à étudier. Autour il y a trois bobines (ou paires de

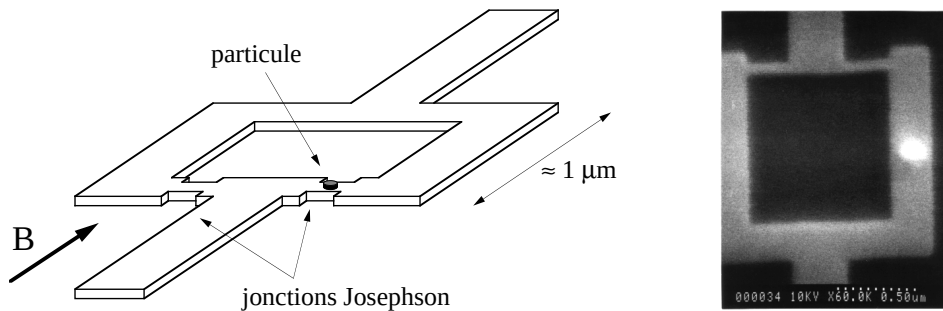


FIG. 4.1 – Particule magnétique posée sur un micro-SQUID. L'image de droite a été obtenue au microscope électronique à balayage.

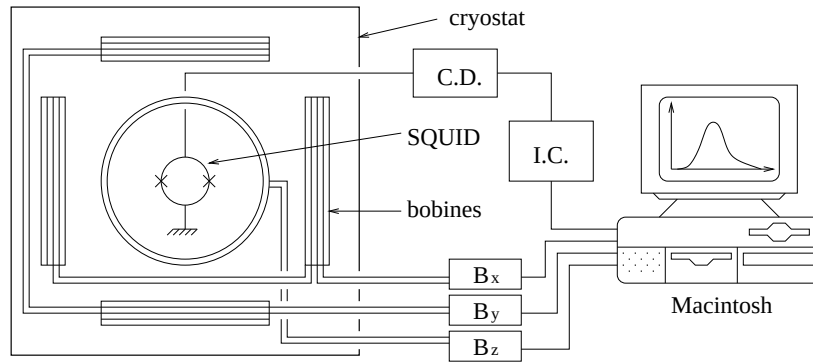


FIG. 4.2 – Schéma général de l'expérience.

bobines) qui permettent de créer un champ $\mathbf{H}$ dans les trois directions de l'espace. Ceci ne contredit pas ce qui a été dit plus haut. Il est nécessaire de pouvoir appliquer un champ dans les trois dimensions pour deux raisons :

- D'une part, il est difficile d'aligner deux bobines (ou paires de bobines) avec leurs axes parfaitement parallèles au plan du SQUID. Or, la tolérance dans cet alignement est très faible, on a donc besoin de pouvoir appliquer un champ qui compense le défaut d'alignement ;
- D'autre part, la méthode de mesure repose sur une boucle de rétroaction. On mesure en fait le champ qu'il faut appliquer perpendiculairement au SQUID pour maintenir le flux qui le traverse constant et égal à une valeur de consigne.

L'ensemble SQUID + bobines se trouve à l'intérieur d'un cryostat à dilution qui permet d'atteindre des températures de l'ordre de 100 mK. À l'extérieur du cryostat on trouve trois générateurs de courant (notés B_x , B_y , B_z) qui alimentent les bobines de champ. On trouve aussi un circuit de détection (C.D.) qui permet d'effectuer la mesure du flux traversant le SQUID. Enfin, une interface de contrôle (I.C.) permet de gérer l'ensemble. Elle est elle même pilotée par un Macintosh qui constitue l'interface utilisateur [43, 44].

4.1 Cryogénie

Les basses températures nécessaires à ces expériences sont obtenues à l'aide d'un cryostat à dilution. Celui qu'on utilise a été mis au point au CRTBT et il a l'originalité d'être monté à l'envers par rapport aux modèles classiques : l'échantillon se trouve au sommet de l'appareil et est donc facilement accessible. Pour cette raison, ce type de cryostat est appelé « Sionludi ». La figure 4.3 est un schéma de principe du Sionludi.

Une circulation d' ^4He assure un premier refroidissement jusqu'à environ 4,2 K.

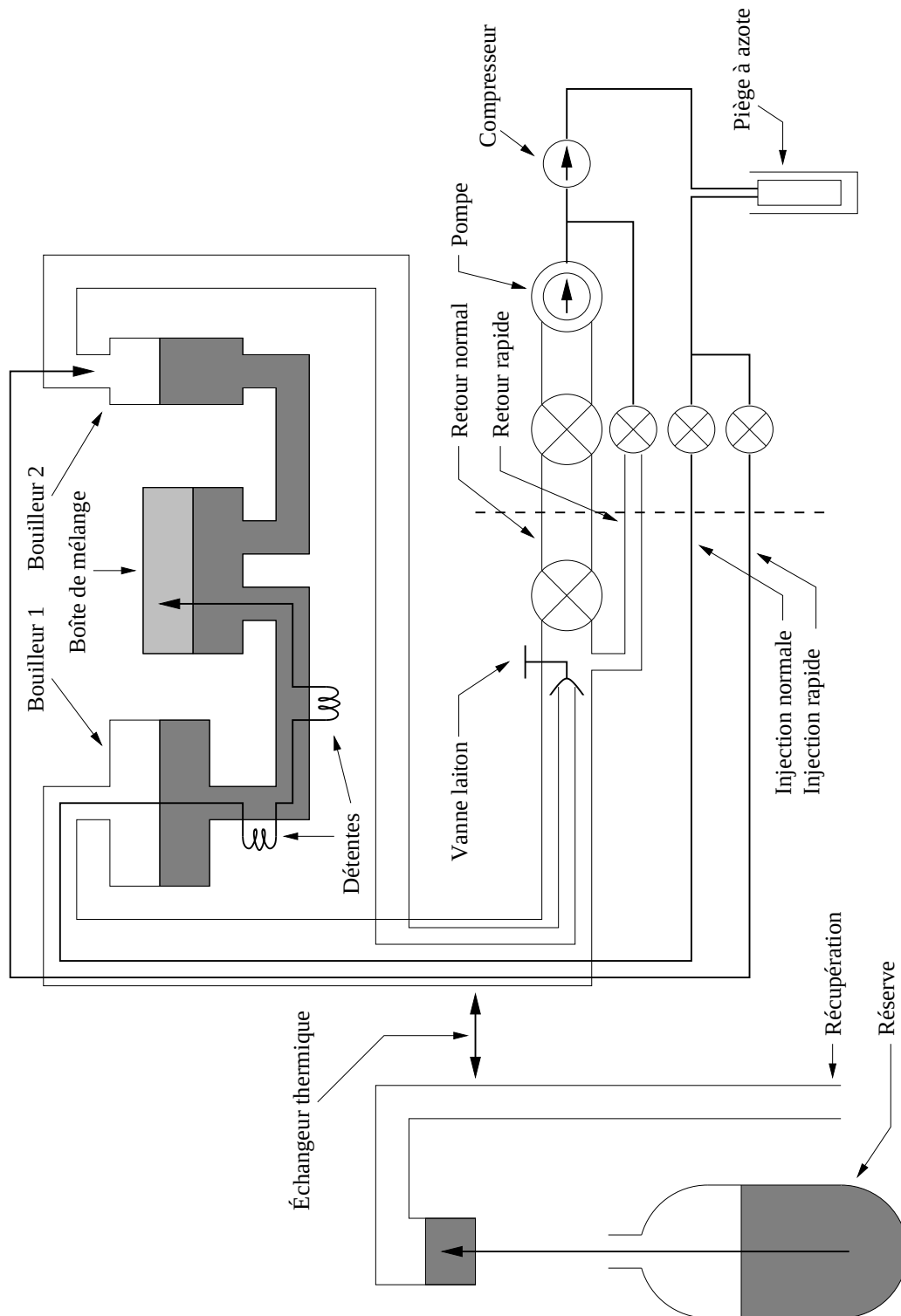


FIG. 4.3 – Schéma de principe du Sionludi.

Les températures inférieures sont obtenues par dilution d'un mélange d' ^3He et de ^4He .

En dessous de 0,87 K le mélange se sépare spontanément en deux phases : la phase concentrée en ^3He surnage au dessus de la phase diluée. L' ^3He étant plus volatile que son isotope lourd, en abaissant la pression au dessus de la phase diluée, celle-ci est appauvrie en ^3He . Il y a alors diffusion de l' ^3He de la phase concentrée vers la phase diluée. Ceci entraîne une diminution de la température [45].

La diffusion entre phases se produit dans la partie la plus froide du cryostat : la « boîte de mélange ». L'échantillon est maintenu en bon contact thermique avec celle-ci. Autour de cet ensemble sont disposés des écrans destinés à limiter les transferts thermiques par radiation. Le premier de ces écrans est maintenu à 4,2 K par une circulation continue d' ^4He . Les autres se thermalisent spontanément à des températures intermédiaires entre 4,2 K et la température ambiante. La température atteinte par la boîte de mélange est comprise typiquement entre 30 et 50 mK sans régulation. Les températures supérieures sont obtenues par un chauffage régulé fourni par une résistance.

Notons qu'un vide poussé règne à l'intérieur du cryostat. Il est produit initialement par une turbopompe. Cependant, lorsque le cryostat est refroidi, la pression tombe bien en dessous de ce que peut produire la turbopompe. C'est l'absorption des molécules résiduelles de gaz par les parois froides du cryostat qui est responsable de cette diminution de pression. Ceci est très utile pour éviter les fuites thermiques par convection.

Chapitre 5

Chaîne de mesure

Cette section concerne plus particulièrement la nouvelle version de l'expérience. Il s'agit d'un système qui est encore en cours de développement et qui devrait nous permettre d'atteindre des performances supérieures à celles accessibles actuellement.

5.1 Micro-SQUID

5.1.1 Une théorie simpliste

Ce qui suit n'a pas la prétention d'être une description précise de ce qui se passe dans un SQUID. C'est seulement un modèle simple, voire simpliste, qui permet de donner une description qualitativement juste de ce qu'on observe [46].

Les SQUID utilisés sont dits hystérétiques puisque leur caractéristique courant-tension présente de l'hystérésis. Pour un flux appliqué à travers la boucle constant, le SQUID est supraconducteur tant que le courant qui le traverse reste inférieur à une valeur critique I_c . Lorsque cette valeur est dépassée, le SQUID est chauffé par effet Joule et par conséquent, il reste à l'état normal (résistif) pour des valeurs du courant inférieures à I_c . C'est seulement lorsque le courant devient inférieur à un autre courant critique $I_c^{\min}$ que l'effet Joule est suffisamment faible pour permettre la transition inverse. La figure 5.1 montre schématiquement la caractéristique courant-tension d'un tel SQUID.

Le SQUID est composé de deux branches notées 1 et 2. Chacune est caractérisée par un coefficient d'induction L_i tel qu'un courant I_i parcourant cette branche engendre un flux $L_i I_i$ à travers la boucle du SQUID. Le courant total qui traverse le SQUID est $I = I_1 + I_2$.

Dans chaque branche se trouve une constriction qu'on appelle « micro-pont » dans laquelle la densité de courant est plus forte que dans le reste du SQUID. Le micro-pont est caractérisé par un courant critique $I_i^{\max}$. Si le courant dans la branche

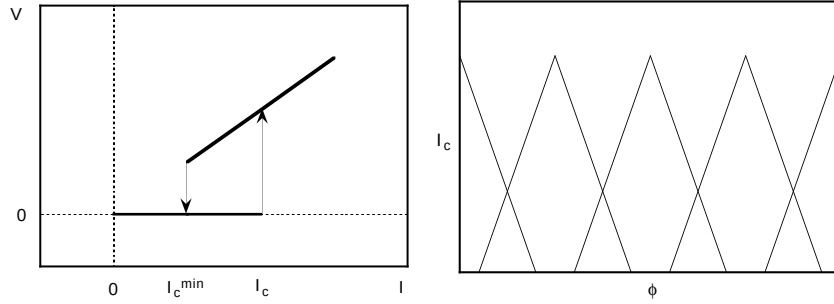


FIG. 5.1 – Caractéristique courant-tension d'un SQUID continu hystérétique et dépendance du courant critique en fonction du flux extérieur appliqué. Les différentes branches du graphe de droite correspondent à des nombres différents de quanta de flux piégés par la boucle.

i dépasse cette valeur critique, la densité de courant dans le micro-pont est telle que celui-ci transite de l'état supraconducteur vers l'état normal. Dès qu'un micro-pont a transité, la chaleur dégagée par effet Joule provoque la transition en avalanche de la totalité du SQUID.

On suppose que le flux magnétique total à travers la boucle du SQUID est quantifié. Il vaut

$$n\Phi_0 = \Phi + L_1 I_1 + L_2 I_2$$

où n est un entier arbitraire, $\Phi_0 = \frac{h}{2e}$ est le quantum de flux et Φ est le flux extérieur appliqué.

La condition pour que la jonction 1 ne transite pas est $I_1 < I_1^{\max}$. En remarquant que

$$n\Phi_0 - \Phi = L_1 I_1 + L_2 I_2 = L_2 I + (L_1 - L_2) I_1$$

cette condition se réécrit :

$$I < \frac{n\Phi_0 - (L_1 - L_2) I_1^{\max}}{L_1} - \frac{\Phi}{L_1}.$$

Le SQUID transite dès que le courant critique est atteint sur l'une des deux branches. Le courant critique du SQUID est alors :

$$I_c = \min \left(\frac{n\Phi_0 - (L_1 - L_2) I_1^{\max}}{L_1} - \frac{\Phi}{L_1}; \frac{n\Phi_0 + (L_1 - L_2) I_2^{\max}}{L_2} - \frac{\Phi}{L_2} \right).$$

Lorsque le SQUID est symétrique on a $L_1 = -L_2$ et $I_1^{\max} = I_2^{\max}$. Alors $I_c(n, \Phi) = I_c(-n, -\Phi)$. Dans la pratique il est plus facile de réaliser $L_1 = -L_2$ que $I_1^{\max} =$

$I_2^{\max}$ puisque la première relation concerne la forme globale du SQUID alors que la deuxième concerne la petite région du micro-pont.

La figure 5.1 montre les variations de I_c en fonction de Φ . On voit qu'il existe plusieurs courants critiques possibles pour un même flux appliqué selon le nombre de quanta de flux piégés par la boucle. En fait les valeurs de n qui ont une probabilité significative sont celles qui donnent des petites valeurs de $|n\Phi_0 - \Phi|$. Ceci correspond aux courants critiques les plus élevés.

Une théorie plus élaborée faisant intervenir des jonctions tunnel donne des résultats sensiblement différents le courant critique est de la forme

$$I_c = I_0 \left| \sin \left(\pi \frac{\Phi}{\Phi_0} \right) \right|.$$

Il est possible que dans une jonction à micro-pont il y ait un petit courant tunnel. Dans ce cas, on s'attend à voir un courant critique ayant une forme intermédiaire entre les deux expressions ci-dessus. C'est exactement ce qu'on observe dans la pratique.

5.1.2 Comportement observé

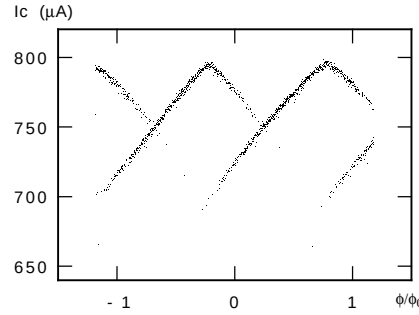


FIG. 5.2 – Courant critique mesuré en fonction du flux appliqué perpendiculairement à la boucle du SQUID. Chaque point représente une mesure individuelle.

La figure 5.2 montre la dépendance du courant critique en fonction du flux appliqué à travers la boucle. On peut remarquer que pour certaines valeurs du flux il y a deux courants critiques possibles suivant le nombre de quanta de flux piégés dans la boucle au moment du refroidissement du SQUID. Ceci est bien en accord avec le modèle simple qui précède. Ce qui est nouveau, c'est que la probabilité de piéger un certain nombre de quanta devient très faible lorsque ce nombre conduirait à un très petit courant critique. Cette propriété permet d'avoir des intervalles de flux pour lesquels une seule branche a une probabilité significative. L'allure générale de la courbe est assez proche de celle prédite plus haut, à ceci près que les arcs ne

sont pas constitués de lignes droites mais sont légèrement incurvés. cette courbure ressemble à une réminiscence de la théorie qui prévoit une dépendance de I_c en $|\sin(\pi\Phi/\Phi_0)|$, mais on n'arrive pas vraiment à l'expliquer.

On peut remarquer que les traces de la figure 5.2 présentent une certaine largeur. Ceci témoigne d'une dispersion des mesures qui est une propriété intrinsèque du SQUID (le SQUID est aussi un système stochastique). La figure 5.3 montre l'histogramme des courants critiques pour une valeur de flux où le nombre de quanta de flux piégés est constant. Autrement cet histogramme serait bimodal.

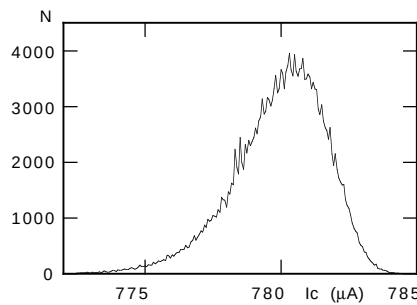


FIG. 5.3 – Histogramme des mesures du courant critique pour un champ appliqué constant. La largeur des canaux de l'histogramme est de 62,5 nA.

5.2 Circuit de détection

Ce circuit, essentiellement analogique, est conçu pour déterminer de manière rapide et précise la valeur du courant critique du SQUID [47]. Il est constitué de deux parties. La partie « source de courant » engendre une rampe de courant comme celle représentée dans la figure 5.4. Le courant est d'abord amené rapidement à un palier qui doit être aussi proche que possible du courant critique. Une fois l'électronique stabilisée, le courant est augmenté lentement jusqu'à atteindre la transition.

La deuxième partie de ce circuit se charge de détecter la transition du SQUID. Celle-ci se manifeste sous la forme d'une brusque augmentation de la tension à ses bornes. Après filtrage passe-bas et amplification, c'est une impulsion de tension qu'il s'agit de détecter. La tension ainsi filtrée est comparée à un *seuil de décision* réglable, ceci permet de distinguer l'impulsion due à la transition du bruit de fond.

Lorsque la transition est détectée, le générateur de courant est court-circuité très rapidement afin de limiter au minimum le temps pendant lequel le SQUID dissipe de la chaleur. En même temps, le générateur de courant est réamorcé et l'interface de contrôle reçoit un signal qui lui indique l'instant où a eu lieu la transition.

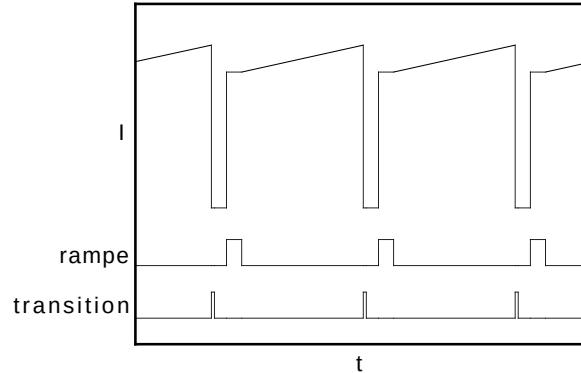


FIG. 5.4 – Chronogramme des signaux utilisés dans l'électronique de détection. I est le courant envoyé sur le SQUID, $rampe$ déclenche un cycle de mesure, $transition$ signale que la transition du SQUID a été détectée.

5.3 Interface de contrôle

C'est un circuit numérique construit autour d'un PLD (*Programmable Logical Device* = circuit logique programmable). Il est chargé de piloter le circuit de détection et de fournir les résultats des mesures à l'ordinateur sous forme numérique.

Dans un premier temps, l'interface se charge de transmettre au circuit de détection les différents réglages nécessaires : hauteur du palier de courant, pente de la rampe, gain de l'amplification dans la chaîne de détection et seuil de décision.

Cette interface envoie un signal logique $rampe$ au circuit de détection pour déclencher un cycle de mesure. Le front montant de $rampe$ déclenche la montée au palier du générateur de courant alors que le front descendant déclenche la rampe de courant.

Le circuit de détection envoie à son tour un signal $transition$ lorsque la transition du SQUID a été détectée. L'interface de contrôle chronomètre le temps écoulé entre le front descendant de $rampe$ et la transition du SQUID, ce qui permet de remonter au courant critique. Le chronométrage se fait à l'aide d'une horloge à 40 MHz.

5.4 Carte d'acquisition

Afin d'interfacer toute cette électronique, une carte d'acquisition a été mise au point. Elle a été réalisée lors d'une collaboration avec le CRTBT. Cette carte est une carte de communication à vocation assez générale et actuellement plusieurs exemplaires sont en service dans diverses expériences, aussi bien au laboratoire Louis Néel qu'au CRTBT. Afin de permettre à d'autres chercheurs de l'utiliser efficacement dans des expériences différentes de la notre, nous en avons rédigé un mode d'emploi

qui se trouve en annexe B. Ce qui suit est une description sommaire de la carte ainsi que de sa place dans notre expérience.

Le rôle de la carte d'acquisition est de faire communiquer l'ordinateur de commande avec le reste de l'électronique à très grande vitesse grâce à une API de très bas niveau. Cette carte est au format PCI et devrait donc être utilisable sur tout ordinateur disposant de ce type de bus actuellement très courant. Cependant, elle n'a été testée à l'heure actuelle que sur Macintosh. La communication se fait dans les deux sens à travers deux bus différents. Le bus d'entrée est large de 18 bits et passe par une mémoire tampon pouvant avoir jusqu'à 32 kmots de profondeur. Cette mémoire est réalisée à l'aide de circuits FIFO intégrés. Le bus de sortie ne fait que 9 bits de large et est non tamponné. L'écriture dans les lignes de sortie est donc immédiate. L'API est on ne peut plus simple. Pour envoyer une donnée en sortie, il suffit de l'écrire dans le registre de sortie. Pour lire une donnée en entrée, il suffit de lire le registre d'entrée. Ces deux registres, ainsi qu'un registre d'état permettant de contrôler le fonctionnement de la carte, sont projetés en mémoire.

Lors des tests de vitesse que nous avons effectués, nous avons pu faire des acquisitions à des cadences allant jusqu'à 1 MHz sans perte de données. À des cadences plus faibles, la profondeur du tampon permet d'assurer une certaine immunité par rapport aux irrégularités du multi-tâche de MacOS. La capacité d'avoir des cadences élevées est utile pour certaines expériences menées au CRTBT. Dans notre expérience, la cadence d'acquisition ne dépasse pas en général 3 kHz. Ainsi, avec un tampon de seulement 8 kmots nous avons un système d'acquisition complètement fiable.

Deux exemplaires de cette carte sont utilisés dans notre expérience. Le premier sert à communiquer avec l'interface de contrôle du micro-SQUID. Il transmet à l'interface les paramètres de fonctionnement et reçoit de sa part les dates de transition du SQUID. Le deuxième exemplaire sert uniquement en sortie pour commander les alimentations de courant des bobines de champ. Pour des raisons économiques il a subi une ablation des mémoires FIFO. Utiliser une telle carte pour transmettre quelques bits à des convertisseurs toutes les 1,5 ms peut sembler une débauche de moyens, mais en fait une fois la carte au point son prix de revient unitaire est très faible.

5.5 Logiciel

Le logiciel de commande qui tourne sur l'ordinateur est probablement la pièce la plus complexe du dispositif expérimental. Il peut être décomposé schématiquement en trois parties : une routine temps réel, une API de bas niveau pour la contrôler et une interface utilisateur. L'ensemble est écrit en C au dessus de la bibliothèque *Manip*. Cette dernière a été écrite au CRTBT et fournit une gestion de haut niveau

des fenêtres de dialogue ainsi que quelques fonctions et structures pour manipuler des données expérimentales. La stabilité de *Manip* étant perfectible, il faut parfois en contourner quelques bogues.

5.5.1 Routine temps réel

Il s'agit d'une fonction qui est appelée périodiquement pendant le fonctionnement de l'expérience. Son appel est déclenché par l'interruption d'horloge, ce qui assure une bonne régularité. La période de l'appel est réglable et on utilise en général 1,5 ms. Lors de chaque appel, son rôle est d'incrémenter le champ (si on veut faire une rampe de champ), de récolter les mesures de courant critique (si le SQUID est dans ce mode de mesure), d'appliquer une rétroaction (si on veut en utiliser une), etc. En fait, cette fonction est en charge de toutes les tâches qui doivent être réalisées en continu pendant l'expérience. Lors de chaque appel, elle réalise une étape élémentaire de chacune de ces tâches. Le fait que l'appel de cette fonction se fasse par une interruption rend son exécution complètement indépendante de la séquence d'exécution du reste du programme. Ceci implémente une sorte de parallélisme.

5.5.2 API de la routine temps réel

Le fonctionnement de la routine temps réel est contrôlé par une API bien définie qui la présente comme un robot programmable. Son utilisation se fait en trois étapes : d'abord on établit le programme, ensuite on met le robot en marche, enfin on l'arrête si le programme ne prévoit pas un arrêt explicite. La récupération des données acquises par le robot peut se faire pendant son fonctionnement ou bien après son arrêt. Le robot comporte un tampon logiciel qui garde les données en attendant qu'elles soient lues par le reste du programme.

Le programme du robot est un tableau de structures donc chacune décrit une étape élémentaire du protocole de mesure. Chaque étape peut impliquer la réalisation d'une rampe de champ magnétique, la mesure périodique du courant critique du SQUID, l'application d'une rétroaction, etc. Les étapes sont de durée variable, potentiellement infinie.

5.5.3 Interface utilisateur

Il s'agit en fait de tout le reste du programme, ce qui est finalement un peu plus qu'une simple interface utilisateur. Plusieurs fenêtres permettent de régler les paramètres des différents modes de mesure. D'autres fenêtres permettent de visualiser et enregistrer les résultats. Cette partie est de loin la plus importante en ce qui concerne le nombre de lignes de code. C'est à ce niveau là que sont définis les différents protocoles de mesure, via des programmations différentes du robot.

Chapitre 6

Protocole de mesure

6.1 Commande du champ

6.1.1 Schéma de principe

L'ordinateur de commande envoie, par le biais d'une carte d'entrées-sorties, une consigne de champ magnétique que nous noterons C . Cette consigne est un signal numérique en binaire 16 bits signé. Un convertisseur numérique analogique transforme ce signal en une tension comprise entre -5 V et $+5\text{ V}$. Ce signal est envoyé vers une alimentation de courant qui est en fait un convertisseur tension-courant. Ce courant sert à alimenter une bobine supraconductrice qui produit un champ H sur l'échantillon. La figure 6.1 montre cette chaîne de commande.

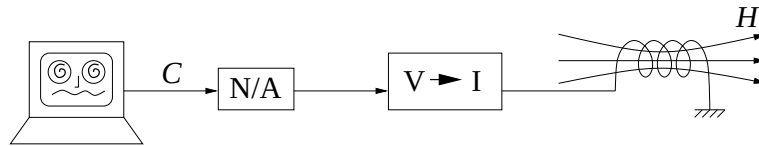


FIG. 6.1 – Système de commande du champ magnétique.

Idéalement, il faudrait qu'on puisse imposer n'importe quelle forme à la consigne $C(t)$ et qu'on ait à tout moment $H = C$. Ceci n'est pas le cas pour plusieurs raisons.

Double discrétisation : Le signal de commande C étant numérique, il ne peut prendre que des valeurs discrètes. Ceci introduit donc un bruit de discrétisation. D'autre part, l'ordinateur de commande a besoin d'un certain temps pour envoyer la consigne. Pour des raisons de commodité, celle-ci est envoyée de façon périodique. Ceci introduit une nouvelle discrétisation, sur la variable temps cette fois-ci.

Temps de réponse du système : L'ensemble convertisseur + alimentation + bobine a un temps de réponse non négligeable. Celui-ci est en fait presque entièrement imputable à l'alimentation de courant. En première approximation, on peut modéliser cette réponse comme celle d'un filtre passe-bas du premier ordre, soit

$$C = H + \tau \frac{dH}{dt}, \quad (6.1)$$

où τ est de l'ordre de 10 ms. Ce modèle est confirmé par des mesures de réponse en fréquence de l'alimentation. Cependant, pour des fréquences très élevées, le comportement cesse d'être du premier ordre. L'équation différentielle a donc des termes supplémentaires. Lorsque nous voudrions voir l'effet de ces termes supplémentaires, nous utiliserons l'équation

$$C = H + \tau \frac{d}{dt} \left(H + \tau' \frac{dH}{dt} \right), \quad (6.2)$$

où $\tau' \ll \tau$.

Alors que τ peut être mesuré assez précisément, τ' ne peut être considéré que comme un moyen de rendre compte qualitativement de l'écart du système à l'équation 6.1. On peut le voir comme une perturbation par rapport au comportement « idéal » défini par l'équation 6.1.

6.1.2 Comment obtenir le champ qu'on veut

Le temps de réponse τ du système entraîne une erreur dans la commande du champ qui est souvent inacceptable. Il faut donc en tenir compte au moment d'élaborer la consigne C de telle sorte que le champ H ait bien les variations qu'on attend de lui. Pour des raisons de simplicité, on recherche en général des variations sous forme de rampes de champ.

Réponse du premier ordre

Nous allons supposer que le système se comporte comme un filtre du premier ordre, c'est à dire conformément à l'équation 6.1. On considère le cas où on commence une rampe à $t = 0$ avec pour condition initiale

$$C = H = \frac{dH}{dt} = 0. \quad (6.3)$$

Si nous appliquons simplement la consigne

$$C = vt,$$

où v représente une vitesse de balayage, la réponse, donnée par intégration de 6.1, est

$$H = v(t - \tau) + v\tau e^{-t/\tau}.$$

La figure 6.2 illustre ce cas.

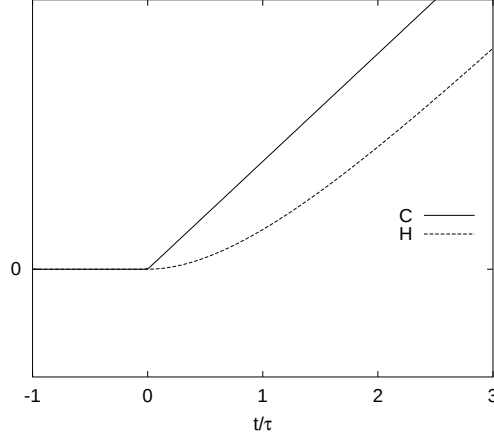


FIG. 6.2 – Réponse à une rampe.

Ce type de réponse est assez désagréable. Nous voulions une rampe de champ à pente constante et nous obtenons une variation non linéaire de $H(t)$. Ceci peut être évité en choisissant comme consigne une superposition d'une rampe et d'un saut. En effet, si nous prenons

$$C = v(t + \tau),$$

l'intégration de 6.1 nous donne

$$H = vt. \quad (6.4)$$

Ceci est illustré dans la figure 6.3.

Réponse du second ordre

Comme dit précédemment, l'équation 6.1 est une excellente approximation de la réponse du système. Le terme en τ' de l'équation 6.2 est en général négligeable. Cependant, la solution trouvée précédemment pour H présente une discontinuité de $\frac{dH}{dt}$ en $t = 0$. Le terme $\tau\tau'\frac{d^2H}{dt^2}$ ne peut donc être négligeable puisque $\frac{d^2H}{dt^2}$ diverge en $t = 0$. Nous allons alors calculer la réponse qu'on obtient en prenant l'équation 6.2 et en appliquant la même consigne.

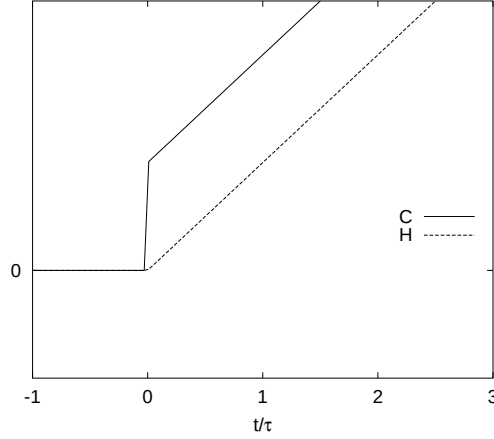


FIG. 6.3 – Réponse à la somme d'une rampe et d'un saut.

Équation homogène : La solution générale de l'équation sans second membre

$$H + \tau \frac{d}{dt} \left(H + \tau' \frac{dH}{dt} \right) = 0$$

est

$$H = Ae^{-t/\tau_1} + Be^{-t/\tau_2},$$

où A et B sont des constantes arbitraires et les constantes de temps τ_1 et τ_2 sont données par

$$-\frac{1}{\tau_i} = \frac{-\tau \pm \sqrt{\tau^2 - 4\tau\tau'}}{2\tau\tau'}.$$

Compte tenu du fait que $\tau' \ll \tau$, on peut développer l'expression ci-dessus au premier ordre en τ'/τ et on obtient tout simplement

$$\tau_1 = \tau \quad \text{et} \quad \tau_2 = \tau'.$$

Solution complète : La réponse sous la forme (6.4) est toujours une solution particulière pour $t > 0$. Sous les conditions initiales (6.3), la solution est

$$H = vt - v \frac{\tau\tau'}{\tau - \tau'} \left(e^{-t/\tau} - e^{-t/\tau'} \right).$$

La figure 6.4 montre cette solution comparée à celle qu'on obtient en négligeant $\tau\tau' \frac{d^2H}{dt^2}$.

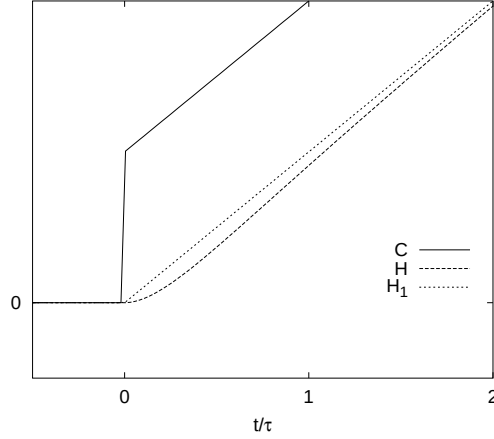


FIG. 6.4 – Réponse à la somme d'une rampe et d'un saut. H_1 est la solution obtenue précédemment avec l'équation 6.1. H a été calculé avec $\tau' = \tau/5$.

Sur cette figure, on peut voir clairement l'effet des deux temps de réponse. τ' correspond au temps nécessaire pour que $\frac{dH}{dt}$ atteigne la pente v . Ce temps entraîne une « erreur » de H de l'ordre de $v\tau'$ par rapport au cas où l'équation différentielle est du premier ordre. Il faut ensuite un temps τ pour que cette erreur se résorbe.

Freiner en douceur

En général, on cherche à avoir une variation du champ qui soit une succession de rampes, c'est à dire une fonction affine par morceaux du temps. Pour tenir compte du temps de réponse τ du système, le signal de commande présente une discontinuité $\Delta H = \tau \Delta v$ à chaque rupture de pente Δv . Il reste après chaque rupture de pente une erreur de l'ordre de $\tau' \Delta v$ qui se résorbe avec la constante de temps τ . Souvent cette erreur n'est pas gênante puisqu'elle entraîne seulement une incertitude sur la valeur du champ juste après la rupture de pente. Il existe pourtant un cas dans lequel cette erreur est inacceptable.

Supposons qu'on veuille mesurer le champ de retournement d'une nanoparticule pour des vitesses de balayage du champ extrêmement faibles (quelques $\mu\text{T/s}$). Pour ne pas perdre du temps en un balayage interminable, le champ est balayé à grande vitesse (de l'ordre de 1 T/s) jusqu'aux environs immédiats du champ de retournement. Au moment où on arrive au début de la plage de champ qu'on veut explorer, la vitesse de balayage est diminuée de plusieurs ordres de grandeur. Dans ce cas, le terme d'ordre deux dans l'équation différentielle 6.2 est responsable d'un pic de champ qui peut très bien être suffisamment élevé pour provoquer le retournement de l'aimantation de la particule. La figure 6.5 illustre ce problème.

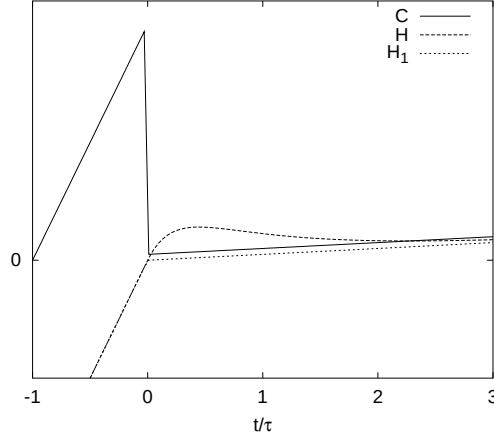


FIG. 6.5 – Freinage raté : le champ présente un pic au début du balayage lent. H_1 est la solution qu'on obtient avec l'équation 6.1.

Pour éviter cette situation, il a fallu mettre au point une technique pour ralentir le champ sans qu'il ne dépasse sensiblement la valeur qu'on attend de lui. Pour minimiser l'effet du terme d'ordre deux, nous n'allons pas chercher à imposer des discontinuités dans $\frac{dH}{dt}$. Il n'y aura donc pas de discontinuité dans la consigne C et on pourra négliger le terme d'ordre deux de l'équation différentielle. La figure 6.6 montre cette façon de procéder.

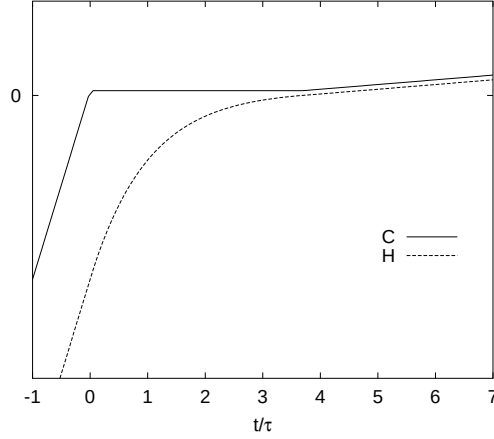


FIG. 6.6 – Freinage réussi : le champ ne va pas trop loin.

Notons H_0 le champ auquel on veut commencer le balayage lent à la vitesse v_2 et v_1 la vitesse du balayage rapide. Cette fois-ci, on introduit une phase de relaxation exponentielle entre les deux rampes de champ. La consigne fait un palier à la hauteur

$H_0 + v_2\tau$ pendant une durée $\tau \ln(v_1/v_2)$. En prenant comme origine des temps le début du plateau, l'intégration de (6.1), montre que pendant ce temps le champ relaxe exponentiellement suivant

$$H = H_0 + v_2\tau - v_1\tau e^{-t/\tau} \quad \text{pour} \quad 0 < t < \tau \ln(v_1/v_2) \quad (6.5)$$

et sa dérivée tombe de v_1 à v_2 . À la fin du palier, le champ vaut précisément H_0 et sa dérivée vaut v_2 . On peut donc enchaîner la deuxième rampe en gardant la continuité du champ et de sa dérivée.

Le terme du second ordre quand on freine en douceur

Avec ce procédé de freinage en douceur, on évite de faire ressortir la deuxième constante de temps en ne faisant pas diverger $\frac{d^2 H}{dt^2}$. On peut toutefois se demander si ce terme peut avoir une influence. Pour répondre à cette question, on peut regarder son effet en le considérant comme une perturbation à l'équation 6.1.

Nous allons poser

$$H = H_1 + H_2$$

où H_1 est la solution de l'équation 6.1 et H_2 l'effet du terme perturbatif $\tau\tau'\frac{d^2 H}{dt^2}$. En négligeant le terme $\tau\tau'\frac{d^2 H_2}{dt^2}$, l'équation 6.2 se réécrit

$$H_2 + \tau \frac{dH_2}{dt} = -\tau\tau' \frac{d^2 H_1}{dt^2}.$$

La solution à cette équation est

$$H_2(t) = \int_{-\infty}^t \frac{dt'}{\tau} e^{-(t-t')/\tau} \left(-\tau\tau' \frac{d^2 H_1}{dt'^2}(t') \right).$$

Dans le cas qui nous intéresse, en dérivant deux fois l'équation 6.5 on a

$$\frac{d^2 H_1}{dt^2} = -\frac{v_1}{\tau} e^{-t/\tau}$$

pour $0 < t < \tau \ln(v_1/v_2)$, la dérivée seconde de H_1 étant nulle en dehors de cet intervalle. L'expression intégrale de H_2 devient donc

$$\begin{aligned} H_2 &= \frac{v_1\tau'}{\tau} \int_0^t dt' e^{-t'/\tau} \\ &= \frac{v_1\tau'}{\tau} t e^{-t/\tau}. \end{aligned}$$

En particulier, au moment où on commence la deuxième rampe ($t = \tau \ln \frac{v_1}{v_2}$), cette perturbation vaut

$$H_2 \left(t = \tau \ln \frac{v_1}{v_2} \right) = v_2 \tau' \ln \frac{v_1}{v_2},$$

ce qui est très faible comparé à ce qu'on obtenait en essayant de freiner le champ brutalement (perturbation de l'ordre de $v_1 \tau'$).

6.1.3 Effet de la discrétisation du temps

Nous ne nous attarderons pas sur la discrétisation du champ dont le seul effet est d'introduire une incertitude sur le champ appliqué qui vaut

$$\delta H = 2^{-16} H_M \approx 1,5 \cdot 10^{-5} H_M,$$

H_M étant le champ maximal qu'on peut demander comme consigne. Étant donné que H_M est réglé pour être de l'ordre du champ de retournement des nanoparticules, nous avons toujours une incertitude relative de l'ordre de $1,5 \cdot 10^{-5}$ sur les champs de retournement mesurés.

L'effet de la discrétisation du temps peut être plus inquiétant. Étant donné que l'ordinateur envoie une consigne aux convertisseurs avec une périodicité $T_c \approx 1,5$ ms, la consigne C n'est jamais une vraie fonction affine du temps mais plutôt une fonction en escalier. On peut considérer que cette fonction est la somme de la rampe qu'on désire de pente v et d'une fonction en dents de scie de période T_c et d'amplitude crête à crête vT_c . Lors d'un balayage rapide ($v \approx H_M/1$ s), cette amplitude est de l'ordre de $1,5 \cdot 10^{-3} H_M$, soit deux ordres de grandeur au dessus du bruit de discrétisation du champ.

On voit à présent que le temps de réponse relativement long des alimentations n'est pas un défaut mais plutôt une qualité. La fréquence $1/T_c$ du bruit en dents de scie est sensiblement supérieure à la fréquence de coupure $1/2\pi\tau$ des alimentations. Celles-ci vont donc atténuer fortement ce bruit. Étant donné qu'à la fréquence $1/T_c$ on peut considérer qu'on a affaire à un intégrateur, la réponse du système au bruit en dents de scie va être un signal formé d'une succession de paraboles. L'amplitude crête à crête de ce signal vaut $vT_c^2/8\tau$. La figure 6.7 montre le bruit de discrétisation du temps sur C et sa répercussion sur H .

Si on considère une nouvelle fois un balayage rapide à $v \approx H_M/1$ s, l'amplitude crête à crête du bruit sur H est de $2,8 \cdot 10^{-5} H_M$, c'est à dire pas beaucoup plus grand que le bruit de discrétisation du champ. Cette estimation est même probablement pessimiste. En effet, la fréquence $1/T_c$ du bruit n'est probablement pas très loin de la deuxième fréquence de coupure $1/2\pi\tau'$ de l'alimentation. Quoi qu'il en soit, dès qu'on diminue un peu la vitesse de balayage, le bruit de discrétisation du temps tombe en dessous du bruit de discrétisation du champ.

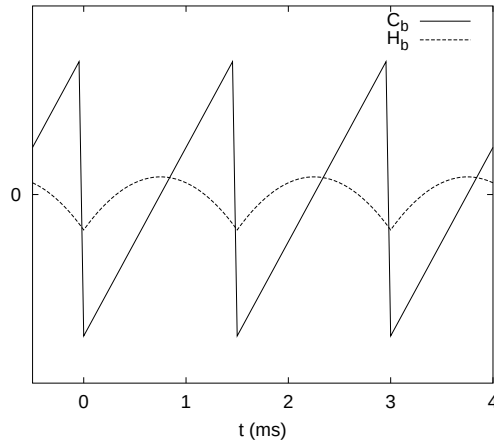


FIG. 6.7 – Bruit de discrétisation du temps dans la consigne C_b , et son effet sur le champ H_b . Ce dernier a été exagéré pour la clarté de la figure.

6.2 Modes de mesure

Le micro-SQUID peut être utilisé de plusieurs manières très différentes pour effectuer des mesures de flux. Nous expliquons ici les principales méthodes de mesure.

6.2.1 Cycle d'hystérésis

Pour des flux plus grands que $\Phi/2$, on ne mesure plus directement le flux, il est en effet beaucoup plus facile de compenser le flux créé par la particule qui traverse le SQUID avec un champ de compensation qui annule toutes ses variations. Pour cela, on définit un point de fonctionnement sur la courbe $I_c(H)$.

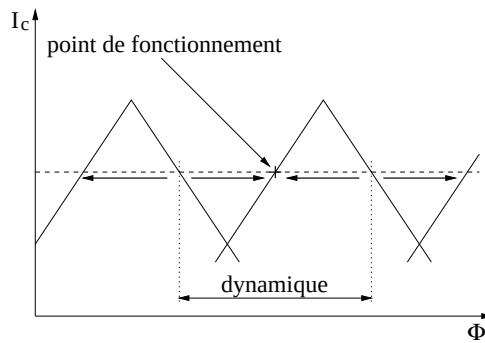


FIG. 6.8 – Choix du point de fonctionnement du SQUID pour avoir une grande dynamique. Les flèches horizontales indiquent l'effet de la rétroaction.

Le point de fonctionnement est défini de manière à avoir le maximum de dynamique. On voit sur la figure 6.8 qu'on a $\pm\Phi/2$ si on le place au milieu de la zone linéaire.

6.2.2 Mode froid

La difficulté principale de l'étude du champ de retournement avec un SQUID tient au fait que celui-ci chauffe par effet Joule dès que le courant critique est atteint. Une fois que le SQUID a transité, il dissipe pendant environ 100 ns, ce qui réchauffe la particule magnétique qui lui est couplée. Le mode froid constitue une solution à ce problème, car on utilise alors le SQUID comme un trigger [48].

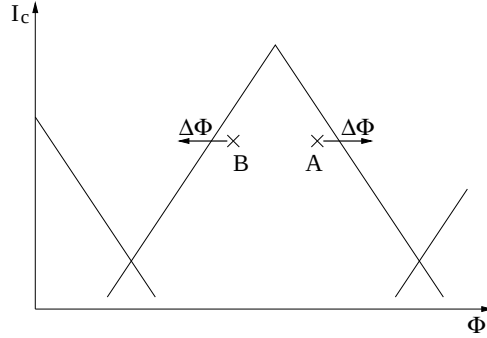


FIG. 6.9 – Polarisation du SQUID en mode froid. Les points A et B correspondent à un saut de flux respectivement positif et négatif.

Dans l'état supraconducteur, le SQUID est polarisé juste sous la valeur du courant critique, avec un champ appliqué perpendiculairement au plan du SQUID, afin qu'il soit dans l'un des états A ou B de la figure 6.9. Ces deux états sont adaptés respectivement à un saut de flux positif ou négatif induit par un retournement d'aimantation. C'est ce retournement qui déclenche la transition du SQUID dans l'état normal. Grâce à cette méthode, l'échantillon est réchauffé seulement après le retournement d'aimantation. C'est cette méthode de mesure que nous appelons le mode froid [49]. On peut ainsi détecter un saut d'aimantation dans un intervalle de quelques dizaines de nanosecondes.

6.2.3 Mode aveugle

Lorsqu'on s'approche, le long de la courbe critique, d'un point super-critique, l'amplitude du saut d'aimantation tend vers zéro. Par conséquent, la méthode de mesure précédente est incapable de détecter ce saut d'aimantation quand on se trouve trop près du point super-critique. Une méthode de mesure indirecte a été développée afin de surmonter ce problème, nous l'avons baptisée *mode aveugle*. Cette

méthode est illustrée sur la figure 6.10 qui montre la courbe critique au voisinage du point super-critique I . On va supposer qu'au point G le saut d'aimantation est trop petit pour pouvoir être directement détecté par le SQUID. En revanche, le saut en E est suffisamment grand pour être ainsi détecté. Le champ appliqué va parcourir le chemin $A \rightarrow B \rightarrow C \rightarrow B \rightarrow D \rightarrow F$, où C est un point du segment $[C_1C_2]$.

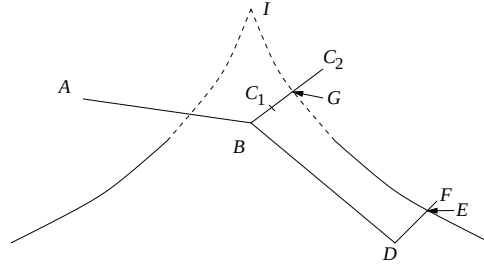


FIG. 6.10 – Les champs critiques qui sont indétectables directement sont déterminés de manière indirecte à l'aide d'une famille de trajets du champ appliqué.

Il y a deux points dans ce chemin où il peut y avoir un saut d'aimantation : les points G et E . Le fait que le chemin parte du point A nous assure que le système est bien à l'origine dans un puits de potentiel qui disparaît en G et en E . Le saut d'aimantation a donc lieu quand le puits en question disparaît pour la première fois, c'est à dire au moment où le champ atteint pour la première fois l'un de ces deux points. On voit sur la figure que si C se trouve dans le segment $[C_1G]$, le champ n'atteindra pas le point G et le saut aura lieu en E où il sera mesuré. Si C se trouve dans le segment $[GC_2]$, le saut aura lieu en G et on ne détectera pas de saut en E . Il en résulte que, en regardant s'il y a ou non un saut au point E , on peut savoir la position relative de G et de C . On arrive ainsi, en répétant cette procédure pour plusieurs choix de C , à déterminer la position de G avec la précision voulue.

L'algorithme de recherche le plus naturel est la dichotomie. Il arrive cependant qu'on ait des erreurs de mesure lorsqu'on veut déterminer s'il y a ou non un saut au point E . Ces erreurs de mesures peuvent aller dans les deux sens. D'une part, si le SQUID est polarisé très près de son courant critique (ce qui est nécessaire quand on mesure des particules donnant des petits signaux), il peut transiter spontanément à cause du bruit ambiant. Ceci peut nous faire croire à tort qu'il y a eu un saut au point E . D'autre part, il peut arriver aussi qu'il y ait un saut d'aimantation et que le SQUID ne transite pas. C'est aussi la faiblesse du signal magnétique de la particule qui est à l'origine de ce problème. De telles erreurs, même si elles sont peu fréquentes, peuvent rendre illusoire une recherche dichotomique. Il faut en effet un nombre non négligeable d'essais dans une telle recherche (typiquement de l'ordre de huit à dix) et une seule erreur suffit à fausser le résultat. On peut aboutir à des estimations très fausses de la position de G si l'erreur s'est produite lors des

premiers pas de la recherche. Nous verrons au chapitre suivant que ce problème n'est pas insurmontable et qu'on arrive, par un raffinement de l'algorithme, à une détermination raisonnablement fiable des champs de retournement par la méthode aveugle.

Le mode aveugle a été développé avec pour première motivation de pouvoir déterminer les champs de retournement au voisinage des points super-critiques. Cependant, cette méthode nous a apporté un « bonus » qui semble *a posteriori* bien plus important que la motivation initiale. Elle rend en effet possible la mesure des champs de retournement en trois dimensions ! Pour se convaincre de cette possibilité, il suffit de reconsidérer la figure 6.10 et de remarquer que nous n'avons besoin de pouvoir détecter un retournement d'aimantation que sur le trajet $D \rightarrow F$. Autrement dit, il n'y a que sur ce trajet que le SQUID doit être allumé. Nous savons que le SQUID ne fonctionne que si le champ appliqué n'a pas de composante significative perpendiculaire à sa boucle, et c'est bien cela qui nous limitait à des mesures bidimensionnelles. Pourtant, dans le mode aveugle, seul le trajet $D \rightarrow F$ est soumis à cette contrainte, le reste du chemin pouvant se trouver n'importe où dans l'espace des champs. Nous pouvons ainsi faire balayer au segment $[C_1C_2]$ toutes les directions de l'espace et déterminer de cette façon l'ensemble de la surface critique d'une nanoparticule.

Pour pouvoir avoir un aperçu de cette surface en même temps qu'on la mesure, on a développé une bibliothèque de dessin de surfaces tridimensionnelles. C'est en effet très utile de pouvoir adapter les paramètres de contrôle du SQUID en fonction du bruit sur la mesure. De plus, comme ces mesures sont longues, ils faut pouvoir vérifier en temps réel que cela fonctionne.

Chapitre 7

Problème de l'oracle qui a bu

Comme nous l'avons vu précédemment, la mesure des astroïdes est basée sur une recherche dichotomique. Il arrive malheureusement que la recherche conduise à un résultat erroné à cause d'erreurs de mesure. Ceux-ci sont de deux natures : il peut arriver que le SQUID transite sur un bruit, ce qui est interprété à tort comme une transition de la particule. Il peut arriver aussi que la particule retourne son aimantation sans que le SQUID transite, puisque le signal magnétique de celle-ci est à la limite de la sensibilité du SQUID.

Pour ces raisons il est important d'avoir un algorithme de recherche qui soit plus tolérant aux erreurs qu'une simple dichotomie. Ceci devient d'autant plus vrai que nous sommes maintenant en mesure de mesurer des particules beaucoup plus petites qu'auparavant, mais au prix d'un taux d'erreurs beaucoup plus grand pour un pas individuel de la dichotomie.

Pour rendre l'approche de ce problème plus ludique, et pour pouvoir facilement en discuter avec des gens qui ne sont pas au courant des nos méthodes expérimentales, j'ai formalisé ce problème sous le nom de *problème de l'oracle qui a bu*. Ce chapitre expose la description formelle du problème ainsi que les différentes approches envisagées pour le résoudre.

7.1 Posons le problème

Soit x un réel compris entre 0 et 1. On veut connaître sa valeur avec une certaine précision fixée à priori (mettons $1/256$ pour fixer les idées). Pour ce faire, on peut interroger un oracle qui connaît la réponse. On choisit alors un nombre a et on lui demande « Ô grand oracle, dis moi si a est plus grand que x », et l'oracle répond. C'est d'ailleurs le seul type de question à laquelle il puisse répondre, on ne peut donc pas se contenter de lui demander ce que vaut x ...

Problème 1 : Trouver une stratégie qui nous permette de connaître x avec la pré-

cision souhaitée et en posant le moins de questions possible à l'oracle.

Solution 1 : On fait une dichotomie. En posant p questions on détermine x avec un précision de 2^{-p} .

Ceci est vraiment trop simple. Il va falloir maintenant compliquer l'affaire !

L'oracle a un peu bu. Quand on lui pose une question il peut nous donner la bonne réponse mais il peut aussi se tromper. La probabilité pour qu'il se trompe est égale à son taux d'alcoolémie. La statistique de ses erreurs est sympathique : le fait qu'il se trompe une fois est statistiquement indépendant du fait qu'il se trompe une autre fois. Si on ne supporte pas l'idée d'un oracle porté sur la bouteille, on peut dire que l'oracle est sobre mais qu'il nous parle à travers une ligne bruitée, et il arrive qu'on interprète mal ses réponses.

Dans ces conditions, on s'aperçoit que l'algorithme simple de la dichotomie nous donne souvent le mauvais résultat. Même si l'oracle ne fait que 1 % d'erreurs, avec 8 pas de dichotomie on est soi même à 7,7 % d'erreurs. Il faut maintenant trouver le moyen de trouver x avec un niveau de certitude meilleur que ça. Évidemment on n'aura jamais une certitude absolue, mais on est content si on n'a que (par exemple) 0,1 % de chances de se tromper. Évidemment on devra lui poser plus de 8 questions, mais on voudrait quand même lui en poser le moins possible. Pour simplifier on va dire qu'on connaît le taux d'alcoolémie de l'oracle.

Problème 2 : Trouver une stratégie optimale pour interroger l'oracle qui a bu. C'est à dire trouver un algorithme qui minimise *en moyenne* le coût (nombre de questions posées), tout en assurant la précision et le niveau de certitude voulus.

7.2 Éléments de réponse

Je n'ai malheureusement pas la solution au problème, mais j'ai trouvé un certain nombre de pistes plus ou moins utiles pour piloter l'expérience.

7.2.1 Reposer la question pour confirmer la réponse

Il ne s'agit évidemment pas de poser deux fois chaque question (trop coûteux). Il s'agit de suivre l'algorithme de dichotomie ordinaire et, à la fin, quand on sait que $A \leq x < B$ (avec $B - A$ la précision souhaitée), on lui redemande

- $A > x$? (il devrait répondre non) ;
- $B > x$? (il devrait répondre oui).

Si l'oracle confirme ce qu'on soupçonnait déjà, on considère le résultat valable, sinon on recommence tout...

... Non, en fait il n'est pas nécessaire de tout recommencer, il y a des choses qu'on sait déjà. En effet, si lors de la première dichotomie l'oracle nous a dit que

- $C > x$;
- $D > x$;

et que par ailleurs on sait que $C > D$, alors on peut prendre pour acquis que $C > x$ (on a deux réponses qui l'affirment). En procédant comme ça on est sûr de faire la deuxième dichotomie sur un intervalle plus petit que la première. On est aussi sûr que l'algorithme donne un résultat en un temps fini même dans le pire des cas. On peut remarquer que l'idée de base de cet algorithme est la même que celle qu'on trouve dans beaucoup de codes auto-correcteurs d'erreurs, à savoir on néglige la probabilité que l'oracle se trompe deux fois.

Cet algorithme marche assez bien. C'est une variante de ceci qui a été longtemps utilisé pour les mesures avec des bons résultats. Pourtant, avec les particules de seulement quelques nanomètres, nous sommes maintenant confrontés à des oracles qui boivent vraiment trop, et là l'algorithme avoue ses limites :

- si on recommence souvent, le coût varie comme $O(p^2)$ au lieu de $O(p)$ (pour une précision de p bits) ;
- le taux d'erreurs devient beaucoup trop grand.

C'est pour ça que je me suis posé le problème de trouver une solution optimale à ce problème pour une alcoolémie quelconque (enfin, inférieure à 1/2 quand même, on verra pourquoi).

7.2.2 Demander n confirmations

C'est la même chose que précédemment, sauf que la première étape (faire une dichotomie simple) est remplacée par une dichotomie avec $n - 1$ confirmations. J'ai écrit cet algorithme sous forme récursive et il s'avère assez élégant sous cette forme. La dichotomie simple et la dichotomie avec une confirmation exposée ci dessus sont des cas particuliers de cet algorithme plus général.

Afin d'estimer les performances de ces algorithmes, je les ai mis à l'épreuve d'une simulation Monte-Carlo. Pour chaque taux d'alcoolémie de l'oracle, l'algorithme testé doit résoudre le problème 40000 fois. La figure 7.1 montre les performances comparées de ces algorithmes. Elle a été dessinée pour une précision souhaitée de 1/256 (8 pas dans une dichotomie ordinaire). À titre de comparaison, l'algorithme du double questionnement a été représenté aussi. Celui-ci est une simple dichotomie où chaque question est posée deux fois, et une troisième fois si les deux premières réponses diffèrent.

J'ai essayé d'estimer, en fonction de l'alcoolémie a de l'oracle, le nombre minimum de questions qu'on peut espérer poser pour obtenir le résultat. Cette estimation se base sur l'idée suivante : dire que l'oracle se trompe avec une probabilité a équivaut à dire qu'il lui arrive, avec une probabilité $2a$, de répondre au hasard, et avec une probabilité $1 - 2a$, de donner la bonne réponse. L'information contenue dans chaque réponse est donc en moyenne de $1 - 2a$ bits. Pour obtenir les 8 bits d'information

dont on a besoin pour trouver x , il nous faut donc au moins $8/(1 - 2a)$ réponses. C'est ce qui est représenté dans la courbe notée *minimum* sur le graphe du coût. On voit ainsi pourquoi quand $a = 1/2$ on n'a plus d'espoir : les réponses de l'oracle, complètement aléatoires, ne contiennent aucune information.

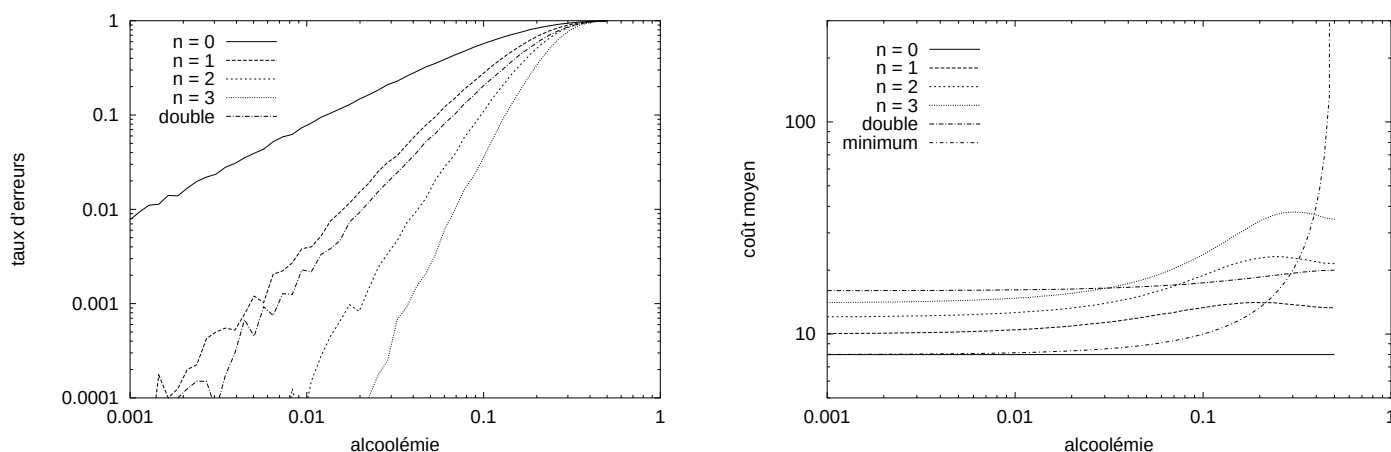


FIG. 7.1 – Performances de l'algorithme de dichotomie avec n confirmations. La courbe $n = 0$ correspond à une dichotomie ordinaire. Celle marquée *double* correspond à l'algorithme du double questionnement.

On peut observer sur ces courbes un certain nombre de choses dont on pouvait se douter à priori :

- Lorsque a est petit, la dichotomie avec n confirmations donne un taux d'erreurs qui varie comme a^{n+1} ; en effet, il faut $n + 1$ erreurs de l'oracle pour qu'on aboutisse à un mauvais résultat ;
- toujours lorsque a est petit, le coût quand on demande n confirmations est $8 + 2n$: 8 questions pour la dichotomie et n confirmations de chaque borne de l'intervalle trouvé à la fin ;
- quand a augmente, le coût augmente puisqu'on doit parfois recommencer une partie de la dichotomie, sauf évidemment pour $n = 0$;
- quand a tend vers $1/2$, le taux d'erreurs tend vers 1 puisque l'information contenue dans les réponses de l'oracle tend vers zéro.

Mais on voit aussi un certain nombre de choses pas forcément aussi évidentes :

- Quand a tend vers $1/2$, le taux d'erreurs tend vers 1 avec une tangente horizontale, ce qui est assez désagréable si on a affaire à des taux d'erreurs élevés ;
- l'algorithme du double questionnement est très mauvais ; il produit à peine moins d'erreurs que $n = 1$ tout en étant toujours sensiblement plus cher.

Le problème de cette famille d'algorithmes est qu'ils ne garantissent pas un taux d'erreurs maximal. Celui-ci dépend en effet de l'alcoolémie de l'oracle. Un moyen de

contourner cet obstacle serait de choisir le nombre n de confirmations à demander en fonction de cette alcoolémie, de manière à maintenir le taux d'erreurs en dessous d'un seuil prédéfini.

7.2.3 Maximiser l'information soutirée à l'oracle

Plus l'oracle a bu, moins ses réponses contiennent de l'information. Je me suis dit que ce serait peut-être une bonne idée de maximiser l'information contenue dans ses réponses. Pour cela, on peut tracer à chaque étape la densité de probabilité de présence de x dans l'intervalle $[0, 1]$. Au départ c'est une distribution plate (probabilité à priori uniforme). À chaque fois qu'il répond, cette probabilité monte du côté de la réponse et descend de l'autre côté. De l'alcoolémie de l'oracle dépend de combien cela monte et descend. Si à chaque étape on l'interroge sur la médiane de cette densité de probabilité, il y aura une chance sur deux pour chaque réponse. On maximise donc ainsi l'information qu'il nous donne.

La figure 7.2 montre quoi peut ressembler cette densité de probabilité après un certain nombre d'itérations. On voit que la bonne solution (ici $x = 0,767578$) a une probabilité très proche de un. Les marches dans l'histogramme correspondent aux valeurs sur lesquelles on a interrogé l'oracle.

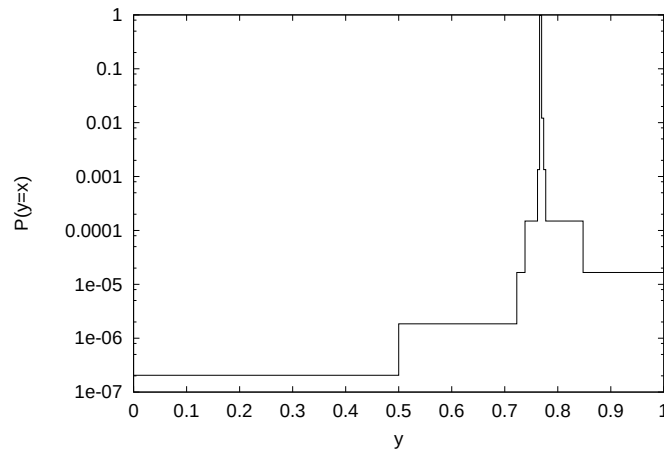


FIG. 7.2 – Probabilité de présence de x dans l'intervalle $[0, 1]$ après 12 itérations de l'algorithme de maximisation de l'information. L'intervalle a été discrétisé en 256 canaux. L'alcoolémie est de 0,1.

Cas particulier : Si l'oracle est sobre, on se retrouve en train de faire une dichotomie ordinaire.

En discrétisant l'intervalle $[0, 1]$ en canaux ayant pour largeur la précision souhaitée, cette méthode finit par concentrer toute la probabilité sur un seul canal. En

effet, à cause de la discrétisation, quand la médiane se trouve dans le canal le plus probable, on questionne l'oracle pour l'un des bords du canal (le plus proche de la vraie médiane). Si le canal contient la bonne réponse, sa probabilité va alors tendre vers 1. L'algorithme se termine quand la probabilité d'un canal atteint le niveau de certitude souhaité. La figure 7.3 montre les performances de cet algorithme pour trois taux d'erreurs admis différents.

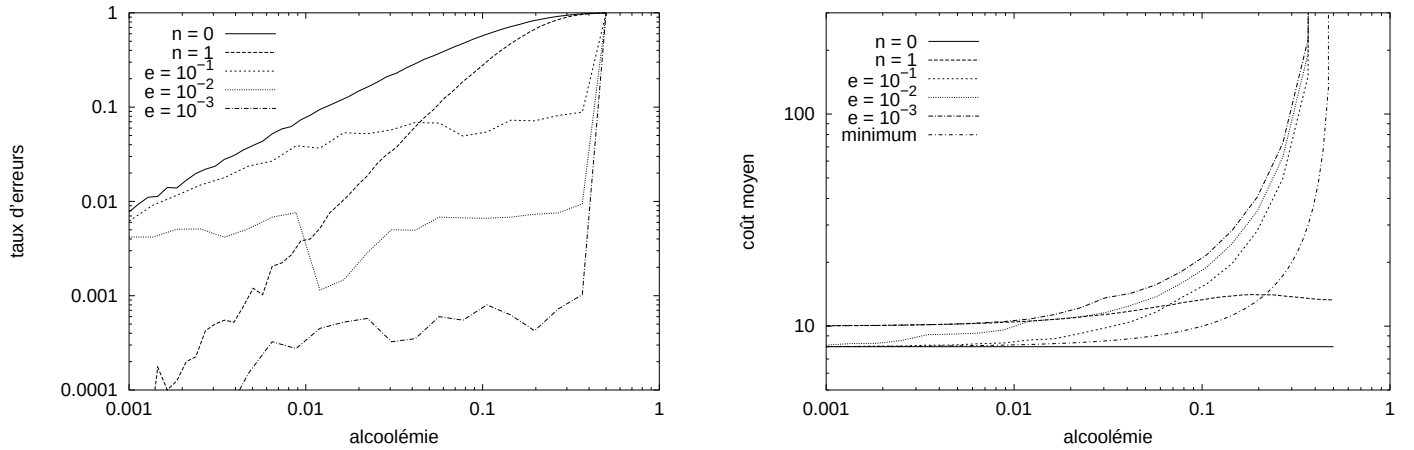


FIG. 7.3 – Performances de l'algorithme de maximisation de l'information pour trois taux d'erreurs e admis. À titre de comparaison, les performances des dichotomies simple ($n = 0$) et avec une confirmation ($n = 1$) ont été superposées.

On remarque sur ces courbes que l'algorithme tient ses promesses : le taux d'erreurs constaté est toujours inférieur au taux accepté à priori, même lorsque celui-ci est très bas et l'alcoolémie très forte. On voit aussi que, de façon très prévisible, le coût diverge quand a tend vers $1/2$.

On peut voir sur ces courbes que cet algorithme est bien meilleur que le précédent. En effet, si on regarde les croisements des courbes, on constate que :

- à taux d'erreurs égal, l'algorithme de maximisation de l'information est toujours moins cher ;
- à coût égal, il a toujours un taux d'erreurs inférieur.

Ceci apparaît de façon particulièrement claire dans le croisement des courbes des taux d'erreurs $n = 1$ et $e = 10^{-2}$. Juste avant le croisement, pour $a < 10^{-2}$, l'algorithme de maximisation de l'information avec $e = 10^{-2}$ a un coût sensiblement inférieur à celui de $n = 1$. Juste après le croisement, les coûts sont pratiquement égaux mais l'algorithme de maximisation de l'information donne un taux d'erreurs sensiblement inférieur.

Une bonne surprise est que le coût est peu dépendant du niveau de certitude choisi. Il n'est donc pas beaucoup plus coûteux d'accepter un taux d'erreurs de 10^{-3}

que de se contenter de 10^{-2} .

7.2.4 Autres pistes

Puisque les performances de l'algorithme ci-dessus sont satisfaisantes, je ne me suis pas donné la peine de chercher d'autres solutions. Je soupçonne aussi un peu l'algorithme de maximisation de l'information d'être optimal, mais je ne sais pas comment montrer cette optimalité. À titre purement indicatif, voici une autre idée qui pourrait valoir la peine d'être explorée si l'algorithme précédent s'avère non optimal.

Rétrécir progressivement l'intervalle de certitude Appelons intervalle de certitude le plus petit intervalle dans lequel on est sûr (au taux d'erreurs qu'on accepte près) que se trouve x . *In fine* il faut que la largeur de cet intervalle soit inférieure ou égale à la précision souhaitée. Cet intervalle se déduit de la densité de probabilité évoquée plus haut.

Une approche qui pourrait être rentable serait de trouver, à chaque étape, la question qui maximise l'espérance de rétrécissement de cet intervalle. C'est ce qu'on appelle un algorithme glouton : on optimise successivement chaque étape.

En fait je ne crois pas beaucoup en cette approche. Lorsqu'on doit gravir un sommet au relief compliqué, la ligne de plus grande pente n'est pas forcément le chemin le plus court. De même, dans les problèmes un peu subtils, il arrive souvent que l'optimisation *locale* faite par l'algorithme glouton (suivre la ligne de plus grande pente) ne conduise pas à une optimisation *globale* (chemin le plus court).

Troisième partie

Résultats

Chapitre 8

Mesures quasi-statiques

Lorsqu'on opère à des températures suffisamment faibles, les champs de retournement mesurés deviennent indépendants de la température. On peut donc considérer qu'on observe le comportement à température nulle du système. Les mesures décrites ci-dessus ont été effectuées sans régulation de la température du cryostat. Celui-ci était donc à sa température minimale qui est de l'ordre de 50 mK. Dans ces conditions, on peut considérer que le retournement de l'aimantation a lieu quand la hauteur de la barrière d'énergie autour de l'état métastable devient nulle. Nous sommes donc dans les conditions étudiées théoriquement au chapitre 1.

8.1 En deux dimensions

Étant donné que le SQUID ne fonctionne que quand le champ appliqué est dans son plan, nous sommes contraints à appliquer le champ dans le plan de sa boucle lorsqu'on utilise le mode froid. Le mode aveugle ne présente pas cette limitation, mais il a lui aussi été utilisé dans un premier temps pour compléter des mesures bidimensionnelles. En tout état de cause, les mesures d'une particule en deux dimensions sont un préalable aux mesures en 3D, ne serait-ce que pour pouvoir choisir le point qui servira à tester si le saut d'aimantation a eu lieu. Nous présentons ci-dessus quelques mesures qui se sont révélées intéressantes.

8.1.1 Particule de Cobalt

Voici la particule magnétique qui s'est avérée la plus intéressante pour la fin de la thèse de Wolfgang et le début de la mienne. Cette situation lui a valu le sobriquet de « particule du siècle ». Il s'agit ici d'une nanoparticule de cobalt enrobée dans du carbone de quelques 20 nm de diamètre. Elle a été préparée à l'aide d'un plasma d'arc électrique [50] par l'équipe de Pascard à l'École Polytechnique [51]. La figure 8.1

montre une *autre* particule préparée de la même manière. On y distingue les plans atomiques.

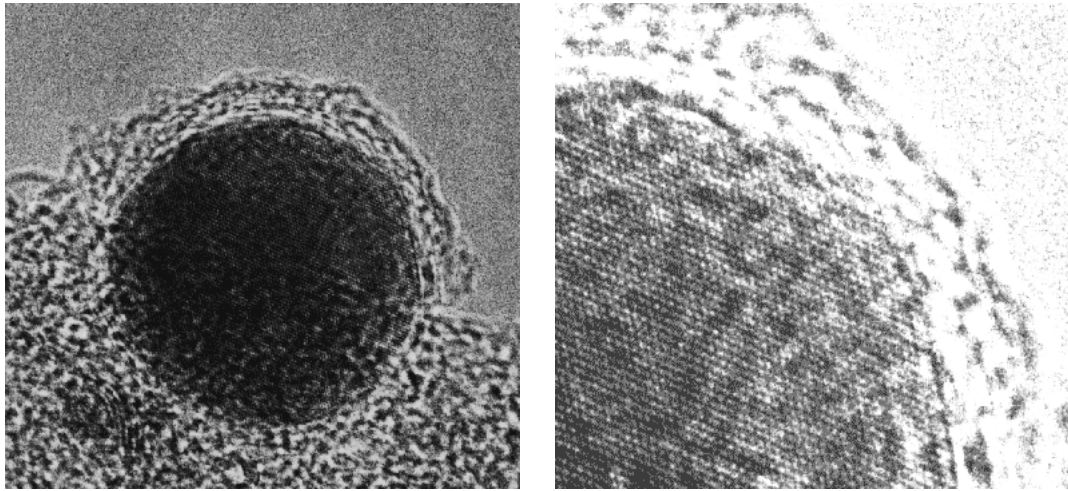


FIG. 8.1 – Particule de cobalt vue au microscope électronique par transmission. La figure de droite est un agrandissement éclairci du quart supérieur droit de la particule.

Les particules sont dispersées sur la surface du SQUID. Les seules qui arrivent à être mesurables sont celles qui sont le mieux couplées avec celui-ci. La figure 8.2 montre des images du SQUID concerné prises au microscope électronique à balayage. On y distingue les particules les mieux couplées au niveau du micro-pont de gauche.

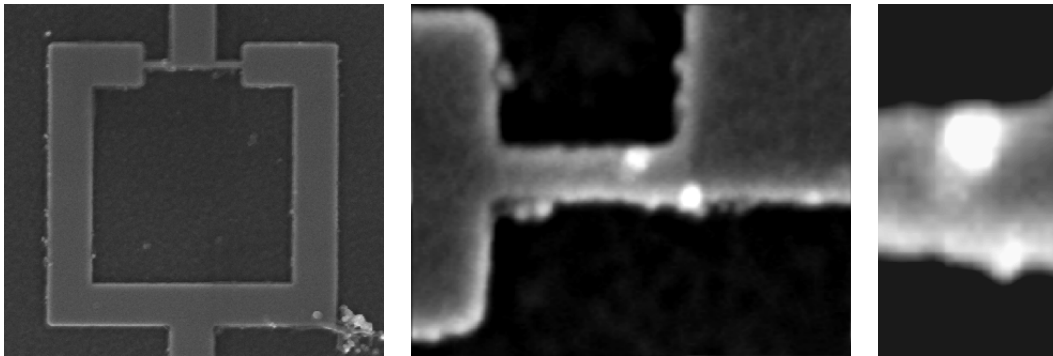


FIG. 8.2 – Zooms successifs sur les particules de cobalt posées sur le SQUID vues au microscope électronique à balayage. Deux particules semblent fortement couplées sur le micro-pont de gauche.

Ici on a affaire à des particules donnant un signal suffisamment fort pour pouvoir tracer un cycle d'hystérésis. Celui-ci est représenté sur la figure 8.3.

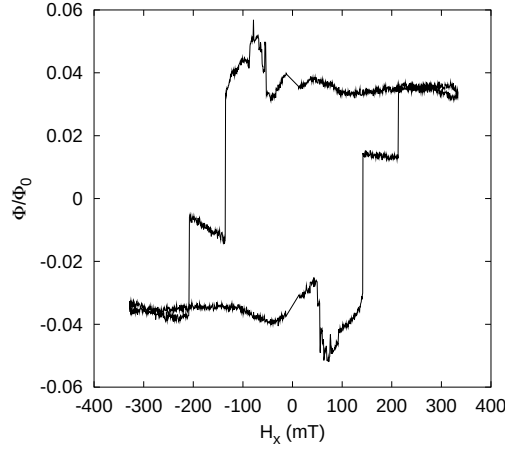


FIG. 8.3 – Cycle d’hystérésis du SQUID comportant la particule du siècle. On remarque deux sauts de flux correspondant aux sauts d’aimantation de deux particules bien couplées au SQUID.

Il est intéressant de remarquer qu’on voit deux sauts importants dans ce cycle alors que justement, dans la figure 8.2 on avait remarqué deux particules fortement couplée à un micro-pont. On est tout naturellement conduits à penser que chacun de ces sauts correspond à l’une des deux particules. Cependant, pour s’assurer que ces deux sauts proviennent bien de deux particules indépendantes et non d’un seul et même système, il faut étudier la dépendance angulaire de ces sauts. Il est toutefois plus commode de faire cette étude en mode froid (où on ne voit que les sauts d’aimantation) plutôt qu’en traçant des multiples cycles. La figure 8.4 montre la mesure en mode froid sur ce SQUID.

La figure 8.4 montre les mesures brute et dépouillée des champs de retournement en mode froid. Le champ a été balayé radialement suivant une direction lentement variable et on a marqué d’un point chaque champ où une transition du SQUID est détectée. On peut remarquer plusieurs choses ici :

- on ne mesure pas une mais plusieurs particules simultanément ;
- certaines transitions semblent être des bruits de mesure ;
- beaucoup de transitions sont concentrées à l’origine.

Les deux sauts d’aimantation du cycle de la figure 8.3 ont pu être associés aux particules notées 1 (particule du siècle) et 2 de la figure 8.4.

Les transitions du SQUID au voisinage de l’origine sont en fait provoquées par la transition supraconducteur $\rightarrow$ normal des fils de connexion du SQUID. Ces fils sont en aluminium et expulsent le champ par effet Meissner lors de leur transition. Ceci produit une variation importante du flux à travers la boucle du SQUID, d’où la transition observée. Dans un cas comme celui-ci, où il y a plusieurs particules sur

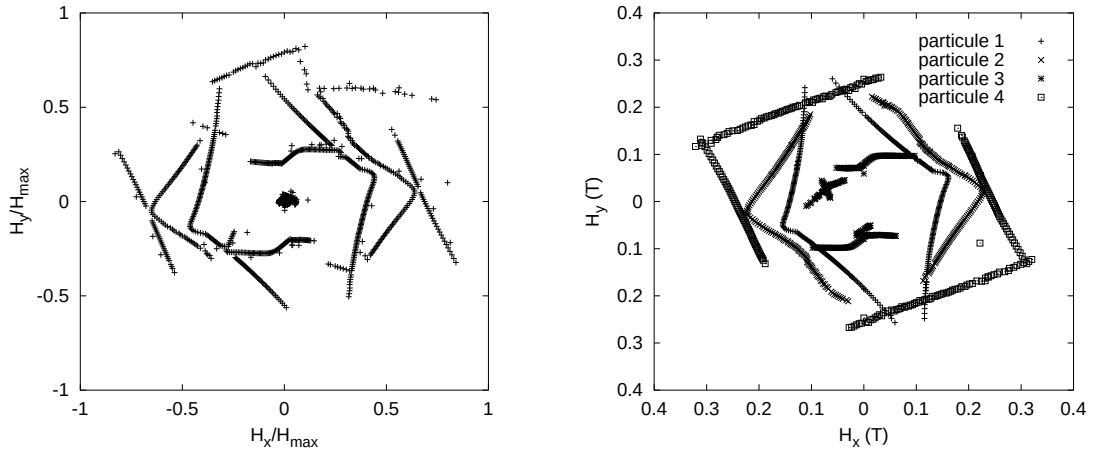


FIG. 8.4 – Mesures brutes des champs de retournement sur des particules de cobalt. La figure de droite montre une autre mesure du même SQUID où on a trié les points de mesure pour les associer aux différentes particules.

le SQUID, on choisit en général une particule qui semble facile à mesurer. Ici ce sera la seule qui montre une trace très régulière et très complète. Les autres peuvent être éventuellement mesurées en mode aveugle.

Ces mesures d'astroïdes 2D permettent non seulement de confirmer qu'on a affaire à des particules indépendantes, mais aussi d'estimer leur interaction. En effet, si on est dans une situation où deux particules sont susceptibles de se retourner sous l'influence du champ appliqué, le retournement de la première d'entre elles modifie par couplage dipolaire le champ vu par l'autre. Le champ de retournement d'une particule n'est donc pas le même suivant qu'elle se retourne avant ou après sa voisine. Ceci peut être mis en évidence en regardant l'endroit où les courbes critiques des deux particules se croisent, ce qui est montré dans la figure 8.5. On constate sur cette figure que les champs de retournement pour chacune des deux particules sont décalés d'environ 2 mT de part et d'autre du croisement. Ce décalage donne l'ordre de grandeur du champ de fuite que chacune crée sur sa voisine. On voit aussi que le saut d'aimantation dans une particule retarde le saut dans l'autre. On peut donc dire que chacune des particules produit, dans son état le plus stable, un champ qui s'oppose au retournement de l'autre.

Nous avons essayé d'ajuster l'astroïde mesurée à un modèle de retournement uniforme en deux dimensions avec deux constantes d'anisotropie. Le résultat de cet ajustement est présenté sur la figure 8.6.

Au moment où nous avons fait cet ajustement, nous ne nous doutions pas encore que des mesures tridimensionnelles au micro-SQUID étaient possibles. Au niveau théorique, la généralisation du modèle de retournement uniforme à trois dimensions

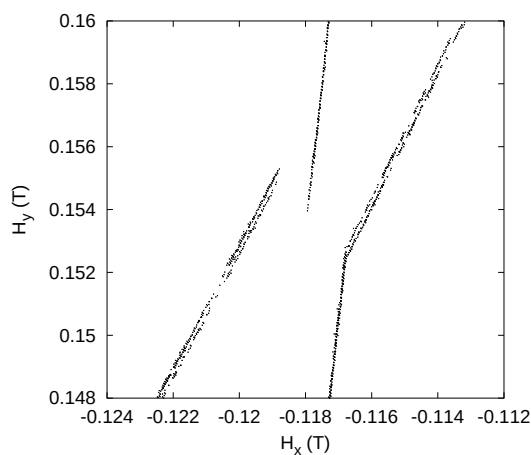


FIG. 8.5 – Détail du croisement des courbes critiques des particules 1 et 2.

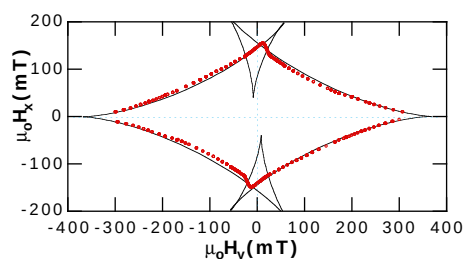


FIG. 8.6 – Ajustement d'une « astroïde » avec deux constantes d'anisotropie sur la mesure précédente.

n'avait pas encore été développé. Même le fait de considérer une deuxième constante d'anisotropie était peu courant et le fait de choisir deux axes différents pour les deux constantes franchement original. Avec le recul qu'on a à présent, on s'aperçoit que notre approche était très simpliste. On s'est aussi aperçu que l'ajustement considéré est franchement très mauvais quand on regarde cette astroïde en trois dimensions !

8.1.2 Particule de fer-cuivre

La figure 8.7 montre la courbe critique mesurée sur une particule de FeCu d'environ 15 nm de diamètre. Cette particule a été préparée à Jussieu dans l'équipe de Mme Pileni à l'aide de la méthode des micelles inverses. Cette particule possède une anisotropie cristalline de symétrie cubique ainsi qu'une anisotropie de forme. La désorientation des axes des deux termes de l'anisotropie est responsable de l'apparente complexité de la courbe obtenue. Les champs de retournement dans cette figure ont été mesurés à environ 100 mK.

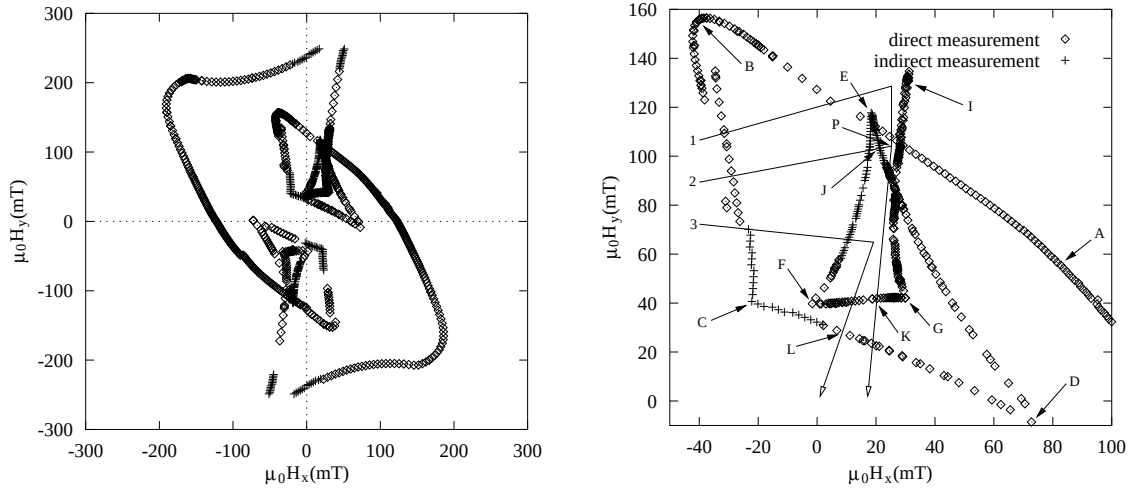


FIG. 8.7 – Courbe critique mesure sur une particule de FeCu d'environ 15 nm de diamètre. La figure de droite est un agrandissement de la région intéressante.

Cette mesure a été l'occasion de vérifier que ce que prévoit la théorie pour les courbes critiques présentant des croisements est bien vérifié dans des systèmes physiques réels [52]. En particulier en ce qui concerne les bifurcations.

Pour pouvoir mesurer la totalité de cette courbe, il a fallu apporter une attention particulière au choix des chemins suivis par le champ. Il y a deux bifurcations super-critiques dans la région agrandie de la figure 8.7 : les points C et E . En supposant que l'aimantation est à l'origine orientée vers la gauche, on peut se demander en quels points il y a un saut d'aimantation lorsque le champ $\mathbf{H}$ traverse la région agrandie. Si $\mathbf{H}$ passe en dessous de C , le saut a lieu sur la branche ABC . S'il passe au dessus de C , le saut a lieu quelque part sur $CDEFGI$. Dans ce dernier cas, il faut savoir de quel côté passe $\mathbf{H}$ par rapport au point E . S'il passe en dessous de E , le saut a lieu dans CDE , autrement il a lieu dans $EFGI$. Le champ de retournement ne dépend pas, par exemple, du fait que $\mathbf{H}$ passe au dessus ou en dessous de B , par conséquent B ne représente pas une bifurcation super-critique. Pour la même raison, D , F et G ne sont pas des bifurcations super-critiques.

À titre d'exemple, pour montrer le type de choix que nous avons eu à faire, considérons les trois chemins représentés dans la figure 8.7. Ils passent tous les trois au dessus de C , le saut a donc lieu sur $CDEFGI$. Le chemin 1 passe au dessus de E , il y a donc un saut en K . Le chemin 3 passe en dessous de E , le saut a donc lieu en L . Le chemin 2 est plus intéressant : puisqu'il passe en dessous de E , il y a un saut en J . Cependant, il s'agit d'un tout petit saut entre deux puits métastables qui se fondent en E , tout comme le saut en B_1 pour le chemin $A \rightarrow B_2 \rightarrow C \rightarrow B_1 \rightarrow A$ sur la figure 1.6 du chapitre 1. Quand le champ atteint le point P , l'aimantation est dans le même état qu'elle aurait eu en suivant le chemin 1, elle fait donc un

deuxième saut en K .

En fait, bien que le chemin 2 présente deux sauts, seul le deuxième peut être mesuré directement. Puisque le saut en J donne un signal magnétique beaucoup trop faible pour être détecté par le SQUID, il a fallu utiliser le mode aveugle pour obtenir ce point, comme tous les autres points au voisinage d'une bifurcation super-critique. Tous les points représentés par des croix sur la figure 8.7 ont été obtenus de cette manière.

8.2 En trois dimensions

8.2.1 Particules de cobalt

Particule du siècle

Dans la figure 8.8 nous retrouvons la particule que nous avons vue précédemment en deux dimensions, et uniquement pour constater à quel point nous étions loin de la réalité en ajustant un modèle bidimensionnel sur nos mesures.

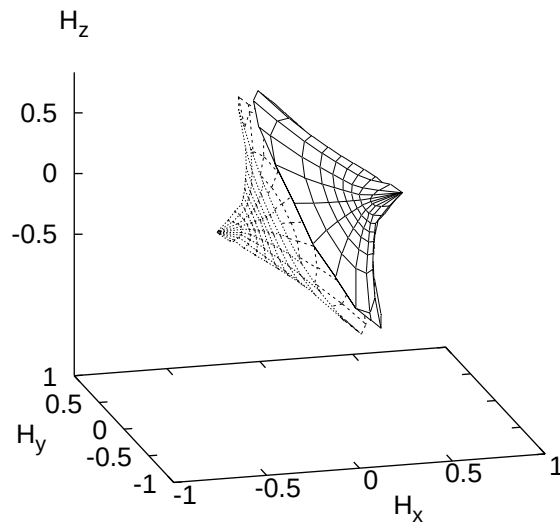


FIG. 8.8 – Surface critique de la particule du siècle.

Dans l'ajustement de la figure 8.6 nous voyons des croisements dans la surface critique qui correspondent à l'existence en trois dimensions de plusieurs nappes qui se raccordent le long de fronces. Comme les mesures sont faites avec un balayage radial

du champ, on devrait mesurer une surface incomplète qui doit apparaître « fendue » là où le raccord se fait mal. Il est possible que le hasard fasse que cette fente ne soit pas apparente dans un plan particulier, mais les mesures en trois dimensions nous permettent d'assurer que cette fente n'existe pas. Ceci peut paraître incertain au vu de la résolution de la figure 8.8, mais celle-ci ne représente, pour des raisons de lisibilité, qu'une petite fraction de tous les points mesurés. Les mesures complètes confirment que la surface n'est pas fendue.

Taillefer

La figure 8.9 montre la surface critique d'une autre particule de cobalt de la même série que la précédente. Son surnom vient du fait que, vue d'un certain angle, cette surface critique évoque une montagne du même nom près de Grenoble.

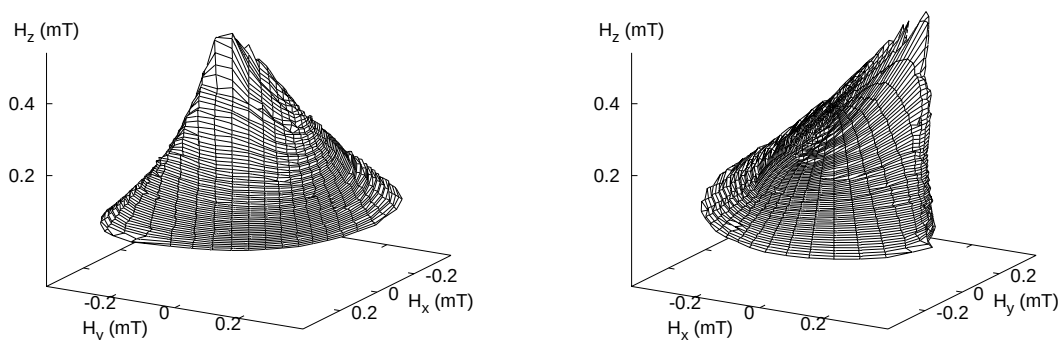


FIG. 8.9 – Surface critique de la deuxième particule de cobalt. Seule la moitié de la surface est représentée, l'autre étant strictement symétrique.

La mesure d'une surface critique demande un temps très long, puisque chaque point de mesure est le résultat d'une recherche dichotomique qui demande un trajet important dans l'espace des champs. Pour cette raison nous ne mesurons qu'une seule moitié, l'autre moitié étant déduite par symétrie. Le fait que les surfaces critiques soient symétriques a été toujours vérifié à deux dimensions, lorsqu'on en mesure des coupes planes. On pourrait vouloir mesurer la totalité de la surface pour avoir encore une confirmation expérimentale de cette symétrie, mais le coût en temps de mesure est trop important pour justifier ce qui ne serait qu'une confirmation de plus parmi tant d'autres.

8.2.2 Particule de ferrite de baryum

La figure 8.10 montre la surface critique mesurée sur une particule de ferrite de baryum d'environ 20 nm de diamètre [53]. La formule chimique de cette particule est $\text{BaFe}_{12-2x}\text{Co}_x\text{Ti}_x\text{O}_{19}$ avec $x = 0,8$, mais nous nous contenterons de noter BaFeO par simplicité. Les mesures ont été effectuées à quelques 50 mK. Bien que cette particule ait une surface critique assez simple et proche de l'astroïde de Stoner-Wohlfarth, elle ne peut pas être engendrée par la rotation d'une astroïde bidimensionnelle. Cependant, en prenant en compte l'anisotropie de forme et l'anisotropie cristalline de la particule, André Thiaville a pu ajuster la surface mesurée, comme montré sur la figure 8.11.

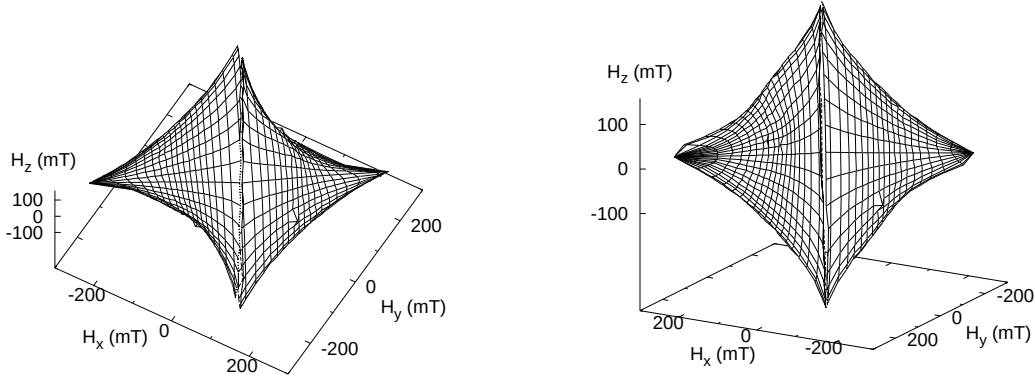


FIG. 8.10 – Surface critique d'une nanoparticule de BaFeO mesurée en trois dimensions. Les deux images correspondent à deux points de vue dans des directions perpendiculaires. Bien que quelques 3700 points aient été mesurés, seulement quelques centaines sont représentés sur cette figure.

La fonction d'anisotropie utilisée pour ajuster cette surface est :

$$\frac{2V_0(\mathbf{r})}{vM_s} = -Ar_x^2 + Br_y^2 + C(r_{x'}^2 + r_{y'}^2)^2 + D(r_{x'}^2 + r_{y'}^2)^3.$$

où (r_x, r_y, r_z) sont les coordonnées de $\mathbf{r}$ dans la base qui diagonalise la partie quadratique de $V_0(\mathbf{r})$ et $(r_{x'}, r_{y'}, r_{z'})$ ses coordonnées dans une base ayant pour axe z' l'axe c du cristal. Les constantes A , B , C et D sont les constantes d'anisotropie. La partie cristalline de cette anisotropie a une symétrie de révolution par rapport à z' . Ceci est cohérent avec la symétrie du cristal qui est hexagonale. Nous aurions pu considérer un terme moins symétrique ayant seulement une symétrie hexagonale (et non de révolution). Ceci n'a pas semblé utile puisque la surface mesurée ne montre

constante	valeur (mT)
A	347,2
B	3,5
C	29,5
D	-7,0

TAB. 8.1 – Constantes d'anisotropie de la particule de BaFeO.

pas la moindre trace de symétrie d'ordre six. Le tableau 8.1 montre les valeurs des constantes telles qu'obtenues par l'ajustement.

Comme en témoignent la figure 8.10 et le tableau 8.1, l'anisotropie de cette particule est essentiellement quadratique. Cependant, les termes d'ordre supérieur peuvent jouer un rôle important dans certaines situations telles que la détermination du splitting tunnel des niveaux d'énergie quantiques. Ils doivent donc être pris en compte si on s'intéresse à la possibilité d'avoir un retournement de l'aimantation par effet tunnel.

Des observations au microscope électronique des cristaux de BaFeO montrent qu'ils ont une forme de prisme droit à base hexagonale. Les faces du cristal sont parallèles aux plans cristallographiques et il y a une très forte anisotropie cristalline qui rend l'axe c facile. Si le cristal étudié avait une telle forme avec pour base un hexagone régulier, l'anisotropie aurait l'axe c du cristal comme axe de symétrie. L'absence d'une telle symétrie fait penser que la forme n'est pas si régulière. L'écart à la symétrie qui paraît le plus naturel consiste à avoir une base en forme d'hexagone non régulier, où les côtés forment entre eux des angles de 60° . La figure 8.12 montre une telle forme.

Un tel écart à la symétrie dans la forme du cristal pourrait expliquer l'absence de symétrie de révolution de la surface critique. Cependant, la forme qu'on considère ici présente un plan de symétrie indiqué par les pointillés sur la figure 8.12. La surface critique d'un tel cristal devrait donc avoir le plan difficile comme plan de symétrie, ce qui est loin d'être le cas au vu de la figure 8.10. On doit donc conclure que le cristal présente un écart à la symétrie parfaite plus important que ça. L'absence de symétrie par rapport au plan difficile pourrait s'expliquer par des faces latérales non perpendiculaires aux faces supérieure et inférieure (ce qui semble étonnant) ou alors par des faces incomplètes.

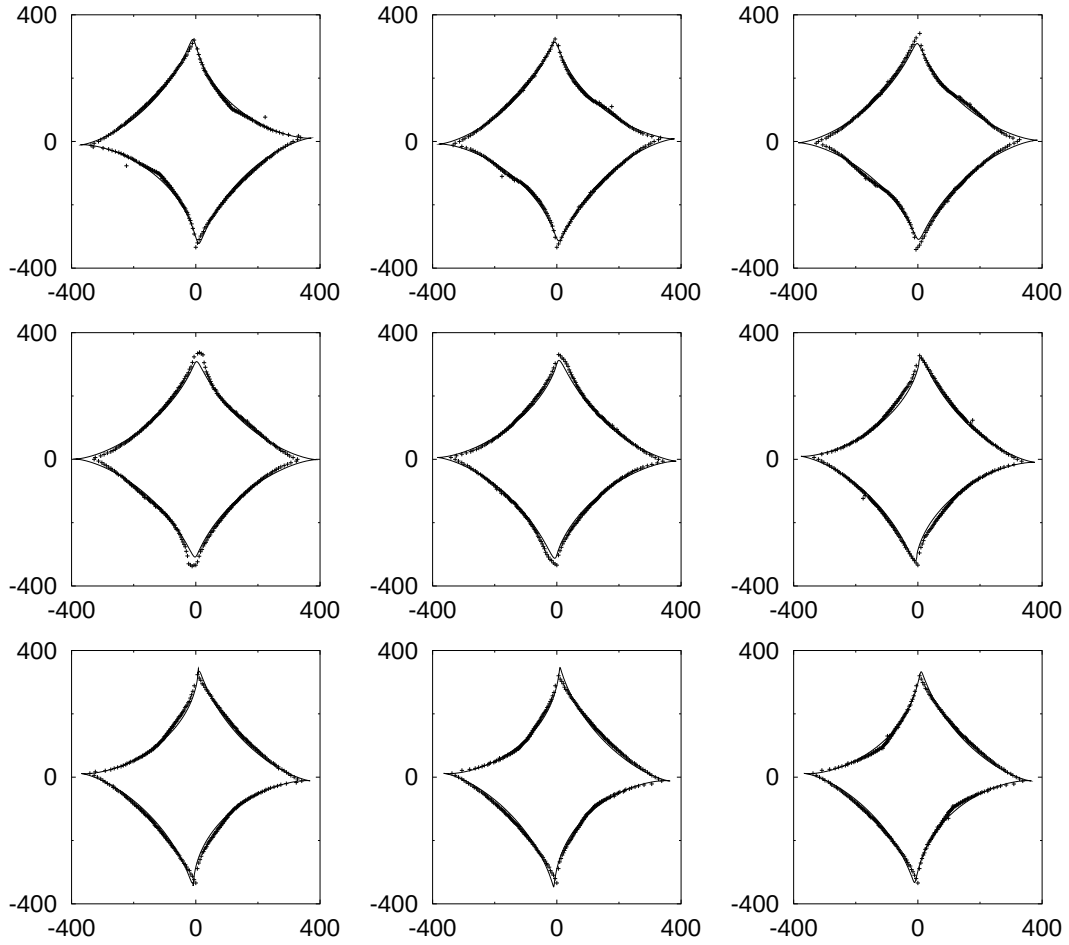


FIG. 8.11 – Neuf sections de la surface critique montrant les points de mesure (les mêmes que pour la figure 8.10) et la courbe ajustée. Toutes les sections droites partagent le même axe vertical y_0 aligné avec la pointe de la surface critique. Les champs sont en mT. Les plans de coupe successifs sont séparés de 20° .

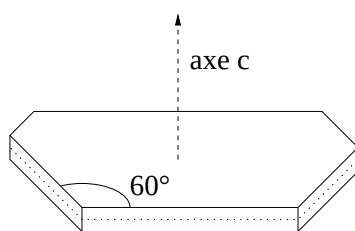


FIG. 8.12 – Cristal de forme hexagonale non régulier où les facettes sont parallèles aux plans cristallographiques. La ligne pointillée le long des facettes latérales représente le plan d'intersection du cristal avec son plan de symétrie.

Chapitre 9

Mesures dynamiques

Afin de déterminer le rôle de l'activation thermique dans le retournement de l'aimantation, nous effectuons plusieurs mesures de champ de retournement en changeant la vitesse de balayage du champ ainsi que la température. Toutes ces mesures se font pour une même trajectoire du champ appliqué. Autrement dit, nous testons toujours le même point de la surface critique. Les variations observées dans le champ de retournement sont faibles : en général inférieures à 1 %. Nous pouvons toutefois déterminer les champs de retournement avec suffisamment de précision pour extraire de leurs variations les informations sur la dynamique du renversement.

9.1 Champs de retournement

La figure 9.1 montre des mesures dynamiques effectuées sur une particule de ferrite de baryum en mode froid. La dépendance du champ de retournement en fonction de la température et la vitesse de balayage a été ajustée à la loi de Kurkijärvi démontrée au chapitre 2. De ces mesures et de leur ajustement on a pu tirer une estimation de l'ordre de grandeur de τ_0 : 26 ns. On obtient aussi le préfacteur E_0 de l'équation 2.8 : $E_0/k_B = 6,7 \cdot 10^5 \text{ KT}^{-3/2}$. À partir de ce préfacteur, on peut extrapoler la hauteur de la barrière à champ nul qui est d'environ $E_0 H_{sw}^{3/2} \approx 37000 \text{ K}$.

Comme nous avons vu dans la section 6.1, le champ appliqué accuse toujours un retard de quelques 10 ms par rapport à la consigne. Ce retard est en principe corrigé, malheureusement nous n'avons aucun moyen d'en connaître la valeur précise. Ainsi, une sous- ou sur-correction de ce retard se traduit par une dépendance linéaire par rapport à la vitesse de balayage des champs de retournement mesurés. La vitesse de balayage étant tracée en échelle logarithmique, cette erreur est visible sur la figure comme une divergence exponentielle des champs mesurés pour les vitesses de balayage les plus élevées. Heureusement la loi de Kurkijärvi prédit une dépendance pratiquement logarithmique par rapport à la vitesse de balayage, ce qui rend cette

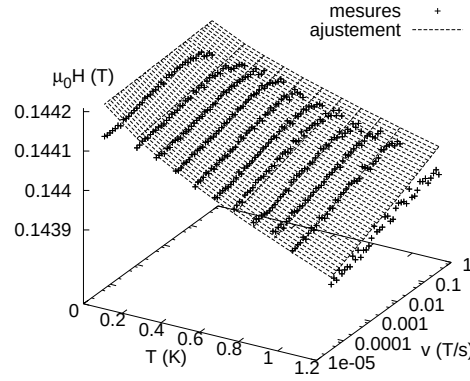


FIG. 9.1 – Champs de retournement de la particule du siècle en fonction de la température et de la vitesse de balayage du champ.

dépendance facile à séparer de l'artefact dû au retard du champ. Ainsi, on peut s'affranchir de ce problème en rajoutant un terme linéaire en v dans la loi qu'on ajuste.

Les champs représentés sur la figure 9.1 sont en fait des valeurs médianes calculées sur un ensemble de réalisations du saut d'aimantation. On utilise la médiane plutôt que la moyenne car la première est plus robuste vis-à-vis des événements rares. Étant donné que les distributions de probabilité de retournement ne sont pas gaussiennes, une telle robustesse est appréciable. Une estimation empirique de cette distribution peut être obtenue en traçant un histogramme de champs de retournement pour une température et une vitesse de balayage fixés. La figure 9.2 montre de tels histogrammes mesurés sur une particule de ferrite de baryum.

On observe que plus la température est basse, plus le champ de retournement moyen est proche de H_c et plus la distribution est étroite. On voit aussi que ces distributions sont asymétriques, bien plus étendues du côté des grands δH . Sur cette particule on voit aussi que la distribution ne change que très peu en dessous de 0,2 K.

9.2 Dépendance par rapport au point de la surface critique

Cette étude dynamique peut être reconduite pour chaque point de la surface critique [54, 55]. Du moins dans le plan du SQUID car les mesures dynamiques en trois dimensions sont plus longues et difficiles à réaliser. Nous avons effectué une

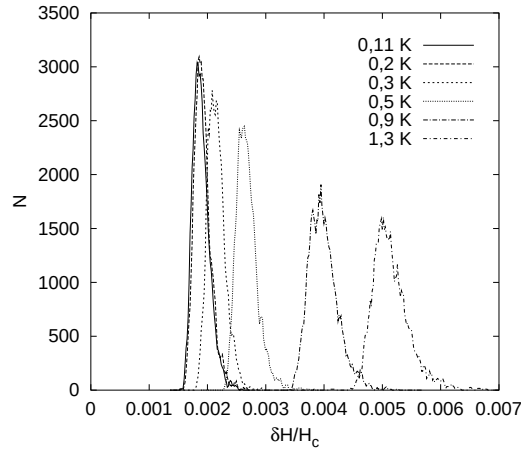


FIG. 9.2 – Histogrammes des champs de retournement à plusieurs températures d’une particule de ferrite de baryum. La grandeur en abscisse ($\delta H/H_c$) est un champ réduit.

telle étude sur une nanoparticule de ferrite de baryum. La figure 9.3 montre les différentes valeurs de τ_0 mesurées. Sur cette figure l’abscisse représente le point de la courbe critique considéré. Cette courbe n’est autre que l’intersection de la surface critique avec le plan du SQUID.

Théoriquement, τ_0 dépend des courbures du potentiel, et donc du point considéré le long de la courbe critique. À partir de la détermination expérimentale du potentiel, nous avons calculé le préfacteur géométrique de τ_0 suivant les deux hypothèses possibles pour le régime d’amortissement. La figure 9.4 montre ces calculs.

9.3 Bruit télégraphique

Lorsqu’on applique un champ très proche d’un point à bifurcation supercritique, le potentiel présente un double puits séparé par une toute petite barrière. Ceci est mis en évidence dans la représentation du potentiel C de la figure 1.7. En choisissant convenablement la position du champ, on peut avoir une barrière dont la hauteur permet au système de la franchir en un temps mesurable. Dans ce cas on s’attend à voir le système sauter spontanément la barrière dans un sens et dans l’autre, suivant une statistique de Poisson. Ce « bruit télégraphique » a déjà été observé sur des assemblées de particules [56]. Nous l’avons mesuré sur une particule de Cobalt individuelle en utilisant le point supercritique correspondant à la direction de difficile aimantation. La figure 9.5 montre un échantillon de cette mesure. L’histogramme des temps d’attente dans l’un des états montre que cette distribution de temps est bien exponentielle, confirmant ainsi la statistique qu’on attend.

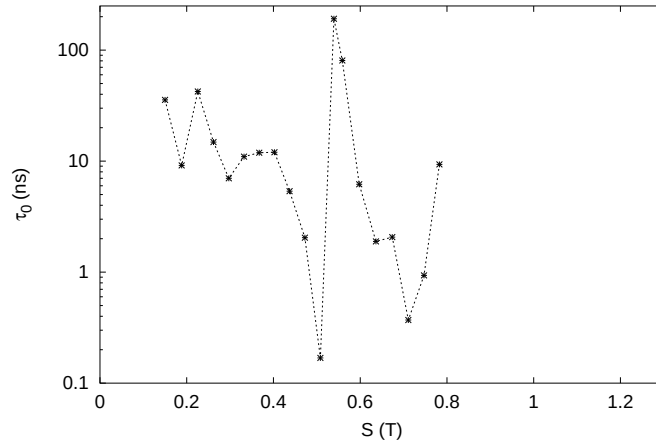


FIG. 9.3 – τ_0 mesuré sur une particule de BaFeO. L'axe des abscisses représente l'abscisse curviligne s du point considéré le long de la courbe critique.

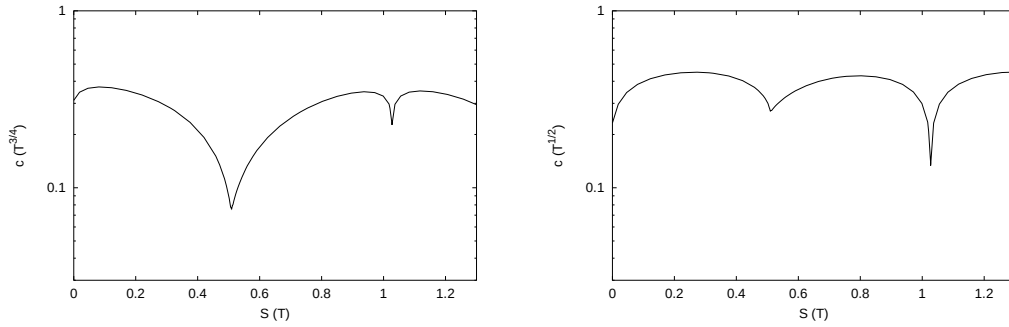


FIG. 9.4 – τ_0 théorique suivant le régime d'amortissement. À gauche pour amortissement faible. À droite pour un amortissement moyen ou fort.

Nous avons vu que près d'un point super-critique les sauts d'aimantation sont très petits. Nous avons donc du procéder de manière indirecte pour obtenir cette mesure. La méthode est légèrement différente au mode aveugle. Elle tire profit du fait que le courant critique dans le SQUID dépend de manière sensible de l'état de l'aimantation dans la particule. Du moins quand on n'est pas trop près du point super-critique. La procédure de mesure est donc la suivante.

- avec le SQUID éteint, le champ est positionné près du point super-critique ;
- il est laissé sur place pendant une seconde ;
- le champ est ramené à une valeur où les deux états d'aimantation sont discernables avec le SQUID ;
- on allume le SQUID pour déterminer l'état de l'aimantation ;
- on l'éteint et on recommence.

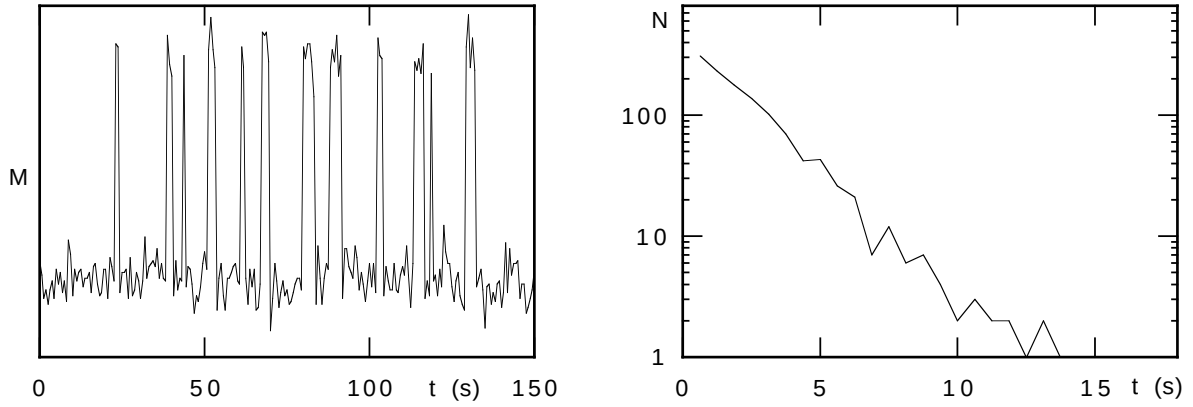


FIG. 9.5 – À gauche : bruit télégraphique sur la particule du siècle au voisinage d'un point super-critique. À droite : histogramme des temps d'attente dans l'état haut.

Ainsi, le SQUID est éteint pendant le temps où on permet au système de franchir la barrière par activation thermique, ce qui évite de le chauffer et de le perturber.

Finalement, cette mesure présente surtout un caractère d'illustration du comportement près d'un point super-critique. Lorsqu'on mesure la dynamique sur toute une série de points le long de la courbe critique, la mesure du bruit télégraphique ne semble pas très intéressante. Ce n'est en effet rien qu'un point de mesure de plus, et de plus un point particulièrement cher en temps de mesure. Nous n'avons donc pas continué l'étude du bruit télégraphique au delà de cette illustration.

Chapitre 10

Discussion

10.1 Interactions avec le milieu

Les particules que nous mesurons sont posées sur des SQUID. Dans le mode de mesure froid, qui est celui utilisé pour les mesures dynamiques, un courant électrique traverse le SQUID pendant la durée de la mesure. On peut s'interroger sur l'effet de ce courant, et du champ qu'il crée, sur le retournement de l'aimantation. Il semble dans la pratique que cet effet soit assez faible. On s'attend à ce que cet effet se manifeste au premier ordre juste comme un décalage des champs de retournement mesurés correspondant précisément au champ créé par le SQUID au niveau de la particule. Quand on fait des mesures on observe bien un tel décalage qui se manifeste par le fait que le centre de symétrie des astroïdes mesurées ne se trouve pas exactement à champ nul. De plus, ce décalage semble être proportionnel au courant qui traverse le SQUID, ce qui confirme cette hypothèse. Pourtant, le décalage est toujours très petit, de l'ordre du millitesla. Il est futile de le corriger quand on réalise des mesures quasi-statiques, mais on en tient compte lors des mesures dynamiques. Si on regarde au delà du premier ordre, il faudrait tenir compte de la non homogénéité du champ au niveau de la particule. Toutefois, cet effet au deuxième ordre doit être sensiblement plus faible que celui au premier ordre, et on a déjà constaté que l'effet au premier ordre est très faible.

Une autre forme d'interaction avec le milieu se manifeste sous la forme de l'anisotropie de surface. L'origine physique de cette anisotropie est le fait que l'échange est différent entre les atomes du cœur de la particule et ceux à sa surface, à cause de la différence de nature des voisins. L'anisotropie de surface est quelque chose de difficile à traiter, heureusement, on a des raisons de croire que ses effets ne sont pas très importants. Tout d'abord, dans une particule de 20 nm de diamètre, la proportion des atomes de surface est de l'ordre de 8 %. C'est une proportion respectable mais encore petite par rapport à la proportion des atomes de cœur. Ensuite, si l'écart à la

sphéricité de la forme de la particule est faible, l'anisotropie de surface doit s'annuler au premier ordre. Enfin, cette anisotropie est souvent modélisée comme un terme quadratique en l'aimantation. Si elle était vraiment quadratique on n'aurait pas à en tenir compte, puisqu'elle serait équivalente à une contribution supplémentaire à l'anisotropie de forme. Il n'y a pas pourtant d'argument théorique pour dire qu'elle ne comporte pas des termes d'ordre supérieur. Tant que ces termes sont petits par rapport au terme principal, il n'est pas gênant de les négliger. Si on venait à montrer que ces termes peuvent être importants, il faudrait en tenir compte dans les expressions de l'anisotropie utilisées. Ceci pourrait devenir très compliqué car il n'y a aucune contrainte sur la symétrie de ces termes.

10.2 Hypothèse de l'aimantation uniforme

On a fait l'hypothèse de l'aimantation uniforme. C'est certainement faux, les spins de surface s'écartent certainement au moins un peu de la direction de ceux du centre, ne serait-ce qu'en raison de l'anisotropie de surface. Par conséquent, le module de l'aimantation n'est pas constant et dépend certainement de son orientation et du champ appliqué.

Pourtant ce n'est pas très grave tant que ces écarts ne sont pas importants. Dans un modèle beaucoup plus général, on peut considérer dans l'espace de toutes les configurations possibles de l'aimantation, le lieu des « configurations critiques » qui représentent un équilibre marginal pour un certain champ. Dans le modèle du retournement uniforme, l'espace des configurations possibles est assimilé à la sphère unité qui est aussi le lieu des configurations critiques. En fait c'est ce dernier qui est le plus important car la surface critique (lieu des champs critiques) est l'image de ce lieu par une application continue. Le modèle du retournement uniforme est équivalent à ce dernier quand l'échange est infini. Si on diminue progressivement l'échange, le lieu des configurations critiques va progressivement se déformer par rapport à la sphère unité. Tant que la déformation n'est pas trop grande, ceci est sans conséquences car on peut toujours associer bijectivement cette surface sur la sphère unité en associant chaque point au vecteur normal à la surface. Ceci n'est évidemment possible que pour des écarts suffisamment faibles. Si ces écarts sont trop grands, un tel mappage bijectif devient impossible.

Une telle théorie reposant sur des arguments topologiques reste encore à développer. Je ne serais pas étonné pourtant si elle s'avérait équivalente au modèle du retournement uniforme tant que le lieu des configurations critiques est homotope à une sphère¹. Ceci exclut bien sûr des circonstances dans lesquelles des modes de retournement tels que le curling peuvent avoir lieu.

¹c'est à dire image d'une sphère par une bijection bicontinue.

10.3 Spins nucléaires et dissipation

La dissipation joue un rôle fondamental dans le comportement dynamique des nanoparticules. On la voit sous la forme du coefficient α dans l'équation de Gilbert. Ce lien est mis en évidence par le théorème fluctuation-dissipation d'après lequel les fluctuations thermiques qui permettent le franchissement de la barrière de potentiel sont directement liées à la dissipation. Ce sont en effet deux manières de voir l'interaction du système avec le thermostat.

Dans l'équation de Gilbert, la dissipation apparaît comme un terme *ad hoc* qui apporte un amortissement du premier ordre dans la dynamique de l'aimantation. Ce terme correspond à des fluctuations se présentant comme une force aléatoire avec un temps de corrélation nul. Dans la réalité, le temps de corrélation du bruit thermique n'est pas nul, mais relié aux constantes de temps des interactions avec le réseau cristallin et avec les électrons de conduction. Quand ce temps devient de l'ordre de τ_0 , on ne peut plus le négliger.

Parmi les différentes sources de dissipation qui viennent à l'esprit, il y a l'interaction avec les phonons, les courants de Foucault dans un conducteur et le champ hyperfin. Pour les premiers, le temps de corrélation est de l'ordre de la période des ondes élastiques dans la particule. Celle-ci est majorée par le temps que met le son à traverser deux fois la particule, soit environ 10^{-11} s. Les temps de corrélation des ondes électroniques sont encore plus courts. On voit donc que, pour ces deux sources de dissipation, il est légitime de négliger le temps de corrélation. Reste le cas du champ hyperfin dont les variations se font dans des temps allant de quelques nanosecondes à une microseconde. Pour cette source de dissipation on ne peut plus négliger le temps de corrélation et le traitement de Fokker-Planck est inadapté.

Pourtant, quand les temps de corrélations sont assez longs, on peut se placer dans une autre limite. Si on a une force aléatoire avec un temps de corrélation τ_b très long devant τ_0 , on peut considérer que cette force est constante dans l'échelle de temps correspondant à la relaxation du système à l'intérieur du puits métastable. Cette force peut alors être considérée comme faisant partie du potentiel effectif et le traitement de Fokker-Planck devient applicable avec ce potentiel modifié. Les équations 2.6 et 2.7 du chapitre 2 sont à nouveau valables. On a donc la probabilité de retournement

$$-\frac{dP}{P} = \frac{dt}{\tau(E)}$$

où

$$\tau(E) = \tau_0 e^{E/k_B T},$$

avec pour seul changement le fait que E et donc τ sont maintenant des variables aléatoires.

Si le temps de corrélation du bruit τ_b est petit devant τ , on peut intégrer la première équation sur une échelle de temps intermédiaire δt sur laquelle P varie très peu et E a le temps d'explorer toute sa distribution. On a alors

$$\begin{aligned} -\frac{\delta P}{P} &= \int_{\delta t} \frac{dt}{\tau(E)} \\ &= \delta t \left\langle \frac{1}{\tau(E)} \right\rangle, \end{aligned}$$

où on suppose la moyenne temporelle identique à la moyenne sur toutes les réalisations de E . Cette équation est identique à (2.7) à ceci près qu'elle s'applique sur une échelle de temps plus grande et que τ est remplacé par la moyenne harmonique de $\tau(E)$. Ceci revient à avoir une barrière statique avec une hauteur efficace E^* donnée par

$$\begin{aligned} \frac{1}{\tau(E^*)} &= \left\langle \frac{1}{\tau(E)} \right\rangle \\ e^{-E^*/k_B T} &= \left\langle e^{-E/k_B T} \right\rangle. \end{aligned}$$

Si on suppose la distribution de E gaussienne de moyenne $\langle E \rangle$ et d'écart type ΔE , alors on a

$$E^* = \langle E \rangle - \frac{(\Delta E)^2}{k_B T}.$$

On voit donc que l'effet de ce bruit n'est pas sensible à haute température mais il se manifeste à basse température comme un abaissement de la hauteur effective de la barrière. Ceci pourrait se manifester expérimentalement par des champs de retournement qui, à basse température, sont plus petits que ceux prévus par la relation de Kurkijärvi en absence de fluctuations de la barrière. C'est donc un comportement à étudier si on veut le discerner clairement d'un éventuel effet tunnel qui se manifeste de la même manière. Malheureusement, comme le font remarquer Dube et Stamp [57], ces phénomènes sont difficiles à traiter car à basse température ils sont dominés par les événements rares qui constituent la queue de la distribution de E .

Bibliographie

- [1] J. L. Dormann, D. Fiorani, and E. Tronc. *Adv. Chem. Phys.*, 98 :283, 1997.
- [2] M. Ledermann, S. Schultz, and M. Ozaki. *Phys. Rev. Lett.*, 73 :1986, 1994.
- [3] W. Rave and J. Miltat et al. *IEEE Trans. Mag.*, 30(6) :4473, 1994.
- [4] V. Gros, Shan-Fab Lee, G. Faini, A. Cornette, A. Hamzic, and A. Fert. *J. Magn. Magn. Mat.*, 165 :512, 1997.
- [5] J. E. Wegrowe, S. E. Gilbert, D. Kelly, B. Doudin, and J.-Ph. Ansermet. *IEEE Trans. Mag.*, 34 :903, 1998.
- [6] A. D. Kent, S. von Molnar, S. Gider, and D. D. Awschalom. *J. Appl. Phys.*, 76 :6656, 1994.
- [7] J. G. S. Lok, A. K. Geim, J. C. Maan S. V. Dubonos, L. Theil Kuhn, and P. E. Lindelof. *Phys. Rev. B*, 58 :12201, 1998.
- [8] S. Wirth, M. Field, D. D. Awschalom, and S. von Moln/'ar. *Phys. Rev. B*, 57 :R14028, 1998.
- [9] S. A. Majetich and Y. Jin. Magnetization directions of individual nanoparticles. *Science*, 284 :470, 1999.
- [10] W. Wernsdorfer, B. Doudin, D. Mailly, K. Hasselbach, A. Benoit, J. Meier, J.-Ph. Ansermet, and B. Barbara. *Phys. Rev. Lett.*, 77 :1873, 1996.
- [11] S. Guéron, Mandar M. Deshmukh, E. B. Myers, and D. C. Ralph. *cond-mat/9904248*, 1999.
- [12] M. E. Schabes. *J. Magn. Magn. Mat.*, 95 :249, 1991.
- [13] W. Rave, K. Fabian, and A. Hubert. *J. Magn. Magn. Mat.*, 190 :332, 1998.
- [14] P. C. E. Stamp, E. M. Chudnovsky, and B. Barbara. *Int. J. Mod. Phys.*, B6 :1355–1478, 1992.
- [15] E. C. Stoner and E. P. Wohlfarth. *Philos. Trans. London Ser. A*, 240 :599, 1948. reprinted in *IEEE Trans. Magn. MAG-27*, 3475 (1991).
- [16] L. Néel. *C. R. Acad. Science*, 224 :1550, 1947.
- [17] Ching-Ray Chang. *J. Appl. Phys.*, 69 :2431, 1991.

- [18] T. Chang and J. G. Chu. *J. Appl. Phys.*, 75 :5553, 1994.
- [19] P. L. Fulmek and H. Hauser. *J. Appl. Phys.*, 76 :6561, 1994.
- [20] Paul Glendinning. *Stability, instability and chaos*. Cambridge University Press, 1994.
- [21] A. Thiaville. Coherent rotation of magnetization in three dimensions : a geometrical approach. to be published in *Phys. Rev. B*, 1999.
- [22] David Mukamel, Michael E. Fisher, and Eytan Domany. *Phys. Rev. Lett.*, 37 :565, 1976.
- [23] W. T. Coffey. *Adv. Chem. Phys.*, 103 :259, 1998.
- [24] P. Hänggi, P. Talkner, and M. Borkovec. *Rev. Mod. Phys.*, 62 :251, 1990.
- [25] L. Néel. *Ann. Geophys.*, 5 :99, 1949.
- [26] W. F. Brown. *J. Appl. Phys.*, 30 :130S, 1959.
- [27] W. F. Brown. *J. Appl. Phys.*, 34 :1319, 1963.
- [28] W. F. Brown. *Phys. Rev.*, 130 :1677, 1963.
- [29] W. F. Brown. *IEEE Trans. Mag.*, 15 :1196, 1979.
- [30] L. J. Geoghegan, W. T. Coffey, and B. Mulligan. *Adv. Chem. Phys.*, 100 :475, 1997.
- [31] E. D. Boerner and H. Neal Bertram. *IEEE Trans. Mag.*, 33 :3052, 1997.
- [32] R. Arias and H. Neal Bertram. *J. Magn. Magn. Mat.*, 171 :209, 1997.
- [33] D. Garcia-Pablos, P. Garcia-Mochales, and N. Garcia. *J. Appl. Phys.*, 79 :6021, 1996.
- [34] H. L. Richards, S. W. Sides, M. A. Novotny, and P. A. Rikvold. *J. Appl. Phys.*, 79 :5749, 1996.
- [35] W. T. Coffey, D. S. F. Crothers, J. L. Dormann, Yu. P. Kalmykov, E. C. Kennedy, and W. Wernsdorfer. *Phys. Rev. Lett.*, 80 :5655, 1998.
- [36] A. Thiaville. *J. Magn. Magn. Mat.*, 182 :5–18, 1998.
- [37] J. Kurkijärvi. *Phys. Rev. B*, 6 :832, 1972.
- [38] R. H. Victora. *Phys. Rev. Lett.*, 63 :457, 1989.
- [39] R. H. Victora. Springer proceedings in physics. volume 50, page 11, 1990.
- [40] A. Garg. *Phys. Rev. B*, 51 :15592, 1995.
- [41] M. Ketchen, D. J. Pearson, K. Stawiaasz, C-H. Hu, A. W. Kleinsasser, T. Brunner, C. Cabral, V. Chandrasekhar, M. Jaso, M. Manny, K. Stein, and M. Bhushan. *IEEE Applied Superconductivity*, 3 :1795, 1993.
- [42] D. Mailly, C. Chapelier, and A. Benoit. *Phys. Rev. Lett.*, 70 :2020, 1993.

- [43] G. Cernicchiaro, K. Hasselbach, D. Mailly, W. Wernsdorfer, and A. Benoit. Nato asi series. In Bernard Kramer, editor, *Quantum Transport in Semiconductor Submicron Structures*, volume 326 of *Series E : Applied Sciences*. 1996.
- [44] G. Cernicchiaro. PhD thesis, Joseph Fourier University, Grenoble, 1997.
- [45] O. V. Lounasmaa. *Experimental principles and methods below 1 K*. Academic Press, London, 1974.
- [46] T. Van Duzer and C. W. Turner. *Principles of superconductive devices and circuits*. Edward Arnold, London, 1981.
- [47] A. Benoit, D. Mailly, M. El-Khatib, and P. Perrier. *Quantum Coherence in Mesoscopic Systems*, page 254 and 237. Plenum Press, New York, 1991.
- [48] W. Wernsdorfer, K. Hasselbach, D. Mailly, B. Barbara, A. Benoit, L. Thomas, and G. Suran. *J. Magn. Magn. Mat.*, 145 :33, 1995.
- [49] W. Wernsdorfer. PhD thesis, Joseph Fourier University, Grenoble, 1996.
- [50] C. Guerret-Piécourt, Y. Le Bouar, A. Loiseau, and H. Pascard. *Nature*, 372 :761, 1994.
- [51] N. Demoncy, O. Stéphan, N. Brun, C. Colliex, A. Loiseau, and H. Pascard. *Eur. Phys. J. B*, 4 :147, 1998.
- [52] E. Bonet, W. Wernsdorfer, B. Barbara, K. Hasselbach, A. Benoît, and D. Mailly. *IEEE Trans. Mag.*, 34 :979, 1998.
- [53] E. Bonet, W. Wernsdorfer, B. Barbara, A. Benoît, D. Mailly, and A. Thiaville. Three dimensional magnetization reversal measurements in nanoparticles. to be published in *Phys. Rev. Lett.*, 1999.
- [54] W. Wernsdorfer, E. Bonet Orozco, K. Hasselbach, A. Benoit, D. Mailly, O. Kubo, H. Nakano, and B. Barbara. *Phys. Rev. Lett.*, 79 :4014, 1997.
- [55] W. Wernsdorfer, E. Bonet, B. Barbara, A. Benoît, D. Mailly, N. Demoncy, H. Pascard, O. Kubo, and H. Nakano. *IEEE Trans. Mag.*, 34 :973, 1998.
- [56] F. Coppinger, J. Genoe, D. K. Maude, U. Genner, J. C. Portal, K. E. Singer, P. Rutter, T. Taskin, A. R. Peaker, and A. C. Wright. *Phys. Rev. Lett.*, 75 :3513, 1995.
- [57] M. Dube and P. C. E. Stamp. *J. Low T. Phys.*, 113 :1085, 1998. cond-mat/9908423.
- [58] I. M. L. Billas, A. Châtelain, and W. A. de Heer. *J. Magn. Magn. Mat.*, 168 :64, 1997.
- [59] O. B. Zaslavskii. *Phys. Rev. B*, 42 :992–993, july 1990.
- [60] Nato asi. In L. Gunther and B. Barbara, editors, *Quantum Tunneling of Magnetization – QTM'94*, volume 301 of *Series E : Applied Sciences*, pages 227–241. 1995.

- [61] M. C. Miguel and E. M. Chudnovsky. *Phys. Rev. B*, 54 :388–394, july 1996.
- [62] Gwang-Hee Kim and Dae Sung Hwang. *Phys. Rev. B*, 55 :8918–8927, april 1997.
- [63] L. Thomas, F. Lioni, R. Ballou, D. Gatteschi, R. Sessoli, and B. Barbara. *Nature*, 145 :383, 1996.
- [64] W. Wernsdorfer and R. Sessoli. *Science*, 284 :133, 1999.
- [65] A. J. Freeman and R. Wu. *J. Magn. Magn. Mat.*, 100 :497, 1991.

Conclusion

Pour décrire le retournement de l'aimantation dans les particules nanométriques, le modèle de Stoner-Wohlfarth est encore très largement employé. Ce modèle, d'une simplicité exemplaire, se généralise naturellement par la considération d'une anisotropie arbitraire en trois dimensions. Ceci reste encore assez simple puisque la particule est encore entièrement décrite par un potentiel défini sur la sphère unité et qu'on choisit en général très régulier. On a pu voir cependant que, malgré cette simplicité, on peut trouver des lieux des champs de retournement formant des surfaces critiques de formes complexes. La surface critique est toujours formée de plusieurs nappes et celles-ci peuvent se croiser en donnant lieu à des hystérésis très compliquées. D'un autre côté, l'étude de ces surfaces critiques est riche en enseignements puisqu'elle permet de remonter à l'anisotropie. En connaissant celle-ci on peut connaître, pour un champ appliqué donné, la hauteur des barrières d'énergie ainsi que les courbures du potentiel au fond des puits et dans les cols.

Ces informations sur la forme du potentiel sont intéressantes pour des études dynamiques. À température non nulle, le modèle de Néel-Brown décrit l'activation thermique au dessus de la barrière de potentiel à l'aide de ces paramètres. Comme il est expérimentalement peu commode d'utiliser une barrière d'énergie de hauteur fixe, on a recours à la formule de Kurkijärvi qui donne le champ moyen de retournement à partir des paramètres du modèle de Néel-Brown. On peut ainsi, en comparant des mesures à cette formule, vérifier si on est bien dans le régime d'activation thermique tel que décrit par ce modèle et trouver la valeur du temps τ_0 caractéristique de la dynamique du système.

Afin de pouvoir comparer des particules réelles à ces prédictions théoriques, la technique de magnétométrie à micro-SQUID a été développée. L'informatisation de l'expérience permet maintenant une très grande souplesse dans les modes de mesure. Une telle souplesse dans le choix des chemins du champ appliqué était nécessaire pour pouvoir étudier des particules à anisotropie complexe. La possibilité d'effectuer des mesures en trois dimensions ouvre la possibilité de déterminer expérimentalement des anisotropies de particules individuelles.

Au vu des résultats, on a la confirmation que les prédictions les plus surprenantes du modèle de retournement uniforme sont vérifiées sur des systèmes réels. Ainsi,

on a vu une surface critique présentant plusieurs intersections et demandant des chemins bien choisis dans l'espace des champs pour être complètement explorée. Un tel cas aurait pu passer pour une simple vue de l'esprit avant cette confirmation expérimentale. Le développement théorique trouve ainsi sa justification. On a vu aussi qu'on arrive à mesurer des surfaces critiques complètes en trois dimensions et à partir de là à déterminer l'anisotropie effective de la particule mesurée. Au niveau dynamique, le modèle de Néel-Brown décrit très bien nos systèmes pour des températures assez élevées. À très basse température, des déviations par rapport à ce modèle font penser à un possible franchissement de la barrière par effet tunnel. Il faut toutefois être prudent car il faut séparer cet effet des possibles fluctuations non quantiques de la hauteur de la barrière qui produisent des déviations dans le même sens. L'observation d'un effet tunnel résonnant serait la meilleure confirmation qu'on voit bien un effet tunnel mésoscopique.

Les études des nanoparticules au micro-SQUID sont loin d'avoir dit leur dernier mot. En mettant les nanoparticules à l'intérieur même du micro-pont plutôt qu'au dessus de celui-ci, il est possible d'atteindre des sensibilités beaucoup plus élevées. Cette expérience vient d'être réussie dans notre équipe et on a vu qu'il est possible maintenant de faire des mesures magnétiques sur des particules de cobalt de seulement 3 nm de diamètre. À cette échelle de taille les nanoparticules se voient en général attribuer le nom d'*agrégats*. Des études du magnétisme des agrégats ont déjà été faites sur des assemblée [58] et la possibilité d'étudier des agrégats individuels ouvre des perspectives très intéressantes pour la recherche de l'effet tunnel de l'aimantation. Ces particule étant sensiblement plus petites que celles étudiées jusqu'ici, elles devraient aussi être sensiblement plus « quantiques » [59, 60, 61, 62]. Avec un moment magnétique de l'ordre de $10^3 \mu_B$, le caractère discret du spectre d'énergie de son aimantation pourrait devenir apparent. Verra-t-on un effet tunnel résonnant à l'image de celui observé dans des systèmes moléculaires tels que Mn_{12} [63] ou Fe_8 [64]? Cette question est pour le moment ouverte et la recherche de sa réponse promet d'être très enrichissante.

L'utilisation de ces nouvelles particules encore plus petites et enfouies à l'intérieur du SQUID pose aussi des nouvelles questions quand à la modélisation du système. Avec près de la moitié des atomes à la surface de la particule, l'anisotropie de surface n'est plus un terme qu'on peut espérer négliger [65]. Même si à l'ordre deux en $\mathbf{r}$ sa contribution ne pose pas de problème, les termes d'ordre supérieur sont très probablement à prendre en compte, puisqu'ils peuvent être comparables et même plus grands que ceux aux mêmes ordres de l'anisotropie cristalline. Une autre question qui se pose est celle des effets de proximité. Nous avons là en effet un petit morceau de matériau ferromagnétique enfoui au cœur même d'un supraconducteur. Près de l'interface, il faut se demander si on peut avoir des paires de Cooper au sein de la particule ferromagnétique ou des électrons non couplés et polarisés au sein du supraconducteur. Autant de questions qu'il faudra se poser pour savoir si le modèle

d'une particule isolée n'interagissant qu'avec un champ extérieur reste pertinent.

Annexes

Annexe A

Compléments mathématiques

A.1 Quelques anisotropies

Nous allons considérer quelques anisotropies simples et, pour chacune d'elles, voir quelles sont ses dérivées. Une anisotropie réelle sera en général la superposition de plusieurs anisotropies simples.

A.1.1 Anisotropie quadratique

Elle s'écrit

$$V_0 = K_x x^2 + K_y y^2 + K_z z^2$$

mais on peut annuler arbitrairement l'un des termes puisque $x^2 + y^2 + z^2 = 1$ est une constante. Ses dérivées sont

$$g = 2 \begin{pmatrix} xK_x \\ yK_y \\ zK_z \end{pmatrix} \quad q = \begin{pmatrix} K_x & 0 & 0 \\ 0 & K_y & 0 \\ 0 & 0 & K_z \end{pmatrix} \quad c = \mathbf{0}$$

A.1.2 Anisotropie cubique

À l'ordre 4, on l'écrit

$$V_0 = K_c(x^2 y^2 + y^2 z^2 + z^2 x^2)$$

dont les dérivées donnent

$$g = 2K_c \begin{pmatrix} x(y^2 + z^2) \\ y(z^2 + x^2) \\ z(x^2 + y^2) \end{pmatrix} \quad q = K_c \begin{pmatrix} y^2 + z^2 & 2xy & 2xz \\ 2yx & z^2 + x^2 & 2yz \\ 2zx & 2zy & x^2 + y^2 \end{pmatrix}$$

$$c = \frac{2K_c}{3} \left(\begin{array}{ccc|ccc|ccc} 0 & y & z & y & x & 0 & z & 0 & x \\ y & x & 0 & x & 0 & z & 0 & z & y \\ z & 0 & x & 0 & z & y & x & y & 0 \end{array} \right)$$

A.1.3 Uniaxial d'ordre 4

Elle s'écrit

$$V_0 = K_4(x^2 + y^2)^2$$

et ses dérivées sont

$$g = 4K_4(x^2 + y^2) \begin{pmatrix} x \\ y \\ 0 \end{pmatrix} \quad q = 2K_4 \begin{pmatrix} 3x^2 + y^2 & 2xy & 0 \\ 2yx & 3y^2 + x^2 & 0 \\ 0 & 0 & 0 \end{pmatrix}$$

$$c = \frac{4K_4}{3} \left(\begin{array}{ccc|ccc|ccc} 3x & y & 0 & y & x & 0 & 0 & 0 & 0 \\ y & x & 0 & x & 3y & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{array} \right)$$

A.1.4 Uniaxial d'ordre 6

Elle s'écrit

$$V_0 = K_6(x^2 + y^2)^3$$

et ses dérivées sont

$$g = 6K_6(x^2 + y^2)^2 \begin{pmatrix} x \\ y \\ 0 \end{pmatrix} \quad q = 3K_6(x^2 + y^2) \begin{pmatrix} 5x^2 + y^2 & 4xy & 0 \\ 4yx & 5y^2 + x^2 & 0 \\ 0 & 0 & 0 \end{pmatrix}$$

$$c = 4K_6 \left(\begin{array}{ccc|ccc|ccc} 5x^3 + 3xy^2 & 3x^2y + y^3 & 0 & 3x^2y + y^3 & x^3 + 3xy^2 & 0 & 0 & 0 & 0 \\ 3x^2y + y^3 & x^3 + 3xy^2 & 0 & x^3 + 3xy^2 & 3x^2y + 5y^3 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \end{array} \right)$$

A.2 Changements de repère orthonormé

Nous avons besoin d'utiliser quatre repères différents :

- celui de l'expérience, dans lequel sont mesurés les points expérimentaux ;
- celui qui diagonalise la partie quadratique de l'anisotropie ;
- celui des axes cristallins de la particule, dans lequel les anisotropies d'ordre supérieur s'écrivent simplement ;
- le repère « local » dans lequel $\mathbf{r}$ est suivant z , qui change pour chaque point calculé.

quelques rappels sur les matrices de changement de base peuvent donc être utiles.

Soient deux bases orthonormées A et B , la matrice de passage p^{AB} de A vers B a pour colonnes les coordonnées dans la base A des vecteurs de la base B . Comme son nom ne l'indique pas, elle permet de retrouver les coordonnées d'un vecteur dans la base A à partir de ses coordonnées dans la base B . En effet, si le vecteur $\mathbf{x}$ a pour coordonnées x^A dans A et x^B dans B , alors, en utilisant la convention d'Einstein de sommation sur les indices répétés,

$$x_i^A = p_{ij}^{AB} x_j^B$$

Comme nous utilisons des bases orthonormées, les matrices de passage sont orthogonales. Leur inverse est simplement leur transposée :

$$p_{ij}^{BA} = p_{ji}^{AB} \quad p_{ij}^{AB} p_{jk}^{AB} = \delta_{ik}$$

A.2.1 Formes linéaires, quadratiques et cubiques

Si g est une forme linéaire et g^A ses coordonnées dans la base A , alors

$$g(\mathbf{x}) = g_i^A x_i^A$$

Les coordonnées de g dans B sont reliées à celles dans la base A par la même formule de changement de base que pour les vecteurs :

$$g_i^A = p_{ij}^{AB} g_j^B.$$

Ceci n'est vrai que parce que p^{AB} est une matrice orthogonale (elle est inverse de sa transposée). En effet, quel que soit $\mathbf{x}$,

$$\begin{aligned} g(\mathbf{x}) &= g_j^B x_j^B \\ &= g_j^B p_{ji}^{BA} x_i^A \\ &= (p_{ij}^{AB} g_j^B) x_i^A \\ &= g_i^A x_i^A \end{aligned}$$

Ceci se généralise naturellement au cas d'une forme quadratique q et une forme cubique c :

$$\begin{aligned} g(\mathbf{x}) &= g_i^A x_i^A & g_i^A &= p_{ij}^{AB} g_j^B \\ q(\mathbf{x}) &= q_{ij}^A x_i^A x_j^A & q_{ij}^A &= p_{ik}^{AB} p_{jl}^{AB} q_{kl}^B \\ c(\mathbf{x}) &= c_{ijk}^A x_i^A x_j^A x_k^A & c_{ijk}^A &= p_{il}^{AB} p_{jm}^{AB} p_{kn}^{AB} c_{lmn}^B \end{aligned}$$

A.2.2 Base propre d'une forme quadratique

Considérons, à deux dimensions, une forme quadratique q avec une valeur propre nulle. Nous avons ses coordonnées dans une base A et nous cherchons la base B qui diagonalise q . Ceci revient à trouver θ tel que

$$p^{AB} = \begin{pmatrix} \cos \theta & -\sin \theta \\ \sin \theta & \cos \theta \end{pmatrix}$$

Dans la base B , la forme quadratique s'écrit

$$q^B = \begin{pmatrix} 0 & 0 \\ 0 & \lambda \end{pmatrix}$$

En utilisant la formule de changement de base pour une forme quadratique, on trouve que dans la base A elle s'écrit

$$q^A = \lambda \begin{pmatrix} \sin^2 \theta & -\cos \theta \sin \theta \\ -\cos \theta \sin \theta & \cos^2 \theta \end{pmatrix}$$

Par identification, on déduit

$$\lambda = q_{11}^A + q_{22}^A \quad \cos 2\theta = \frac{q_{22}^A - q_{11}^A}{\lambda} \quad \sin 2\theta = \frac{-2q_{12}^A}{\lambda}$$

A.2.3 Angles d'Euler

En général on utilise les angles d'Euler pour exprimer les changements de base. Il faut pouvoir passer de ces angles aux matrices de changement de base et réciproquement.

On peut remarquer que la la matrice de passage p^{AB} de A vers B est la même que la matrice $r^{A \rightarrow B}$ qui décrit dans la base A la rotation qui fait passer les axes de A sur ceux de B . En effet, considérons $\mathbf{x}$ un vecteur et $\mathbf{y}$ son image par cette rotation. Notons x^A , x^B , y^A et y^B leurs coordonnées dans A et B . Il est clair que $\mathbf{x}$

a dans A les mêmes coordonnées que $\mathbf{y}$ dans B , donc $X_A = Y_B$. Par ailleurs, d'après les définitions de p^{AB} et $r^{A \rightarrow B}$,

$$\begin{aligned} y_i^A &= p_{ij}^{AB} y_j^B \\ &= r_{ij}^{A \rightarrow B} x_j^A \end{aligned}$$

donc, comme c'est vrai quel que soit $x^A = y^B$, on a

$$p^{AB} = r^{A \rightarrow B}.$$

Supposons que (θ, φ, ψ) sont les angles d'Euler associés à la matrice de passage $P_{\theta, \varphi, \psi} = P_{AB}$. La façon la plus intuitive de décrire ce changement de base est de dire que pour porter le repère A sur le repère B il faut effectuer dans l'ordre les rotations suivantes :

- d'un angle φ autour de l'axe z ;
- d'un angle θ autour du nouvel axe x ;
- d'un angle ψ autour du nouvel axe z .

On peut exprimer ces rotations à l'aide de matrices de rotation. Si on note

$$\tilde{x} = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & -1 \\ 0 & 1 & 0 \end{pmatrix} \quad \tilde{y} = \begin{pmatrix} 0 & 0 & 1 \\ 0 & 0 & 0 \\ -1 & 0 & 0 \end{pmatrix} \quad \tilde{z} = \begin{pmatrix} 0 & -1 & 0 \\ 1 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}$$

alors, une rotation d'un angle θ autour de chacun des axes x, y, z s'écrit

$$\begin{aligned} e^{\theta \tilde{x}} &= \begin{pmatrix} 1 & 0 & 0 \\ 0 & \cos \theta & -\sin \theta \\ 0 & \sin \theta & \cos \theta \end{pmatrix} & e^{\theta \tilde{y}} &= \begin{pmatrix} \cos \theta & 0 & \sin \theta \\ 0 & 1 & 0 \\ -\sin \theta & 0 & \cos \theta \end{pmatrix} \\ e^{\theta \tilde{z}} &= \begin{pmatrix} \cos \theta & -\sin \theta & 0 \\ \sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{pmatrix} \end{aligned}$$

Avec ces notations, on peut exprimer ces rotations de la manière suivante :

- rotation d'un angle φ autour de l'axe z : $e^{\varphi \tilde{z}}$;
- rotation d'un angle θ autour du *nouvel* axe x : $e^{\varphi \tilde{z}} e^{\theta \tilde{x}} e^{-\varphi \tilde{z}}$; la composée de ces deux rotations donne

$$e^{\varphi \tilde{z}} e^{\theta \tilde{x}} e^{-\varphi \tilde{z}} = e^{\varphi \tilde{z}} e^{\theta \tilde{x}}$$

- rotation d'un angle ψ autour du *nouvel* axe z : $(e^{\varphi \tilde{z}} e^{\theta \tilde{x}}) e^{\psi \tilde{z}} (e^{-\theta \tilde{x}} e^{-\varphi \tilde{z}})$; la composée des trois rotations donne

$$\begin{aligned} \mathcal{P}_{\theta, \varphi, \psi} &= e^{\varphi \tilde{z}} e^{\theta \tilde{x}} e^{\psi \tilde{z}} e^{-\theta \tilde{x}} e^{-\varphi \tilde{z}} e^{\theta \tilde{x}} \\ &= e^{\varphi \tilde{z}} e^{\theta \tilde{x}} e^{\psi \tilde{z}} \end{aligned}$$

On peut donc aussi décrire ces rotations en utilisant toujours les mêmes axes (ceux de A) l'ordre des rotations est alors inversé : on tourne

- d'un angle ψ autour de z ;
- d'un angle θ autour de x ;
- d'un angle φ autour de z .

La matrice de passage est donc

$$\begin{aligned}
 \mathcal{P}_{\theta,\varphi,\psi} &= e^{\varphi\tilde{z}} e^{\theta\tilde{x}} e^{\psi\tilde{z}} \\
 &= \begin{pmatrix} \cos \varphi & -\sin \varphi & 0 \\ \sin \varphi & \cos \varphi & 0 \\ 0 & 0 & 1 \end{pmatrix} \begin{pmatrix} 1 & 0 & 0 \\ 0 & \cos \theta & -\sin \theta \\ 0 & \sin \theta & \cos \theta \end{pmatrix} \begin{pmatrix} \cos \psi & -\sin \psi & 0 \\ \sin \psi & \cos \psi & 0 \\ 0 & 0 & 1 \end{pmatrix} \\
 &= \begin{pmatrix} \cos \varphi \cos \psi - \sin \varphi \cos \theta \sin \psi & -\cos \varphi \sin \psi - \sin \varphi \cos \theta \cos \psi & \sin \varphi \sin \theta \\ \sin \varphi \cos \psi + \cos \varphi \cos \theta \sin \psi & -\sin \varphi \sin \psi + \cos \varphi \cos \theta \cos \psi & -\cos \varphi \sin \theta \\ \sin \theta \sin \psi & \sin \theta \cos \psi & \cos \theta \end{pmatrix}
 \end{aligned}$$

Annexe B

Mode d'emploi de la carte d'acquisition

B.1 Description de la carte

Il s'agit d'une carte d'acquisition qui s'enfiche dans un slot PCI. Elle a été prévue pour être utilisée sur un PowerMac, au moins de deuxième génération. Bien que pouvant à priori être utilisée sur n'importe quel ordinateur équipé d'un bus PCI, elle n'a été testée que sur PowerMac.

B.1.1 Schéma général

La figure B.1 représente un schéma de principe de cette carte. Celle-ci communique avec l'extérieur au moyen des bus *Input* et *Util* et avec l'unité centrale à travers le bus PCI.

Les signaux d'entrée arrivent directement par le bus *Input*, de 18 bits de large, sur une mémoire tampon de 8 ou 32 kmots. Ils transitent ensuite par le registre d'entrée avant d'être envoyés à l'unité centrale via le bus PCI, qui est un bus 32 bits.

Les signaux de sortie constituent le bus *Util*, de 9 bits de large. Les lignes de ce bus affichent en permanence le contenu du registre de sortie qui est chargé directement avec les mots en provenance de l'unité centrale.

Il existe un registre *flags*, accessible en lecture et en écriture depuis l'unité centrale. Il permet de contrôler le fonctionnement de la carte. La carte dispose aussi d'un registre *test* qui ne sera pas décrit dans ce document.

Pour connaître la profondeur de votre tampon, regardez la référence des deux circuits intégrés identiques qui sont en haut de la carte : ces circuits constituent la mémoire tampon. Si la référence est *CY7C460*, vous avez un tampon de 8 kmots ; si c'est *CY7C464*, il s'agit d'un tampon de 32 kmots.

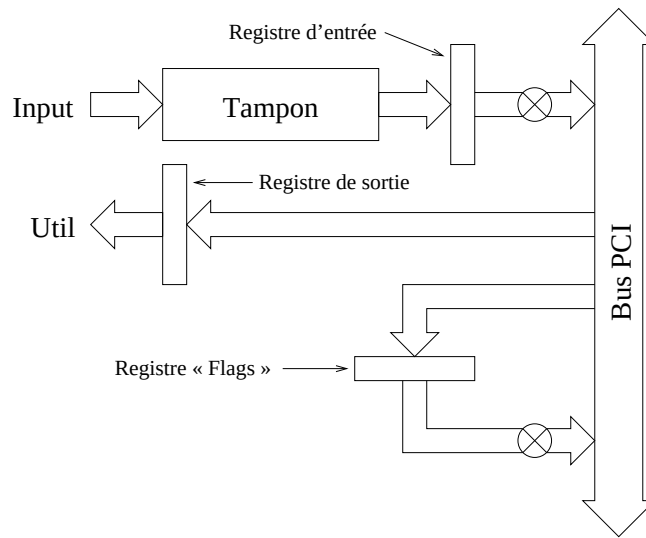


FIG. B.1 – Schéma de principe de la carte PCI.

B.1.2 Aspect physique

La figure B.2 représente la carte équipée de ses différents circuits.

La carte est construite autour d'un circuit logique programmable (PLD). Elle possède un connecteur PCI pour s'interfacer avec l'ordinateur, ainsi qu'un connecteur plat à 34 broches pour communiquer avec l'extérieur. Elle doit être équipée de trois cavaliers qui servent à mettre à la masse les lignes *test* (lignes inutilisées).

La carte peut aussi être équipée de deux mémoires tampon de type FIFO. Il y a à ce niveau trois possibilités :

pas de circuits FIFO : dans ce cas la carte ne peut être utilisée qu'en écriture (lignes *Util*) ;

circuits CY7C460 : ils constituent ensemble un tampon de 8 kmots de profondeur, les mots ayant une largeur de 18 bits ;

circuits CY7C464 : ils constituent un tampon de 32 kmots.

Il faut aussi équiper la carte d'un circuit mémoire PROM qui contient le programme configurant le PLD. Ce programme existe, bien entendu, en plusieurs versions plus ou moins boguées. Comme les PROM coûtent de l'argent et ne sont pas reprogrammables, vous devrez utiliser une version boguée dans la mesure où le bogue en question ne vous gêne pas dans l'utilisation que vous ferez de la carte. Les différentes versions disponibles sont :

FIFO : cette version contient un bogue entraînant parfois des erreurs en lecture, en revanche elle est tout à fait adaptée à une utilisation en écriture seule ;

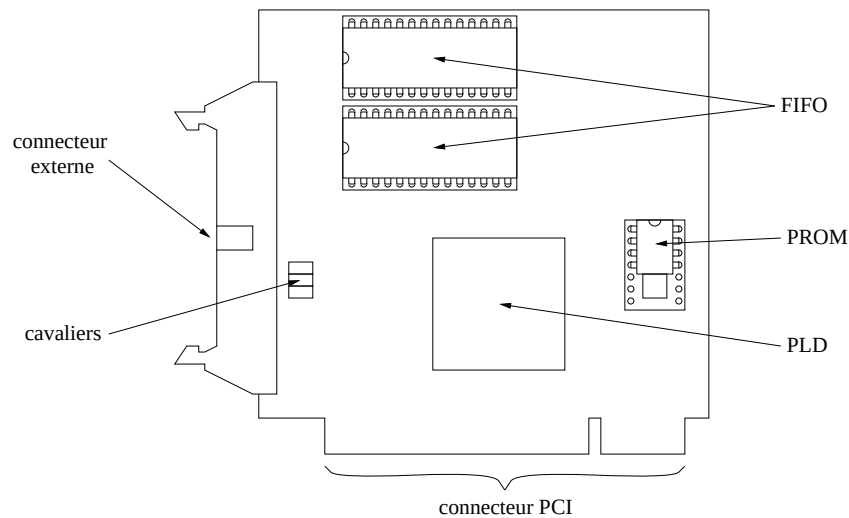


FIG. B.2 – Schéma physique de la carte PCI.

FIFC : celle-ci contient un bogue l'empêchant de cohabiter avec d'autres cartes dans le bus, elle vous conviendra si vous n'insérez qu'une seule carte PCI dans votre Mac ;

FIFD : aucun bogue connu.

La version du programme est écrite sur un autocollant fixé sur le circuit PROM.

B.1.3 Connecteur externe

C'est un connecteur plat 34 broches. Il est composé d'un bus d'entrée *Input* de 18 bits, d'un bus de sortie *Util* de 9 bits, d'un bus d'entrée *Test* de 3 bits qui est inutilisé et doit être mis à la masse avec des cavaliers, ainsi que des deux lignes de contrôle suivantes :

nEcriture est activée pour signaler la présence de données valides dans le bus *Util*. Ce signal est en logique négative.

nLecture ordonne au tampon de lire une donnée dans le bus *Input*. Ce signal est en logique négative. L'entrée du tampon étant du type *latch*, les données sont effectivement prises en compte lorsque ce signal est désactivé.

La figure B.3 montre la numérotation des broches du connecteur externe. Cette figure représente le connecteur de la carte tel que vu de l'arrière de l'ordinateur. Sur cette figure la carte mère est à droite.

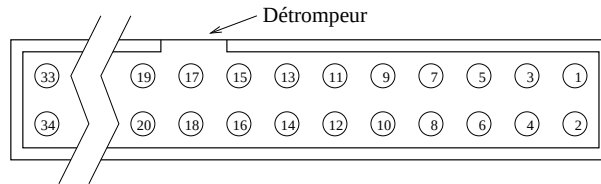


FIG. B.3 – Numérotation des broches du connecteur. La figure représente le connecteur de la carte tel que vu de l'extérieur.

L'affectation des broches est la suivante :

1 : Util1 ;	2 : Util0 ;	3 : Util3 ;	4 : Util2 ;
5 : Util5 ;	6 : Util4 ;	7 : Util7 ;	8 : Util6 ;
9 : Util8 ;	10 : Test0 ;	11 : nEcriture ;	12 : Test1 ;
13 : nLecture ;	14 : Test2 ;	15 : Input16 ;	16 : Input17 ;
17 : GND (masse) ;	18 : GND (masse) ;		
19 : Input14 ;	20 : Input15 ;	21 : Input12 ;	22 : Input13 ;
23 : Input10 ;	24 : Input11 ;	25 : Input8 ;	26 : Input9 ;
27 : Input6 ;	28 : Input7 ;	29 : Input4 ;	30 : Input5 ;
31 : Input2 ;	32 : Input3 ;	33 : Input0 ;	34 : Input1.

B.1.4 Modes de fonctionnement

Le registre d'entrée est un registre de 32 bits qui est chargé avec les mots de 18 bits en provenance du tampon. Il peut être chargé suivant deux modes différents : le mode *non multiplexé* et le mode *multiplexé*. Dans les deux modes, les mots en provenance du tampon sont chargés dans le registre d'entrée dès que cela est possible.

Mode non multiplexé

En mode non multiplexé, les 18 bits de poids faible du registre d'entrée sont chargés avec le mot en provenance du tampon. Les bits suivants contiennent le registre *flags*. On numérote toujours les bits d'un mot en commençant par celui de poids faible qui est numéroté zéro. Avec cette convention, la répartition des bits est la suivante :

bits 0 à 17 :	mot en provenance du tampon ;
bit 18 :	1 ;
bits 19 à 25 :	registre <i>flags</i> ;
bits 26 à 31 :	0.

En aucun cas (à cause du bit 18) le registre d'entrée n'est chargé uniquement

avec des zéros. Ceci permet de contrôler, lorsqu'on lit ce registre, qu'il contient bien des données pertinentes. En effet, si le registre d'entrée n'a pu être chargé (le tampon est vide), une tentative de lecture sur ce registre renvoie zéro.

Mode multipléxé

C'est un mode qui permet, au prix de certaines contraintes sur les mots lus, de multiplier par deux le débit de lecture du tampon en regroupant les mots deux par deux dans le registre d'entrée. Dans ce mode, les deux bits de poids fort des mots en provenance du tampon sont toujours ignorés. Ainsi, deux mots consécutifs du tampon sont chargés respectivement dans les 16 bits de poids fort du registre d'entrée puis dans les 16 bits de poids faible.

Outre le fait que les deux bits de poids fort des mots du tampon sont considérés comme dépourvus de sens, une hypothèse supplémentaire est faite sur les 16 bits de poids faible : les valeurs 0x0000 et 0x0001 sont considérées comme illégales. Lors du chargement du registre d'entrée, les mots valant 0x0000 sont transformés en 0x0001. Ainsi, ce registre n'est jamais chargé avec une valeur nulle.

Lors d'une tentative de lecture du registre d'entrée, si celui-ci n'a pas pu être complètement chargé (il ne contient pas deux mots), la lecture renvoie zéro.

B.1.5 Registre *flags*

Le registre *flags* est composé des bits suivants :

bit 19 :	Fifo_full ;
bit 20 :	nFifoef ;
bit 21 :	nFifo hf ;
bit 22 :	nFifo ff ;
bit 23 :	Plein0 ;
bit 24 :	Plein1 ;
bit 25 :	Multiplex.

Tous ces bits sont accessibles en lecture, seul le dernier (*Multiplex*) est accessible en écriture. À la lecture, les bits qui ne sont pas dans cette liste renvoient zéro. À l'écriture, tous les bits sauf le bit 25 (*Multiplex*) sont ignorés.

Les bits 20 à 22 renseignent sur l'état de remplissage courant du tampon. Ces trois drapeaux sont en logique négative et indiquent respectivement que le tampon est vide, qu'il est au moins à moitié plein et qu'il est plein. Dans ce dernier cas, il y a probablement eu perte de données arrivant dans le bus *Input*.

Le bit 19 (*Fifo_full*) permet de savoir si le tampon a été plein depuis la dernière lecture du registre *flags* : c'est un drapeau qui est levé dès que le tampon est plein et qui reste levé quand celui-ci cesse d'être plein. Il est baissé automatiquement après chaque lecture du registre *flags*, ainsi qu'après réinitialisation du tampon.

Les bits 23 et 24 permettent de savoir s'il y a des données dans le registre d'entrée. En mode multipléxé, le drapeau *Plein0* indique qu'au moins un mot a été chargé dans ce registre et le drapeau *Plein1* indique que deux mots ont été chargés. En mode non multipléxé, le drapeau *Plein1* indique qu'un mot a été chargé, l'autre n'a pas de signification. Dans les deux modes, le drapeau *Plein1* indique que le registre d'entrée est disponible pour lecture.

Lorsque le registre *flags* est lu à travers le registre d'entrée (bits de poids fort en mode non multipléxé), le drapeau *Plein1* n'a pas de signification. En effet, si le registre d'entrée n'a pu être chargé, la lecture renvoie zéro. Dans le cas contraire, le drapeau *Plein1* est forcément levé au moment où on le lit.

Le bit 25 (*Multiplex*) est le seul accessible en écriture. Le chargement du registre d'entrée se fait en mode multipléxé lorsque ce drapeau est levé et en mode non multipléxé sinon.

B.2 Interface C

L'interface est fournie sous la forme du fichier *carte_pci.h* qu'il suffit d'inclure dans un source C pour pouvoir utiliser facilement la carte. Ce fichier définit les fonctions et les constantes décrites à continuation. Ces fonctions sont en fait des macros (`#define`). Toutes sauf la dernière (`fast_read()`) ont été écrites de sorte qu'elles puissent être utilisées exactement comme des fonctions C. Les mots de 32 bits qui circulent à travers le bus PCI seront toujours considérés dans un programme C comme étant de type `unsigned long`.

B.2.1 Description des fonctions

```
unsigned long read_fifo(void)
```

renvoie le contenu du registre d'entrée, soit un ou deux mots extraits du tampon suivant que la carte travaille en mode multipléxé ou non. S'il n'y a aucune donnée à lire dans le registre d'entrée, cette fonction renvoie zéro.

```
unsigned long write_util(unsigned long)
```

écrit sur les lignes de sortie *Util* les neuf bits de poids faible de la valeur donnée en argument. Renvoie cette même valeur. Une impulsion d'écriture de 30 ns de durée est envoyée sur la ligne *nEcriture* quand les données sont stables.

```
unsigned long read_flags(void)
```

renvoie le contenu du registre *flags*. Chaque drapeau peut être testé en faisant le *et logique* de la valeur de retour avec l'une des constantes suivantes : `FLAG_FIFO_FULL`, `FLAG_NFIFOEF`, `FLAG_NFIFOHF`, `FLAG_NFIFOFF`, `FLAG_PLEINO`, `FLAG_PLEIN1`, `FLAG_MULTIPLEX`. L'appel de cette fonction abaisse le drapeau *Fifo_full*.

```
unsigned long write_flags(unsigned long)
```

écrit dans le registre *flags* la valeur donnée en argument. Renvoie cette même valeur. On utilise en général comme argument `FLAG_MULTIPLEX` si on veut un fonctionnement en mode multiplexé ou 0 (zéro) dans le cas contraire.

```
void reset_fifo(void)
```

réinitialise (vide) le tampon ainsi que le registre d'entrée. Cette fonction abaisse le drapeau *Fifo_full*.

```
unsigned long read_test(void)
```

renvoie la valeur contenue dans le registre *test*. Cette fonction n'est utile que pour tester le bon fonctionnement de la carte.

```
unsigned long write_test(unsigned long)
```

écrit dans le registre *test* la valeur donnée en argument et renvoie cette même valeur. Cette fonction n'est utile que pour tester le bon fonctionnement de la carte.

```
void wait_read(unsigned long *)
```

renvoie à l'adresse pointée par l'argument le contenu du registre d'entrée. Cette fonction attend qu'il y ait une donnée valide dans le tampon d'entrée, elle ne renvoie donc jamais zéro. **Attention** : s'il n'y a aucun circuit pour écrire dans le tampon, l'appel à cette fonction peut produire une boucle infinie et donc bloquer l'ordinateur. Syntaxiquement, cette macro et la suivante sont des instructions et non des expressions, on ne peut donc pas les utiliser comme expressions de contrôle d'une boucle *for* par exemple.

```
void fast_read(unsigned long)
```

cette macro est une alternative à la précédente qui ne respecte pas la syntaxe des fonctions C puisqu'elle modifie la variable donnée en paramètre. Elle est utile quand il faut faire une lecture très rapide. `fast_read(x)` est équivalent à `wait_read(&x)` à condition que l'évaluation de `x` ne produise pas d'effet de bord. Les utilisations de la forme `fast_read(tableau[i++])` sont donc à proscrire. En fait, cette macro est définie tout simplement par

```
#define fast_read(x)    {while (!((x) = read_fifo())) {}}
```

B.2.2 Utilisation de plusieurs cartes

Si vous devez utiliser plusieurs de ces cartes à la fois, vous devez taper

```
#define plusieurs_cartes
```

avant d'inclure `carte_pci.h`. Ceci a pour effet de créer la déclaration suivante

```
static int carte_pci = 1;
```

Vous pourrez ensuite modifier cette variable pour la faire pointer sur la carte avec laquelle vous voulez communiquer. Ne tapez pas cette définition si vous n'utilisez qu'une seule carte, elle a pour effet secondaire de ralentir légèrement les fonctions définies ci-dessus.

Les différentes cartes présentes sur le bus sont numérotées à partir de 1. La numérotation se fait de gauche à droite dans mon Mac mais ce ne sera pas forcément le cas sur le votre. Attention : un slot libre ne compte pas dans la numérotation, si vous avez une seule carte dans n'importe quel slot, elle sera numérotée 1. Évitez de faire pointer `carte_pci` sur une carte qui n'existe pas, autrement vous planterez le système à la première transaction.

B.2.3 Exemples de code

Supposons qu'on ait un appareil de mesure qui enregistre ses résultats dans le tampon. On veut lancer une acquisition de `NB_POINTS` de mesure qui seront rangés dans le tableau `donnees`.

Mode non multiplexé

On peut lire des mots de 18 bits, le tableau `donnees` sera du type `unsigned long`.

```
#include <stdio.h>
#include "carte_pci.h"
#define NB_POINTS 1024

unsigned long donnees[NB_POINTS];
int i; /* nombre de données lues */

write_flags(0); /* mode non multiplexé */
reset_fifo(); /* vide le tampon et abaisse le drapeau Fifo_full */
for (i = 0 ; i < NB_POINTS ; i++) {
    wait_read(&donnees[i]); /* lecture d'un mot */
    donnees[i] &= 0x3ffff; /* efface les bits 18 à 31 */
}
if (read_flags() & FLAG_FIFO_FULL) /* test du drapeau Fifo_full */
    fprintf(stderr, "Débordement du tampon\n");
```

Mode multiplexé

MacOS est un système gros-boutiste, ce qui correspond à la façon dont sont agencés les mots en mode multiplexé. Le plus pratique pour utiliser ce mode est d'écrire les données lues dans un tableau de type `unsigned long` et de les lire à l'aide d'un pointeur de type `unsigned short` pointant sur la même adresse.

```
#include <stdio.h>
#include "carte_pci.h"
#define NB_POINTS 1024 /* doit être pair */

unsigned long lecture[NB_POINTS / 2];
unsigned short *donnees = (unsigned short *) lecture;
int i;

write_flags(FLAG_MULTIPLEX);
reset_fifo(); /* vide le tampon et abaisse le drapeau Fifo_full */
for (i = 0 ; i < NB_POINTS / 2 ; i++)
    fast_read(lecture[i]); /* lecture d'un mot */
if (read_flags() & FLAG_FIFO_FULL) /* test du drapeau Fifo_full */
    fprintf(stderr, "Débordement du tampon\n");
```

Ensuite on peut utiliser `donnees` comme un tableau de longueur `NB_POINTS` contenant les données lues avec le type `unsigned short`. La figure B.4 illustre le rôle respectif de `lecture` et de `donnees`.

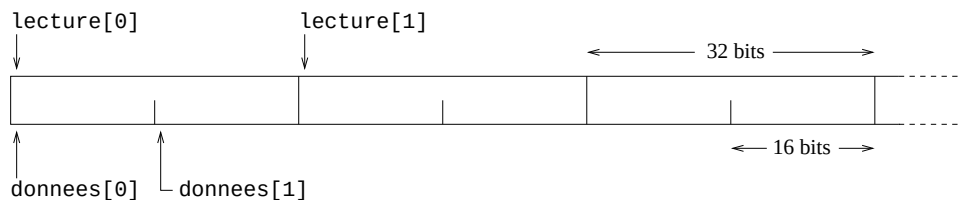


FIG. B.4 – Les pointeurs `lecture` et `donnees` pointent sur la même adresse mémoire. Ils y voient respectivement des mots de 32 et de 16 bits.